



Titre: Modèle bayésien d'agrégation des avis d'experts en exploitation
Title: d'équipements, application à l'optimisation de la disponibilité

Auteur: Bannour Souilah
Author:

Date: 2014

Type: Mémoire ou thèse / Dissertation or Thesis

Référence: Souilah, B. (2014). Modèle bayésien d'agrégation des avis d'experts en exploitation d'équipements, application à l'optimisation de la disponibilité
Citation: [Mémoire de maîtrise, École Polytechnique de Montréal]. PolyPublie.
<https://publications.polymtl.ca/1479/>

 **Document en libre accès dans PolyPublie**
Open Access document in PolyPublie

URL de PolyPublie: <https://publications.polymtl.ca/1479/>
PolyPublie URL:

Directeurs de recherche: Mohamed-Salah Ouali
Advisors:

Programme: Génie industriel
Program:

UNIVERSITÉ DE MONTRÉAL

MODÈLE BAYÉSIEN D'AGRÉGATION DES AVIS D'EXPERTS EN EXPLOITATION
D'ÉQUIPEMENTS, APPLICATION À L'OPTIMISATION DE LA DISPONIBILITÉ

BANNOUR SOUILAH
DÉPARTEMENT DE MATHÉMATIQUES ET DE GÉNIE INDUSTRIEL
ÉCOLE POLYTECHNIQUE DE MONTRÉAL

MÉMOIRE PRÉSENTÉ EN VUE DE L'OBTENTION
DU DIPLÔME DE MAÎTRISE ÈS SCIENCES APPLIQUÉES
(GÉNIE INDUSTRIEL)
AOÛT 2014

UNIVERSITÉ DE MONTRÉAL

ÉCOLE POLYTECHNIQUE DE MONTRÉAL

Ce mémoire intitulé :

MODÈLE BAYÉSIEN D'AGRÉGATION DES AVIS D'EXPERTS EN EXPLOITATION
D'ÉQUIPEMENTS, APPLICATION À L'OPTIMISATION DE LA DISPONIBILITÉ

présenté par : SOUILAH Bannour

en vue de l'obtention du diplôme de : Maîtrise ès sciences appliquées

a été dûment accepté par le jury d'examen constitué de :

Mme YACOUT Soumaya, D.Sc., présidente

M. OUALI Mohamed-Salah, Doct., membre et directeur de recherche

M. ADJENGUE Luc-Désiré, Ph.D., membre

DÉDICACE

*À mes parents,
je vous aime...*

REMERCIEMENTS

Je tiens pour commencer à remercier mon directeur de recherche, Monsieur Mohamed-Salah Ouali, pour son encadrement et sa disponibilité tout au long de ma recherche ainsi qu'au cours de la rédaction de ce mémoire.

Je remercie également les membres de jury, Madame Soumaya Yacout et Monsieur Luc-Désiré Adjengue, qui m'ont honoré en faisant partie du jury de soutenance de mon mémoire.

Je remercie mon ami Oussama Marzouk qui a été pour moi une référence dans l'élaboration de tous les modèles de ce travail. Merci Chef!

Enfin et tout particulièrement, je pense énormément à mes parents, Ahmed et Moufida, pour avoir toujours cru en moi et m'avoir supporté tout au long de mon séjour à Montréal, malgré la distance et les moments difficiles.

RÉSUMÉ

Le travail proposé dans ce mémoire traite le problème de la prise en compte des avis d'experts en exploitation des équipements industriels dans les modèles statistiques de fiabilité, de maintenance et de disponibilité. Les avis d'experts sont formulés sous forme de distributions a priori sur des paramètres inconnus de ces derniers modèles. L'inférence bayésienne est proposée comme une méthode de construction des distributions a posteriori de ces mêmes paramètres en tenant compte des données de vraisemblance collectées lors de l'exploitation d'un équipement. Ainsi, trois modèles sont développés et résolus à l'aide des méthodes de simulation de chaînes de Markov de type MCMC. Le premier modèle concerne l'agrégation des avis d'experts issus d'un même domaine d'expertise formulés sur un seul paramètre. Le second propose un modèle d'agrégation des avis formulés sur deux paramètres, les experts proviennent de deux domaines d'expertises différentes. Ces deux premiers modèles sont validés à l'aide de trois critères statistiques : le DIC (Deviance information criterion), le p-value bayésien et le test non paramétrique de mesure de tendance des données. Le troisième modèle concerne l'optimisation de la disponibilité d'un équipement avec une actualisation du taux de défaillance et du taux de réparation sur plusieurs périodes. Les trois modèles ainsi que les algorithmes de résolution sont programmés sous Matlab. Des données de simulation ont été exploitées afin de tester la logique de modélisation et d'analyser la performance des modèles proposés.

ABSTRACT

The work proposed in this document addresses the problem of combining different opinions of experts about unknown parameters within the reliability, maintenance and availability models. Expert opinions are characterised using a prior distribution functions of these unknown parameters. The Bayesian inference is proposed as a modeling approach to establish the posterior distribution functions of the same unknown parameters with respect to the likelihood data collected during the operation time of the equipment. Thus, three models are developed and solved using Markov Chains Monte Carlo (MCMC). The first model concerns the aggregation of expert opinions from the same area of expertise made on a single unknown parameter. The second provides an aggregation of opinions on two unknown parameters. The experts belong to two different areas of expertise. These first two models are validated using three statistical criteria: DIC (Deviance information criterion), the bayesian p-value and the nonparametric test of Fisher which measures the data trend. The third model concerns the optimization of the availability of equipment with a discount of the failure and repair rates over several periods of time. These three models as well as the solving algorithms are programmed using Matlab software. Simulation data were used to test the logic of modeling and to analyze the performance of the proposed models.

TABLE DES MATIÈRES

DÉDICACE	iii
REMERCIEMENTS	iv
RÉSUMÉ	v
ABSTRACT	vi
TABLE DES MATIÈRES	vii
LISTE DES TABLEAUX	ix
LISTE DES FIGURES	x
LISTE DES ABRÉVIATIONS	xi
LISTE DES ANNEXES	xii
CHAPITRE 1 INTRODUCTION	1
1.1 Problématique	1
1.2 Objectif du travail	3
1.3 Organisation du mémoire	4
CHAPITRE 2 APPROCHES DE MODÉLISATION ET REVUE DE LITTÉRATURE	5
2.1 Approches de modélisation	5
2.1.1 Fiabilité d'un équipement	5
2.1.2 Maintenabilité d'un équipement	7
2.1.3 La disponibilité d'un équipement	9
2.1.4 Classifications des interventions de maintenance	10
2.2 Méthodes classiques d'estimation paramétrique	12
2.3 Méthodes d'estimation bayésienne	12
2.3.1 Algorithme de Metropolis-Hastings	14
2.3.2 Algorithme de Gibbs	14
2.3.3 Qualité d'ajustement d'un modèle bayésien	15
2.4 Revue des modèles de combinaisons des avis d'experts	16
2.4.1 Méthode du mélange	17

2.4.2	Méthode de la moyenne	18
2.5	Stratégies d'optimisation de la disponibilité	19
CHAPITRE 3 MODÈLE D'AGRÉGATION BAYÉSIENNE		22
3.1	Mise en contexte	22
3.2	Modélisation de la vraisemblance et des avis d'experts	23
3.2.1	Calcul de la fonction de vraisemblance	23
3.2.2	Modélisation des avis d'experts	24
3.2.3	Algorithme de résolution	26
3.2.4	Illustration du modèle d'agrégation	29
3.3	Comparaison des modèles de combinaison d'avis d'experts	32
3.4	Validation du modèle d'agrégation bayésienne avec la tendance des données	33
3.4.1	Analyse de la tendance des données de vraisemblance	34
3.4.2	Analyse des avis d'experts selon la tendance des données	35
CHAPITRE 4 OPTIMISATION DE LA DISPONIBILITÉ D'UN ÉQUIPEMENT		38
4.1	Stratégie d'optimisation de la disponibilité	38
4.1.1	Description de la stratégie	38
4.1.2	Stratégie de disponibilité prenant en compte les avis d'experts	40
4.2	Modèle de vraisemblance associé à la disponibilité	44
4.3	Distributions a priori sur λ et μ	46
4.4	Distributions a posteriori sur λ et μ	48
4.5	Exemple d'optimisation de la disponibilité	50
CHAPITRE 5 CONCLUSION ET PERSPECTIVES		56
RÉFÉRENCES		59
ANNEXES		62

LISTE DES TABLEAUX

Tableau 2.1	Relations entre les différentes fonctions de fiabilité	6
Tableau 2.2	Fonctions de fiabilité selon les modèles exponentiel et de Weibull	7
Tableau 2.3	Relations entre les fonctions de maintenabilité	8
Tableau 2.4	Fonctions de maintenabilité selon les modèles exponentiel et log-normal	9
Tableau 2.5	Quelques distributions usuelles conjuguées	13
Tableau 3.1	Durées de bon fonctionnement exprimées en semaine	29
Tableau 3.2	Avis des 2 experts et paramètres des distributions a priori	29
Tableau 3.3	Statistiques des distributions a postérieur des 3 modèles	31
Tableau 3.4	Fonction de masse a postérieur de la variable I	31
Tableau 3.5	Modélisation des 2 avis d'experts	33
Tableau 3.6	Résultats a postérieur de la comparaison	33
Tableau 3.7	Probabilités a postérieur du modèle d'agrégation bayésienne	33
Tableau 3.8	Résultats du test de <i>Fisher</i> pour un seuil de signification $\alpha = 0,05$. .	35
Tableau 3.9	Modélisation des 3 avis d'experts	36
Tableau 3.10	Validation de la tendance des données avec les probabilités a postérieur	36
Tableau 4.1	Avis des experts en fiabilité et paramètres des distributions a priori . .	51
Tableau 4.2	Avis des experts en maintenance et paramètres des distributions a priori	51
Tableau 4.3	Données de vraisemblance	52
Tableau 4.4	Statistiques a postérieur de la première année sur les paramètres λ et μ	52
Tableau 4.5	Fonction de masse a postérieur de la variable J	52
Tableau 4.6	Fonction de masse a postérieur de la variable V	53
Tableau 4.7	Maximisation de la disponibilité de l'équipement	54

LISTE DES FIGURES

Figure 2.1	Influence de la maintenance parfaite, minimale et imparfaite sur le taux de défaillance	11
Figure 3.1	Distributions a priori des deux experts et distribution a post��riori du mod��le 3	31
Figure 3.2	Distributions a priori des deux experts et distributions a post��riori des 3 mod��les	34
Figure 3.3	Probabilit��s a post��riori des 3 experts	37
Figure 4.1	Maximisation de la disponibilit��	39
Figure 4.2	Strat��gie de disponibilit�� et introduction d'avis d'experts	41
Figure 4.3	Cycle de fonctionnement de l'��quipement	44
Figure 4.4	D��roulement de l'��chantillonnage pour les param��tres de la premi��re ann��e	53
Figure 4.5	Allure du taux de d��faillance	55

LISTE DES ABRÉVIATIONS

MCMC	Markov Chain Monte Carlo
DIC	Deviance Information Criterion
DM	Preneur de décision « Decision maker »

LISTE DES ANNEXES

Annexe A	PROGRAMME MATLAB 1 : AGRÉGATION DES AVIS D'EXPERTS	62
Annexe B	PROGRAMME MATLAB 2 : OPTIMISATION DE LA DISPONIBILITÉ	67

CHAPITRE 1

INTRODUCTION

1.1 Problématique

Dans une étude de fiabilité d'un équipement en exploitation, les durées de vie enregistrées, sont généralement modélisées par des lois de probabilité paramétriques telle que la loi exponentielle, de Weibull ou autre. Ces lois de probabilité caractérisent le comportement passé de l'équipement et permettent de prédire son comportement futur avec un certain degré de confiance, et ce en fonction de la nature des données de durées de vie disponibles. Or, l'estimation des paramètres de ces lois constitue la tâche la plus délicate vu la nature des données à partir desquelles ces paramètres sont estimés. En effet, un nombre limité de données ne permet pas une bonne estimation des paramètres à travers les méthodes classiques d'estimation paramétriques telles que les méthodes des moindres carrés et du maximum de vraisemblance. De plus, des données incomplètes ou encore censurées requièrent un traitement particulier afin de procéder à l'estimation adéquate des paramètres inconnus.

Généralement, lorsque les données de défaillance ne sont pas disponibles en quantité suffisante, l'estimation des paramètres inconnus par les méthodes classiques sera statistiquement biaisée. Ainsi, les fonctions statistiques caractérisant le comportement de l'équipement telles que la fiabilité, le risque de défaillance et le taux de défaillance ne sont pas représentatives. Par conséquent, elles ne peuvent pas être utilisées pour optimiser une stratégie de maintenance telle que l'optimisation des coûts de maintenance d'un équipement ou maximiser sa disponibilité. Dans ce contexte, les gestionnaires de maintenance sont généralement amenés à considérer les lois de défaillance fournies par le constructeur de l'équipement pour planifier les activités de maintenance. Ces lois sont établies sur la base d'essais de laboratoire ou de retour d'expérience du service après-vente. Elles considèrent un usage dit « normal » de l'équipement.

Or, l'utilisation d'un équipement ne se fait pas dans les conditions spécifiées par le constructeur. En effet, l'équipement peut être modifié pour incorporer des fonctions non prévues par le constructeur. De plus, les actions de maintenance que subit l'équipement peuvent changer son comportement futur si elles ne sont réalisées conformément aux recommandations du constructeur (maintenance imparfaite). De ce fait l'usage des modèles de fiabilité et de maintenabilité proposés par le constructeur sera désuet. Les modèles ne refléteront plus

le comportement réel de l'équipement. Or, durant l'exploitation d'un équipement, plusieurs informations sont collectées et analysées par les experts en exploitation d'équipements dans le but de prendre des décisions de maintenance telles que les procédures d'inspection à mettre en place, le remplacement partiel ou total des composants de l'équipement. Ces informations ne sont pas considérées par les modèles statistiques alors que réellement elles ont été mises à contribution par les experts pour prendre des décisions de maintenance qui affectent le comportement de l'équipement. Dans ce cas de figure, l'intégration de ces informations sous forme d'avis d'experts sur les paramètres inconnus des modèles statistiques peut aider à actualiser les modèles statistiques et ainsi à mieux prédire le comportement futur de l'équipement.

De plus, l'estimation paramétrique des paramètres inconnus à l'aide des méthodes classiques d'estimation ne prend pas en compte des avis d'experts. Elle nécessite des méthodes plus avancées afin de répondre à la question principale suivante : comment prédire le comportement futur d'un équipement et optimiser la stratégie de maintenance d'un équipement au fur et à mesure que de nouvelles données de défaillance et des avis d'experts sont collectées. Cette question principale peut se décliner en quatre questions secondaires suivantes :

1. Comment renseigner un modèle statistique par un avis d'expert exprimé sur un ou plusieurs paramètres inconnus du modèle ?
2. Comment combiner plusieurs avis d'experts exprimés sur un même paramètre inconnu ?
3. Comment évaluer la contribution (le poids) de chaque avis d'expert dans l'estimation du paramètre inconnu ?
4. Comment valider les modèles statistiques prenant en compte plusieurs avis d'experts ?
5. Comment optimiser la stratégie de maintenance d'un équipement basé sur une actualisation des modèles statistiques ?

Pour répondre aux trois premières questions (1, 2 et 3), l'une des méthodes proposées dans la littérature est l'inférence bayésienne. Cette méthode probabiliste est basée sur le théorème de Bayes tel que sera expliqué dans la revue de littérature. En fait, l'inférence bayésienne permettra de conjuguer des avis d'experts exprimés sur un ou plusieurs paramètres inconnus avec les données de défaillance déjà enregistrées. Les paramètres inconnus seront ainsi déterminés en considérant les données de défaillance et les avis d'experts. Des développements mathématiques seront nécessaires dépendamment des modèles choisis pour caractériser les données de défaillance et les avis d'experts. Nous les considérons ultérieurement.

Pour adresser la question 4), comme toute approche statistique, nous aurons besoin de tests statistiques pour valider les hypothèses de modélisation. Nous devrons utiliser des tests qui nous permettront de comparer statistiquement la performance des experts entre eux,

c'est-à-dire celui ou ceux qui se rapproche le plus du modèle de défaillance. Aussi, nous aurons besoin de tests de validation des avis d'experts par rapport à la tendance réelle des données de défaillance enregistrées.

En ce qui concerne la dernière question 5), Nous considérons un équipement multi-composant soumis à une stratégie de réparation minimale (remise en état aussi bon que vieux) en cas de défaillance de l'équipement et à une réparation majeure (remise en état aussi bon que neuf) périodiquement. Cette stratégie sera appliquée à l'optimisation de la disponibilité de l'équipement en tenant compte des modèles statistiques combinant les données de défaillance et les avis d'experts en exploitation d'équipements. Des développements mathématiques seront nécessaires afin de permettre cette optimisation.

1.2 Objectif du travail

L'objectif principal de recherche est de proposer un modèle d'agrégation des avis d'experts provenant de plusieurs domaines d'expertises, en particulier les domaines de fiabilité et de maintenabilité d'équipements industriels en exploitation, basé sur l'inférence bayésienne et appliqué à l'optimisation de la disponibilité d'un équipement multi-composant. Cet objectif se décompose en quatre sous objectifs :

1. Développer un modèle paramétrique d'agrégation des avis d'experts, peu importe leur nombre et leur degré de connaissance, formulés à propos d'un paramètre inconnu du modèle de fiabilité (taux de défaillance) ou de maintenabilité (taux de réparation) d'un équipement.
2. Comparer plusieurs modèles d'agrégation des avis d'experts sur la base de critère de performance statistique tels que le DIC (Deviance Information Criterion) et/ou le p-value bayésien.
3. Valider les modèles statistiques a posteriori sur les paramètres d'intérêt en exploitant la tendance des données et les meilleurs avis d'experts retenus par le DIC ou le p-value bayésien.
4. Appliquer le modèle d'agrégation des avis d'experts à l'optimisation de la disponibilité d'un équipement au cours du temps. Cette disponibilité s'exprime en fonction de deux paramètres inconnus : le taux de défaillance et le taux de réparation. Ce sous-objectif nécessite l'actualisation des paramètres inconnus du modèle au fur et à mesure que l'équipement est exploité.

1.3 Organisation du mémoire

Le mémoire est organisé en quatre chapitres. Le deuxième chapitre présente une revue de littérature sur les approches de modélisation de la fiabilité, la maintenabilité, la disponibilité et de classification des interventions en maintenance. Une revue de la littérature sur les méthodes d'estimation classique et bayésienne sera également élaborée. Ce chapitre traite aussi une revue de littérature des méthodes de combinaison des avis d'experts et les moyens de modélisation mathématique. Il présente également une revue sur les stratégies de maintenance permettant la maximisation de la disponibilité d'équipement.

Le troisième chapitre présente un modèle d'agrégation des avis d'experts exprimés sur un paramètre inconnu. Ce modèle exploite les méthodes d'échantillonnage par simulation de type Markov Chain Monte Carlo (MCMC) afin de résoudre le modèle proposé d'agrégation des avis d'experts. Les résultats de ce modèle seront comparés à deux types de modèles de la littérature, le modèle du mélange et le modèle de la moyenne. Il sera ensuite validé sur la base du DIC et le p-value bayésien. Une validation des avis d'experts exprimés sur plusieurs périodes sera faite à l'aide du test non-paramétrique de Fisher appliqué à tendance des données.

Dans le quatrième chapitre, le modèle d'agrégation des avis d'experts sera appliqué à l'optimisation de la disponibilité de l'équipement. Le modèle de disponibilité tiendra compte de deux paramètres inconnus : le taux de défaillance et le taux de réparation. Une adaptation du modèle de remplacement périodique avec réparation minimale en cas de défaillance de Barlow et Proschan (1996) sera proposée afin de prendre en compte les avis d'experts exprimés sur les paramètres inconnus du modèle d'optimisation de la disponibilité.

Le cinquième chapitre présente la conclusion du travail et les perspectives de recherche à partir des résultats obtenus.

CHAPITRE 2

APPROCHES DE MODÉLISATION ET REVUE DE LITTÉRATURE

Dans ce chapitre, nous présentons les approches de modélisation de la fiabilité, la maintenabilité, la disponibilité et de classification des interventions en maintenance. Nous présentons par la suite les méthodes classiques d'estimation des paramètres des modèles de fiabilité, de maintenabilité et de disponibilité. Ensuite, nous présentons une revue sur les méthodes d'estimation bayésienne lorsque les paramètres des modèles sont inconnus. Les critères statistiques qui mesurent la qualité d'estimation des paramètres (qualité d'ajustement des modèles) sont décrits et expliqués. De même, une revue de littérature sur les méthodes de combinaison des avis d'experts sera élaborée. Enfin, nous proposons une revue des stratégies de maintenance permettant de maximiser la disponibilité d'un équipement.

2.1 Approches de modélisation

2.1.1 Fiabilité d'un équipement

Selon la norme AFNOR X06501, « *la fiabilité est l'aptitude d'un système à accomplir une fonction requise dans des conditions d'utilisation et pour une période de temps bien déterminées* ». En pratique, la fiabilité est directement liée à la durée de vie d'un équipement. Cette durée de vie est une variable aléatoire, notée T , caractérisant le passage aléatoire d'un équipement d'un état de fonctionnement à un état de défaillance selon une loi de probabilité qui peut être connue ou inconnue. La durée de vie T s'exprime en unités de temps, en nombre de cycles ou en unités produites par l'équipement.

Le modèle permettant de calculer la probabilité de survie d'un équipement est appelé « Fonction de fiabilité ». Elle s'exprime en fonction de la durée de la mission t , une réalisation de la variable aléatoire T , par : $R_T(t) = P(T > t)$. C'est la probabilité de survie au-delà de la durée t . Comme l'équipement ne peut être qu'en état de fonctionnement ou en état de défaillance, la probabilité que l'équipement tombe en défaillance avant t est donnée par le « Risque de défaillance », noté $F_T(t) = 1 - R_T(t) = P(T \leq t)$.

Connaissant la fonction de densité de défaillance $f_T(t)$, tel que $f_T(t)dt = P(t < T < t+dt)$ calcule la probabilité instantanée de défaillance d'un équipement entre t et $t+dt$, la fonction de fiabilité est donnée par : $R_T(t) = \int_t^{+\infty} f_T(v)dv$.

D'une manière plus pratique, $R_T(t)$ peut s'écrire aussi en fonction du taux de défaillance $\lambda_T(t)$. Le taux de défaillance calcule la probabilité de défaillance d'un équipement par unité de temps pour une mission de durée t . Il s'agit d'une vitesse d'arrivée des défaillances, exprimée en nombre de défaillances par unité de temps. La quantité $\lambda_T(t)dt$ donne la probabilité de défaillance d'un équipement entre t et $t + dt$ sachant que l'équipement a survécu jusqu'à t . Le tableau 2.1 résume les différentes relations entre les fonctions de fiabilité (La variable T a été enlevée par souci de simplification).

Tableau 2.1 Relations entre les différentes fonctions de fiabilité

Fonctions de fiabilité	$f(t)$	$R(t)$	$F(t)$	$\lambda(t)$
$f(t)$		$\int_t^{+\infty} f(v)dv$	$\int_0^t f(v)dv$	$\frac{f(t)}{\int_t^{+\infty} f(v)dv}$
$R(t)$	$-\frac{dR(t)}{dt}$		$1 - R(t)$	$\frac{-\frac{dR(t)}{dt}}{R(t)}$
$F(t)$	$\frac{dF(t)}{dt}$	$1 - F(t)$		$\frac{\frac{dF(t)}{dt}}{1-F(t)}$
$\lambda(t)$	$\lambda(t) \cdot e^{-\int_0^t \lambda(v)dv}$	$e^{-\int_0^t \lambda(v)dv}$	$1 - e^{-\int_0^t \lambda(v)dv}$	

Usuellement, deux grandes approches d'estimation des fonctions de fiabilité sont répertoriées dans la littérature : les approches non paramétriques et les approches paramétriques. Les approches non paramétriques sont basées sur le rang de la défaillance dans un échantillon donnée. Les méthodes de Kaplan-Meier (Kaplan et Meier (1958)) et de Johnson (Herd (1960)) sont les plus utilisées en fiabilité. Ces méthodes donnent une estimation de la probabilité de défaillance d'un équipement pour une mission donnée de durée t . Aucune hypothèse statistique n'est émise ou vérifiée avec ces méthodes. Les approches paramétriques prennent comme hypothèse que les données des durées de vie suivent une loi de probabilité connue. Cette loi est caractérisée par un ou plusieurs paramètres. Les données collectées sont inférées statistiquement afin de calculer des estimateurs des paramètres recherchés. Un test statistique est généralement appliqué pour accepter ou refuser le modèle (Elsayed (2012)).

Parmi les modèles paramétriques les plus rencontrés dans le domaine de l'ingénierie de la fiabilité, nous retrouvons le modèle exponentiel et le modèle de Weibull. Dans le livre de Elsayed (2012), ce taux de défaillance est considéré comme constant faisant référence à ce que l'occurrence des défaillances suit un processus de Poisson homogène de paramètre λ ($\lambda = \text{constante}$). Ce modèle est fréquemment utilisé vu sa simplicité. Il est généralement utilisé pour les équipements électroniques et est parfaitement adapté aux équipements dont les défaillances soudaines surviennent indépendamment de son âge. De plus, le modèle exponentiel est fréquemment utilisé pour les équipements multi-composants dont la défaillance est

causée par la défaillance de l'un des composants qui le constituent (Jardine et Tsang (2013)). Cependant, comme l'équipement vieillit au cours de son usage, le taux de défaillance varie au cours du temps faisant référence à un processus de Poisson non homogène.

Le second modèle est celui de Weibull, qui est aussi un modèle exponentiel. Il est plus flexible et permet de modéliser des modes de défaillance plus complexes (usure, fatigue, etc.) grâce à ses trois paramètres (β, η, γ) , où β , η et γ caractérisent respectivement les paramètres de forme, d'échelle et de position. Le taux de défaillance du modèle de Weibull est bien adapté pour décrire les trois phases de vie d'un équipement suivant son paramètre de forme β (Bacha *et al.* (1998)) :

- $\beta < 1$: C'est la phase de jeunesse. $\lambda_T(t)$ décroît avec le temps.
- $\beta = 1$: C'est la phase de la vie utile. $\lambda_T(t)$ est constant.
- $\beta > 1$: C'est la phase de la vie utile. $\lambda_T(t)$ augmente avec le temps.

Le tableau 2.2 résume les fonctions de fiabilité exprimées selon les modèles exponentiel et de Weibull.

Tableau 2.2 Fonctions de fiabilité selon les modèles exponentiel et de Weibull

Fonctions	$f(t)$	$R(t)$	$F(t)$	$\lambda(t)$
Modèle exponentiel	$\lambda e^{-\lambda t}$	$e^{-\lambda t}$	$1 - e^{-\lambda t}$	λ
Modèle de Weibull	$\frac{\beta}{\eta} \cdot \left(\frac{t-\gamma}{\eta}\right)^{\beta-1} \cdot e^{-\left(\frac{t-\gamma}{\eta}\right)^\beta}$	$e^{-\left(\frac{t-\gamma}{\eta}\right)^\beta}$	$1 - e^{-\left(\frac{t-\gamma}{\eta}\right)^\beta}$	$\frac{\beta}{\eta} \cdot \left(\frac{t-\gamma}{\eta}\right)^{\beta-1}$

Pour un modèle de fiabilité donné, il est possible de calculer la durée de vie moyenne comme étant l'espérance mathématique de la variable aléatoire T :

$$E[T] = \int_0^{+\infty} R(t)dt = MTBF$$

Où *MTBF* désigne « Mean Time Before Failure » pour les équipements non réparables et « Mean Time Between Failure » pour les équipements réparables. La traduction française de *MTBF* est « Moyenne de Temps de Bon Fonctionnement ».

2.1.2 Maintenabilité d'un équipement

Selon la norme AFNOR X60010, « la maintenabilité d'un équipement est son aptitude à être maintenu en un temps donné dans un état et dans des conditions déterminées ». Plus précisément, la maintenabilité est la probabilité que l'équipement soit réparé (remis en état de fonctionnement) en un temps t . Donc, la durée de réparation est considérée comme

une variable aléatoire. C'est à l'aide des méthodes de temps prédéterminés ou de l'historique des réparations que le modèle de maintenabilité pourrait-être établi. L'estimation de la probabilité qu'un équipement se trouve réparé avant un t donné se fait à l'aide de modèles paramétriques. Les principales distributions rencontrées dans la littérature sont les modèles exponentiel et log-normal (Procaccia *et al.* (2011)). Il existe une grande similarité entre les fonctions caractérisant la maintenabilité d'un équipement et ceux associées à la fiabilité.

Soit Y la variable aléatoire définissant la durée de réparation. On définit par $g_Y(t)$ la densité de réparation. $g_Y(t)dt$ calcule la probabilité que l'équipement soit remis en service entre les temps t et $t + dt$. La fonction de maintenabilité est donnée par $M_Y(t) = P(Y \leq t)$. C'est la probabilité que l'équipement soit réparé avant t . Connaissant l'expression de la fonction $M_Y(t)$, la densité du temps de réparation est donnée par $g_Y(t) = \frac{dM_Y(t)}{dt}$.

Par ailleurs, il est possible d'exprimer la fonction de maintenabilité en fonction du taux de réparation $\mu_Y(t)$. Ce taux est défini tel que $\mu_Y(t) \cdot dt$ est la probabilité conditionnelle de fin de réparation entre t et $t + dt$ sachant que l'équipement était non fonctionnel à l'instant t . Le taux de réparation est exprimé en nombre de réparations par unité de temps. Le tableau 2.3 résume les relations entre les différentes fonctions de maintenabilité (la variable Y a été enlevée par souci de simplification).

Tableau 2.3 Relations entre les fonctions de maintenabilité

Fonction de maintenabilité	$g(t)$	$M(t)$	$\mu(t)$
$g(t)$		$\int_0^t g(v)dv$	$\frac{g(t)}{1 - \int_0^t g(v)dv}$
$M(t)$	$\frac{dM(t)}{dt}$		$\frac{\frac{dM(t)}{dt}}{1 - M(t)}$
$\mu(t)$	$\mu(t) \cdot e^{-\int_0^t \mu(v)dv}$	$1 - e^{-\int_0^t \mu(v)dv}$	

Dans le cas de maintenabilité, le temps moyen de réparation $MTTR$ est défini comme étant l'espérance mathématique de la variable aléatoire Y .

$$MTTR = E[Y] = \int_0^{+\infty} t \cdot g(t)dt = \int_0^{+\infty} (1 - M(t))dt.$$

Comme dans le cas de la fiabilité, nous devons choisir des modèles permettant de modéliser la variable aléatoire Y . Le tableau 2.4 résume les fonctions de la maintenabilité pour les modèles exponentiel et log-normal les plus utilisés.

Tableau 2.4 Fonctions de maintenabilité selon les modèles exponentiel et log-normal

Fonctions de maintenabilité	$g(t)$	$M(t)$	$\mu(t)$
Modèle exponentiel	$\mu \cdot e^{-\mu t}$	$1 - e^{-\mu t}$	μ
Model Log-Normal	$\frac{1}{\sigma t \sqrt{2\pi}} \cdot e^{\left(-\frac{[\ln(x)-\mu]^2}{2\sigma^2}\right)}$	$\frac{1}{2} + \frac{1}{2} \operatorname{erf} \left[\frac{\ln(x)-\mu}{\sqrt{2}\sigma} \right]$	$\frac{\frac{dM(t)}{dt}}{1-M(t)}$

2.1.3 La disponibilité d'un équipement

Pour un équipement réparable en exploitation, il est souvent intéressant de savoir quelle est la probabilité que ce dernier soit disponible ou, en d'autres mots, quelle sera la proportion de temps durant laquelle cet équipement sera fonctionnel. La disponibilité $A(t)$ est définie comme étant la probabilité qu'à un instant donné t , l'équipement soit en opération, il ne soit ni défaillant ni en réparation. D'une façon analogue, nous définissons l'indisponibilité comme la probabilité que l'équipement ne soit pas disponible à un instant t donné.

Dans le cas d'un équipement, pouvant se trouver dans un des deux états : en fonctionnement ou en défaillance, dont les taux de défaillance λ et de réparation μ sont constants, il est relativement facile de montrer que la disponibilité instantanée, si l'équipement est fonctionnel à $t = 0$, s'écrit comme suit (Elsayed (2012)) :

$$A(t) = \frac{\mu}{\lambda + \mu} + \frac{\lambda}{\lambda + \mu} \cdot \exp(-(\lambda + \mu) \cdot t).$$

Si l'équipement se trouve dans l'état de défaillance à $t = 0$, la disponibilité instantanée s'écrit comme suit :

$$A(t) = \frac{\mu}{\lambda + \mu} \cdot (1 - \exp(-(\lambda + \mu) \cdot t)).$$

Il est important de mentionner que lorsque l'équipement possède des taux de défaillance et de réparation constants, sa disponibilité tend vers une limite indépendante de son état de départ (en fonctionnement ou en défaillance) :

$$\lim_{t \rightarrow +\infty} A(t) = \frac{\mu}{\lambda + \mu} = \frac{MTBF}{MTBF + MTTR}.$$

La disponibilité peut être aussi traitée comme une variable aléatoire. Dans le chapitre 4 de ce mémoire, nous allons définir une variable aléatoire caractérisant la disponibilité en fonction de deux autres variables aléatoires : le temps de bon fonctionnement et le temps de réparation. La fonction de densité de la variable caractérisant la disponibilité sera déterminée et par la suite utilisée dans l'estimation des paramètres inconnus.

2.1.4 Classifications des interventions de maintenance

Afin de comprendre les effets des interventions de maintenance, nous avons besoin d'introduire tout d'abord les différentes classifications des activités de maintenance. La norme AFNOR X60010 donne une première classification des interventions selon le niveau de complexité de la tâche. Cette complexité est caractérisée par le niveau de compétence requis de l'intervenant (opérateur de maintenance). Plus le niveau de complexité est élevé plus la compétence requise le sera et par conséquent plus le niveau de l'intervention le sera aussi (Ouali (2013)) :

- Niveau 1 : Ce niveau regroupe des activités de maintenance telles que les réglages simples, les nettoyages et les graissages. Ces activités ne nécessitent pas de démontage ou d'ouverture de l'équipement. Ce niveau nécessite un niveau de compétence bas et concerne des éléments facilement accessibles de l'équipement. Un opérateur de production (exploitant) peut assumer les tâches de ce niveau.
- Niveau 2 : Ce niveau regroupe des activités de maintenance telles que les remplacements par échange standard d'éléments ou de sous-ensembles. Un technicien habile peut assumer les tâches de ce niveau. Il s'agit des actions de maintenance préventive.
- Niveau 3 : Les actions de maintenance dans ce niveau nécessitent l'usage de procédures complexes en particulier le diagnostic des défaillances. Les actions de ce niveau sont telles que l'identification de pannes, la réparation par échange de composants fonctionnel et les réparations mécaniques mineures. Un technicien spécialisé peut assumer les tâches de ce niveau. Il s'agit des actions de maintenance correctives.
- Niveau 4 : Ce niveau inclut les actions de maintenance corrective ou préventive à l'exception des rénovations. Ce niveau nécessite la maîtrise d'une technique ou d'une technologie particulière. Un équipe encadrée par un ingénieur ou technicien spécialisé peut assumer les tâches de ce niveau.
- Niveau 5 : Ce niveau regroupe les actions de maintenance telles que les travaux de rénovation, de reconstitution ou de réparations importantes. Ces actions ont des procédures qui impliquent un savoir-faire. Elles font appel à des techniques spécialisées et doivent être assumées par une équipe polyvalente.

La deuxième classification des interventions de maintenance considère le degré de restauration d'un équipement. Wang et Pham (2006b) suggèrent la classification en maintenance parfaite, minimale et imparfaite :

- Maintenance parfaite : Dans ce cas, l'équipement est ramené à l'état neuf, « aussi bon

que neuf » (as good as new). Il s'agit d'un renouvellement de l'équipement. Après une maintenance parfaite, le taux de défaillance repart à son niveau le plus bas ($t = 0$).

- Maintenance minimale : L'équipement est ramené à son état d'opération sans toutefois affecter son taux de défaillance. Il s'agit d'une remise en état qualifiée de « aussi mauvais que vieux » (as bad as old). Cette maintenance est généralement considérée pour caractériser une intervention limitée à quelques composants d'un équipement. Bien que l'état des composants puisse être caractérisé par « aussi bon que neuf », l'état de tout l'équipement reste « aussi mauvais que vieux ». Cela dépend, bien entendu, du nombre et de l'importance des composants ayant subi une intervention de maintenance. Il en résulte que la réparation partielle (un ou plusieurs composants) de l'équipement n'affecte pas le taux de défaillance global de l'équipement.
- Maintenance imparfaite : L'équipement n'est pas ramené à l'état neuf mais sera par contre rajeuni. Cette maintenance se situe entre la maintenance parfaite et la maintenance minimale.

La figure 2.1 suivante montre l'effet de la maintenance parfaite, minimale et imparfaite sur le comportement du taux de défaillance d'un équipement.

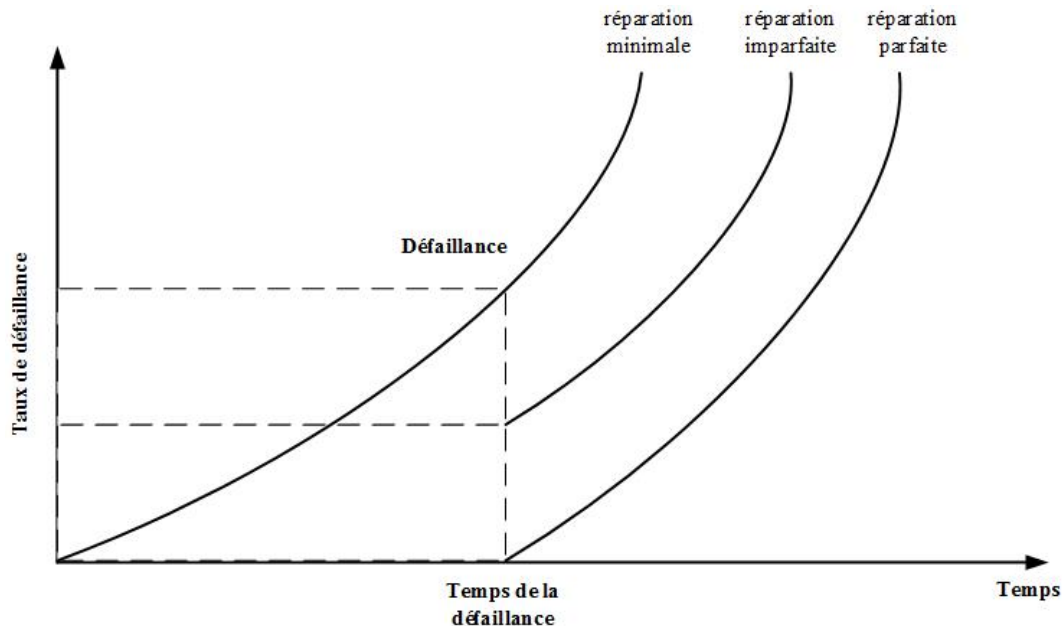


Figure 2.1 Influence de la maintenance parfaite, minimale et imparfaite sur le taux de défaillance

2.2 Méthodes classiques d'estimation paramétrique

Le but ultime de l'usage des outils de la statistique est de produire, à partir des observations liés à un phénomène aléatoire, une inférence portant sur la loi de probabilité paramétrique générant ces observations. Cette loi permet de caractériser ce phénomène, de l'analyser selon l'interprétation donnée aux paramètres et de prédire la probabilité future admettant l'arrivée possible d'une nouvelle observation.

Sans perte de généralité, soit un échantillon de données d'une variable aléatoire X . Cette variable a une densité où seul le vecteur des paramètres θ est inconnu. Les observations sont supposées être la seule source d'information sur le vecteur de paramètres inconnus de la loi de probabilité suivant laquelle ces observations sont distribuées. Pour ce faire, il existe deux méthodes classiques d'estimation des paramètres inconnus (Lejeune (2011)) :

- Estimation paramétrique par intervalle de confiance. Cette estimation permet de donner des valeurs plausibles aux paramètres inconnus sous la forme d'intervalle ou encore « fourchette ».
- Estimation paramétrique ponctuelle. Cette estimation permet de calculer une valeur unique pour chaque paramètre inconnu. La méthode du maximum de vraisemblance est la plus utilisée dans la littérature. À partir des observations $x_1, x_2, \dots, x_n$ d'une variable aléatoire X de densité $f(X|\theta)$ où seul le paramètre θ est inconnu, la fonction de vraisemblance s'écrit $\Pi(X|\theta) = \prod_{i=1}^n f(x_i|\theta)$. L'estimateur de θ au sens du maximum de vraisemblance est $\hat{\theta} = \arg \max_{\theta} (\Pi(X|\theta))$.

Dans le cas où les paramètres sont à leurs tours considérés comme des variables aléatoires, les méthodes classiques d'estimation des paramètres ne peuvent être utilisées. Les méthodes dites d'estimation bayésienne ou d'inférence bayésienne sont généralement proposées dans la littérature. La section suivante décrit ces méthodes et donne leurs usages dans l'ingénierie de la fiabilité.

2.3 Méthodes d'estimation bayésienne

Contrairement aux méthodes classiques, les méthodes d'estimation bayésienne considèrent les paramètres inconnus θ du modèle statistique comme des variables aléatoires. Ces méthodes utilisent le théorème ou la règle de Bayes dans le calcul des probabilités conditionnelles entre événements. Plus particulièrement et sans perte de généralité, chaque paramètre inconnu θ sera modélisé par une densité de probabilité $\Pi(\theta)$ appelée distribution « a priori ». En plus

de l'information provenant des observations, l'estimation bayésienne tiendra compte de cette information supplémentaire provenant de la distribution « a priori » sur le paramètre θ . Une distribution « a posteriori » sur le paramètre θ est calculée en utilisant la combinaison de la fonction de vraisemblance caractérisant les observations $\Pi(x|\theta)$ et l'information a priori $\Pi(\theta)$. L'évaluation de la distribution « a posteriori » est directement donnée par le théorème de Bayes comme suit :

$$\Pi(\theta|x) = \frac{\Pi(x|\theta) \cdot \Pi(\theta)}{\int \Pi(x|\theta) \cdot \Pi(\theta) d\theta}$$

Nous pouvons remarquer facilement que la distribution a posteriori est proportionnelle au numérateur de l'équation précédente. Cette propriété sera d'une grande utilité dans l'évaluation de distribution a posteriori. Ainsi, l'information a priori sur les paramètres inconnus du modèle statistique est actualisée à l'aide des données observées via la fonction de vraisemblance.

Les méthodes bayésiennes ont été fortement utilisées au cours des dernières années dans plusieurs disciplines. Cette croissance est essentiellement due aux développements des méthodes de simulation qui permettent d'évaluer la distribution a posteriori lorsque la forme analytique de la distribution a posteriori n'est pas explicite ou connue. En effet, dans la plupart des situations, il n'est pas toujours évident de connaître la forme analytique de la distribution a posteriori si la fonction de vraisemblance $\Pi(x|\theta)$ et la distribution a priori ne sont pas conjuguées. Par contre, si la fonction de vraisemblance $\Pi(x|\theta)$ est conjuguée avec la distribution a priori $\Pi(\theta)$, la distribution a posteriori sera de même forme que celle de la distribution a priori. Le tableau 2.5 résume quelques distributions a priori conjuguées. La démonstration de la forme analytique de la distribution a posteriori peut être consultée dans Berger (1985).

Tableau 2.5 Quelques distributions usuelles conjuguées

$\Pi(x \theta)$	$\Pi(\theta)$	$\Pi(\theta x)$
<i>Normale</i> (θ, σ^2)	<i>Normale</i> (μ, τ^2)	<i>Normale</i> $(\rho(\sigma^2\mu + \tau^2x), \rho\sigma^2\tau^2)$ où $\rho^{-1} = \sigma^2 - \tau^2$
<i>Poisson</i> (θ)	<i>Gamma</i> (α, β)	<i>Gamma</i> $(\alpha + x, \beta + 1)$
<i>Binomiale</i> (n, θ)	<i>Beta</i> (α, β)	<i>Beta</i> $(\alpha + x, \beta + n - x)$
<i>Gamma</i> (v, θ)	<i>Gamma</i> (α, β)	<i>Gamma</i> $(\alpha + v, \beta - x)$

Lorsque la forme analytique de la distribution a posteriori n'est pas simple, les méthodes de simulation sont mises à contribution pour générer un échantillon suivant la distribution $\Pi(\theta|x)$. Les algorithmes de simulation de type Markov Chain Monte Carlo (MCMC) les plus répandus dans la littérature sont celui de Gibbs (Gelfand *et al.* (1990)) et celui de Metropolis-Hastings (Hastings (1970)). Les procédures des deux algorithmes sont étudiées en détail dans

les sous sections suivantes.

2.3.1 Algorithme de Metropolis-Hastings

L'algorithme de Metropolis-Hastings est une généralisation de la technique de (Metropolis *et al.* (1953)) qui a été utilisée dans la physique. Cet algorithme est basé sur la simulation de chaînes de Markov. Il permet d'obtenir un état de la chaîne à l'état $t + 1$ en simulant un candidat Y à partir d'une distribution dite instrumentale $q(.|X_t)$. Il est indispensable de noter que ce nouveau candidat dépend uniquement de l'état précédent. La loi instrumentale doit satisfaire des conditions dites de régularité (Roberts (1996)). En effet, cette loi doit être irréductible et apériodique (Chib et Greenberg (1995)). L'irréductibilité veut dire qu'il est possible d'atteindre n'importe quel point de la distribution a posteriori avec une probabilité non nulle. L'apériodicité assure que la chaîne n'oscille pas entre deux états différents. Généralement, ces deux conditions sont satisfaites si le support de la loi instrumentale contient celui de la distribution a posteriori. Le nouveau candidat Y est accepté avec une probabilité donnée par :

$$\alpha(X_t, Y) = \min \left(1, \frac{\Pi(Y) \cdot q(X_t|Y)}{\Pi(X_t) \cdot q(Y|X_t)} \right)$$

Si le point Y n'est pas accepté, alors la chaîne ne change pas et $X_{t+1} = X_t$. La procédure de l'algorithme est la suivante :

1. Initialiser la chaîne à X_0 et $t = 0$.
2. Générer un candidat Y à partir de la loi instrumentale $q(.|X_t)$.
3. Générer une valeur aléatoire U à partir de la loi uniforme $[0, 1]$.
4. Si $U \leq \alpha(X_t, Y)$ alors $X_{t+1} = Y$ sinon $X_{t+1} = X_t$.
5. Mettre $t = t + 1$ et recommencer les étapes 2 à 5.

2.3.2 Algorithme de Gibbs

L'algorithme de Gibbs peut être considéré comme un cas particulier de l'algorithme de Metropolis-Hastings (Casella et Robert (1999); Gilks *et al.* (1996)). Cependant, il existe deux différences entre l'algorithme de Gibbs et celui de Metropolis-Hastings. En effet, dans l'algorithme de Gibbs, le nouveau candidat est toujours accepté et il faut connaître toutes les distributions conditionnelles. L'algorithme de Gibbs a été développé essentiellement par (Geman et Geman (1984)) et a été appliqué dans le traitement d'image. Il a été ensuite introduit dans l'analyse statistique des données à travers les articles de Gelfand et Smith (1990) et Gelfand *et al.* (1990).

La procédure de l'algorithme de Gibbs est la suivante :

1. Générer un point initial $X_0 = (X_{0,1}, X_{0,2})$ et initialiser $t = 0$.
2. Générer un point $X_{t,1}$ à partir de la distribution conditionnelle $f(X_{t,1}, X_{t,2} = x_{t,2})$.
3. Générer un point $X_{t,2}$ à partir de la distribution conditionnelle $f(X_{t,2}, X_{t+1,1} = x_{t+1,1})$.
4. Mettre $t = t + 1$ et recommencer les étapes 2 à 4.

2.3.3 Qualité d'ajustement d'un modèle bayésien

Dans la statistique bayésienne, la qualité d'ajustement d'un modèle donné peut être évaluée à partir de plusieurs techniques. Parmi ces techniques, la déviance bayésienne et l'analyse du caractère prédictif du modèle sont les plus performantes et les plus répandues.

Déviance bayésienne

Le DIC (Deviance information criterion) est basé sur la distribution a posteriori de la déviance (Spiegelhalter *et al.* (2002, 1998)). Le DIC est défini comme étant $DIC = \bar{D} + p_D$ avec $\bar{D} = E(-2 \log(\Pi(x|\theta)))$ est l'espérance a posteriori de la déviance qui peut être considérée comme une mesure d'ajustement aux données ($\Pi(x|\theta)$ est la fonction de vraisemblance du modèle et θ désigne tous les paramètres du modèle) et $p_D = \bar{D} - D(\bar{\theta})$ est appelé le nombre effectif des paramètres (Carlin et Louis (2000)). En effet, $D(\bar{\theta}) = -2 \log(\Pi(x|\bar{\theta}))$ avec $\bar{\theta}$ est la moyenne a posteriori des paramètres θ . p_D peut être considéré comme un facteur mesurant la complexité du modèle. En pratique, le calcul du DIC se fait en utilisant les méthodes d'échantillonnage MCMC (algorithme de Gibbs ou celui de Metropolis-Hastings). À partir d'un échantillon $\theta^{(1)}, \theta^{(2)}, \dots, \theta^{(L)}$, il est possible de calculer $\bar{D} = \frac{1}{L} \left(-2 \log \prod_{l=1}^L \Pi(x|\theta^{(l)}) \right)$. L'évaluation des modèle selon le DIC est gérée suivant le principe plus le DIC est faible plus le modèle est meilleur.

Analyse du caractère prédictif du modèle

L'analyse du caractère prédictif du modèle bayésien est une technique très utilisée dans la mesure de la qualité d'ajustement du modèle. Elle est simple à implémenter et à facile à interpréter (Sinharay et Stern (2003); Gelman *et al.* (1996)). L'idée de cette analyse est de simuler un échantillon x^{rep} à partir de la distribution à posteriori de θ et de le comparer avec les données observées x . Si l'échantillon simulé est semblable aux données observées alors il est possible de conclure que la qualité d'ajustement du modèle est bonne. Il faut choisir une certaine statistique $T(y, \theta)$ pour permettre de comparer l'échantillon simulé et les données observées. Il faut calculer $T(y, \theta)$ et $T(y^{rep}, \theta)$ puis les comparer en les traçant sur un même graphique.

Il est possible également de mesurer quantitativement le manque d'ajustement du modèle en introduisant le p-value bayésien (Sinharay *et al.* (2006)). Le p-value bayésien est défini comme étant la probabilité que $T(y^{rep}, \theta)$ soit supérieure ou égale à $T(y, \theta)$. Un p-value bayésien proche de 0 ou 1 indique que le modèle ne prédit pas bien les données observées. Ainsi, l'évaluation des modèles en utilisant le critère du p-value bayésien suit le principe suivant : plus le p-value bayésien est proche de 0,5, meilleur est le modèle. Il ne faut pas confondre le p-value bayésien avec le p-value traditionnel permettant des tests d'hypothèses.

Il est à noter que le DIC permet de comparer au moins deux modèles. En effet, un seul DIC ne permet pas de faire une comparaison. Contrairement au DIC, nous pouvons juger la qualité d'un seul modèle en utilisant son p-value bayésien et ce ceci en le comparant à 0,5.

2.4 Revue des modèles de combinaisons des avis d'experts

Dans cette section, nous allons nous intéresser aux modèles permettant de combiner plusieurs avis d'experts exprimés sur un paramètre inconnu. Comme vu dans la section précédente, une distribution a priori sur un paramètre inconnu jumelé avec la distribution de vraisemblance décrivant les données observées permet d'évaluer la distribution a posteriori de ce paramètre en procédant par l'inférence bayésienne. Cette distribution a priori est généralement construite sur la base de connaissance sur le paramètre d'intérêt par la personne (preneur de décision noté dans la suite DM) qui va prendre une décision ou qui veut prédire le futur.

Supposons maintenant que le DM a des connaissances limitées sur le paramètre d'intérêt. Le DM doit alors consulter une autre personne qui est supposée avoir d'amples connaissances sur ce paramètre, donc considérée comme expert en vue de ce paramètre. Si le DM prévoit que la décision à prendre est d'une grande importance, il serait donc judicieux de consulter plus qu'un expert. Le recours à plusieurs experts est justifié par l'importance du problème ainsi bien que par la confiance que peut attribuer le DM aux différents experts en sa disposition.

Dans le cas où il s'agit d'un seul expert, la construction de la distribution a priori est simple, il suffit d'élucider les informations nécessaires de l'expert et les transformer en une distribution de probabilité. Éliciter l'avis d'un expert est une procédure un peu complexe et qui demande beaucoup d'expertises vu qu'il va falloir poser des questions tangibles à l'expert puis transformer les réponses correspondantes en des quantités statistiques permettant de construire la distribution a priori. Si plusieurs experts sont consultés, il faut éliciter l'avis de chaque expert, pour obtenir une distribution a priori de chacun. Le problème qui se présente est comment utiliser toutes les distributions a priori provenant des experts pour construire

une seule distribution a priori. En effet, l'inférence bayésienne requiert la présence d'une seule distribution a priori pour chaque paramètre inconnu, d'où la nécessité de combiner les avis d'experts.

Dans la littérature, il existe deux façons pour combiner les avis d'experts : la combinaison mathématique des distributions de probabilités et la combinaison se basant sur un accord ou consensus entre les experts appelée combinaison comportementale. La combinaison comportementale consiste à amener un groupe d'experts et puis éliciter leurs avis comme s'ils étaient un seul expert (Wallach *et al.* (1962)). Cette approche de combiner les avis repose essentiellement sur un dialogue qui doit avoir lieu entre les différents experts pour ressortir une distribution a priori unique. La distribution a priori résultante est donc une sorte d'accord ou de consensus entre les experts par rapport au paramètre inconnu. Cependant, cette distribution peut ne pas exister si les experts ont des avis complètement opposés (très différents) et qu'ils n'arrivent pas à trouver un terrain d'entente.

La combinaison mathématique des avis d'experts, quant à elle, a été largement traitée dans la littérature. Nous présentons dans la suite les deux principales méthodes de combinaison de ces avis : la méthode du mélange avec une pondération des avis fixée à l'avance et la méthode de la moyenne.

2.4.1 Méthode du mélange

Supposons que nous avons N experts ($E_1, \dots, E_N$) qui ont été consultés et que nous avons associé à chacun une distribution a priori $\Pi_j(\theta)$ sur un paramètre inconnu θ (ou un vecteur de paramètres). Le problème est de déterminer à partir des $\Pi_j(\theta)$ une distribution a priori unique $\Pi(\theta)$ qui représente ainsi un consensus entre les avis d'experts. La méthode du mélange est considérée comme une approche principale dans les travaux de (O'Hagan *et al.* (2006); Genest et Zidek (1986); Albert *et al.* (2012)). La méthode du mélange est une combinaison linéaire des distributions a priori de chaque expert :

$$\Pi(\theta) = \sum_{j=1}^N w_j \cdot \Pi_j(\theta) \text{ avec } w_j \geq 0 \text{ et } \sum_{j=1}^N w_j = 1$$

Une restriction de cette méthode est que les poids w_i doivent être fixés et connus d'avance avant de procéder au mélange des avis. Leur somme doit être égale à 1. Cependant, rien n'empêche que les distributions a priori choisies par chaque expert soient de natures différentes (plusieurs types de distributions de probabilités peuvent être mélangés). Le seul problème revient à déterminer convenablement le poids w_i de chaque expert dans le mélange.

D'une manière générale, il est raisonnable d'attribuer un poids plus important à $\Pi_j(\theta)$ (donc à E_j) comparé à $\Pi_l(\theta)$ (donc à E_l) si E_j est jugé «meilleur» ou «plus compétent» que E_l . La question qui est le meilleur expert ou l'expert le plus compétent n'admet pas de réponse objective. La «performance» d'un expert doit être vue comme une question subjective. Les travaux de (Cooke et Goossens (2008); Winkler (1968); French (2011)) ont traité la question de l'estimation des poids w_j . Ils précisent que dans plusieurs cas, le DM serait amené à suivre l'une des règles suivantes pour l'attribution d'un poids w_j à chaque expert E_j :

1. Il attribue des poids égaux à $w_j = \frac{1}{N}, i = 1, 2, \dots, N$ pour chaque expert. Dans ce cas, le DM ne pense pas qu'il existe une grande différence entre les experts.
2. Le DM attribue des poids proportionnels à un ordonnancement des experts. Il ordonne les experts suivant leurs 'performances'. Il affecte le poids $\frac{r}{\sum_{r=1}^N r}$ à l'expert dont le rang est r ($r = 1, \dots, N$). Cette règle suppose que le DM est capable d'ordonner correctement les experts suivant leurs performances.
3. Chaque expert affecte à lui-même une note de 1 à g (g désigne une note maximale fixée). Le DM attribue ensuite un poids proportionnel à la note de chaque expert (la constante de proportionnalité est déterminée de telle sorte que les poids somment à 1).
4. Le DM attribue aux experts des poids suivant leurs performances passées.

L'avantage majeur de la méthode du mélange est qu'elle permet de prendre en compte toutes les valeurs possibles prévues par les experts dans leurs estimations sur les paramètres inconnus. Cet avantage qui pourrait dans certains cas être un inconvénient constitue la différence avec la méthode de la moyenne qui sera traitée dans ce qui suit.

2.4.2 Méthode de la moyenne

La méthode de la moyenne a été l'objet de plusieurs travaux dans la littérature (Burgman *et al.* (2011); Lipscomb *et al.* (1998)). L'avantage de cette méthode est qu'elle est simple à implémenter mais son inconvénient est qu'elle peut ne pas tenir compte des incertitudes que les experts prévoient sur les paramètres inconnus. En effet, cette méthode constitue un accord entre les différents experts qui sont consultés. Cette méthode nécessite aussi l'utilisation d'un même type de distribution a priori pour tous les experts consultés.

Pour simplifier la présentation de la méthode, nous allons supposer que la modélisation de la distribution a priori pour chaque expert va être faite avec une loi $beta(a_j, b_j)$. La distribution a priori $\Pi(\theta)$ résultante de la combinaison avec la méthode de la moyenne sera une loi $beta(a, b)$ avec $a = \sum_{j=1}^N w_j \cdot a_j$ et $b = \sum_{j=1}^N w_j \cdot b_j$. Encore une fois, la question des poids se

présente. Cependant, pour la méthode de la moyenne, non seulement les poids doivent être déterminés mais encore leur somme. Winkler (1968) suggère que la somme des poids doit être comprise entre 1 et N (le nombre d'experts). En effet, la somme des poids peut être égale à N si les experts proviennent de domaines d'expertises différents. Si les experts sont du même domaine d'expertise, la somme des poids peut être prise égale à 1. Dans la réalité, le cas le plus répandu est celui où la somme est égale à 1. Cela peut être justifié, car le DM va consulter les experts pour formuler leurs avis sur un paramètre donc seront nécessairement experts en vue de ce dernier et ont la même expertise dans le domaine.

La méthode de la moyenne, bien qu'elle soit simple à implémenter, est plus avantageuse à utiliser si les avis d'experts sont proches. Cependant, l'utilisation de la méthode du mélange est recommandée si les avis d'experts sont différents, car dans ce cas, la méthode de la moyenne ne permet pas de tenir compte de l'incertitude des experts puisqu'elle considère une distribution a priori résultante qui est une référence pour tous les experts.

Dans le chapitre 3, une méthode d'agrégation d'avis d'expert sera élaborée. Cette méthode est une version améliorée de la méthode du mélange. En effet, les poids w_j accordés à chaque expert j seront actualisés avec les données de la vraisemblance. Cette méthode d'agrégation sera comparée aux méthodes du mélange et de la moyenne et sera validée à l'aide du test statistique non paramétrique permettant de vérifier la tendance des données.

2.5 Stratégies d'optimisation de la disponibilité

Dans cette section, nous allons présenter une revue de littérature des stratégies de maintenance qui maximisent la disponibilité d'un équipement. En effet, à part les nombreuses stratégies qui ont pour but de minimiser les coûts de maintenance, certaines stratégies de maintenance traitent le cas de la maximisation de la disponibilité. Ceci s'explique par le fait que l'exploiteur d'un équipement est parfois incapable d'estimer les coûts liés à la réparation ou le remplacement d'un équipement, ou qu'il souhaite profiter le plus de l'utilisation de l'équipement, dans ce cas il serait judicieux de penser à maximiser sa disponibilité.

Jardine et Tsang (2013) traitent deux stratégies de maintenance permettant de maximiser la disponibilité d'un équipement. La première stratégie consiste à trouver le temps optimal pour faire un remplacement préventif d'un équipement qui est sujet à des défaillances et qu'il est remplacé par un neuf (maintenance parfaite) à chaque fois où il tombe en défaillance (stratégie de type bloc). Dans la deuxième stratégie, l'équipement est remplacé par un neuf lorsqu'il tombe en défaillance mais son remplacement préventif se fait dépendamment de son âge (stratégie de type âge).

Kececioglu (2003) présente une stratégie qui permet de déterminer l'intervalle optimal entre deux inspections de l'équipement permettant de maximiser sa disponibilité. À chaque fois où l'équipement est inspecté et qu'il est défaillant, il sera réparé ou remplacé de telle sorte qu'il est restauré à l'état neuf (maintenance parfaite).

Kabak (1969) présente un modèle d'optimisation de la disponibilité d'un équipement. Les durées de vie de ce dernier sont distribuées selon une loi exponentielle et la durée de sa réparation est supposée constante. Le modèle permet de suivre la variation de la disponibilité et de prédire les durées moyennes de bon fonctionnement et de réparation permettant de minimiser le coût d'exploitation de l'équipement.

Barlow et Proschan (1972) présentent un modèle d'optimisation de la disponibilité pour des systèmes en série. Ils supposent dans ce travail que la maintenance est parfaite. Zhao (1994) traite la même problématique, mais il suppose que la maintenance est imparfaite. Wang et Pham (2006a) traitent aussi l'optimisation de la disponibilité des systèmes en série dans le cas où ces systèmes sont sujets à des réparations imparfaites. Plusieurs autres travaux ont traité l'optimisation de la disponibilité des équipements sujets à des stratégies de maintenance imparfaite. Nous pouvons notamment citer les travaux de Wang et Pham (1999); Chan et Shaw (1993).

Dans le cadre de ce mémoire, une stratégie de maintenance maximisant la disponibilité d'un équipement sera élaborée dans le chapitre 4. Cette stratégie est adaptée du modèle d'optimisation des coûts d'une stratégie de réparation minimale de Barlow et Proschan (1996). Cette stratégie sera modifiée pour permettre la maximisation de la disponibilité d'un équipement. Cette stratégie, telle que décrite dans (Kececioglu, 2003), permet de minimiser le coût moyen par unité de temps du remplacement préventif d'un équipement soumis à des réparations minimales lors des défaillances. Le but est de trouver l'intervalle optimal entre deux remplacements préventifs. Plus particulièrement, à chaque fois où l'équipement tombe en défaillance, il subit une réparation minimale de telle sorte que son taux de défaillance reste presque le même qu'avant sa défaillance. Kececioglu (2003) justifie que le taux de défaillance demeure le même après la réparation minimale par le fait qu'il traite un équipement multi-composant, et que la réparation d'un sous ensemble de composants n'affecte pas le taux de défaillance global de l'équipement. Les paramètres de cette stratégie sont les suivantes :

- C_p : Coût moyen du remplacement préventif.
- C_{rm} : Coût moyen d'une réparation minimale.
- $\lambda(t)$: Taux de défaillance de l'équipement.
- t_p : Intervalle entre deux remplacements préventifs (c'est la variable de décision que nous cherchons à déterminer).

Le coût moyen par unité de temps sur un horizon infini s'exprime comme suit :

$$C(t_p) = \frac{C_p + C_{rm} \cdot H(t_p)}{t_p}$$

Où $H(t_p)$ désigne l'espérance du nombre de défaillances suivies par une réparation minimale. $H(t_p)$ se calcule comme suit :

$$H(t_p) = \int_0^{t_p} \lambda(t) dt$$

Il suffit que le taux de défaillance soit strictement croissant pour qu'un minimum unique t_p^* de la fonction $C(t_p)$ existe.

Dans le chapitre 4, cette stratégie servira de base de modélisation de la stratégie de maintenance maximisant la disponibilité d'un équipement. Elle sera modifiée pour tenir compte des temps de réparation minimale et de remplacement préventif. De plus, des avis d'experts seront intégrés dans le but d'estimer les taux de défaillance et de réparation actualisés d'une période à l'autre sur un horizon infini.

CHAPITRE 3

MODÈLE D'AGRÉGATION BAYÉSIENNE

Dans ce chapitre, nous allons proposer un modèle d'agrégation d'avis d'experts exprimés sur un paramètre inconnu. Nous commençons dans un premier temps par présenter une mise en contexte et une explication des données sur lesquelles nous allons travailler. Ensuite, nous présentons un modèle d'agrégation d'avis d'experts ainsi que la démarche de résolution que nous avons adoptée. Nous allons comparer ce modèle avec les modèles du mélange et de la moyenne traités dans la littérature afin de ressortir les avantages et les inconvénients du modèle proposé. La comparaison est faite à l'aide de l'analyse des résultats obtenus avec notre échantillon de données. Enfin, nous allons valider les résultats de ce modèle en utilisant un test non paramétrique vérifiant la tendance des données.

3.1 Mise en contexte

Dans le cadre de ce travail, nous considérons un équipement neuf mis en exploitation à $t = 0$. Nous supposons que les durées de bon fonctionnement de cet équipement suivent une loi exponentielle de paramètre λ . Le fournisseur de cet équipement fournit un taux de défaillance de base noté λ_0 . L'équipement sera exploité pendant plusieurs périodes et les durées de bon fonctionnement ainsi que le nombre de défaillances seront enregistrées pour chacune des périodes. Ainsi, à chaque période, une vraisemblance sera constituée pour l'équipement.

Au début de chaque période, plusieurs experts en fiabilité d'équipements seront consultés pour donner leurs avis sur le taux de défaillance de la période subséquente. Les experts se basent sur les actions qu'avait subit l'équipement au cours de la période précédente ou qu'il subira (maintenance, inspection, réparation...) au cours de la période ultérieure pour formuler leurs avis. Ainsi, une distribution a priori sur le taux de défaillance sera construite pour chaque expert au début de chaque période.

Un modèle d'agrégation d'avis d'experts, qui sera formulé dans la section suivante, va nous permettre d'avoir un seul avis qui serait pris en compte au début de chaque période. La vraisemblance de chacune des périodes sera modélisée par un processus de Poisson homogène de paramètre λ . Une fois un avis commun pour tous les experts soit formulé et que la vraisemblance soit connue, nous serons en mesure d'effectuer une inférence bayésienne sur le paramètre λ .

L'inférence bayésienne sera faite à l'aide de l'algorithme de Metropolis-Hasting. Nous aurons donc une distribution a posteriori sur le paramètre λ caractérisée par sa moyenne et son écart-type. Nous aurons également un classement des avis d'experts qui ont été formulés au début de la période. Ce classement sera validé à l'aide de trois tests statistiques : le DIC (Deviance Information Criterion), le p-value bayésien et le test non paramétrique de la tendance des données lorsque les avis sont exprimés sur plusieurs périodes.

3.2 Modélisation de la vraisemblance et des avis d'experts

Dans cette section, nous allons construire dans un premier temps notre fonction de vraisemblance à partir des données de bon fonctionnement de l'équipement. Par la suite, nous allons modéliser la distribution a priori sur le paramètre λ pour chaque expert en fonction de son avis. Le modèle d'agrégation des avis d'experts sera ensuite formulé à l'aide de l'inférence bayésienne.

3.2.1 Calcul de la fonction de vraisemblance

Pour une période k donnée ($k \in \{1, \dots, K\}$), nous disposons des durées de bon fonctionnement de l'équipement $x_{i,k}$ ($i \in \{1, \dots, n_k\}$). Nous supposons que pour chaque période k nous avons n_k défaillances enregistrées. Nous supposons également que les durées de bon fonctionnement de l'équipement suivent une loi exponentielle de paramètre λ de densité $f(t, \lambda) = \lambda \exp(-\lambda t)$. Nous pouvons donc écrire la fonction de vraisemblance comme suit :

$$\begin{aligned} L_k(X, \lambda) &= \prod_{i=1}^{n_k} f(x_{i,k}, \lambda) \\ &= \prod_{i=1}^{n_k} \lambda \exp(-\lambda x_{i,k}) \quad \forall k \in \{1, \dots, K\} \\ &= \lambda^{n_k} \exp\left(-\lambda \sum_{i=1}^{n_k} x_{i,k}\right) \end{aligned} \tag{3.1}$$

Où X désigne l'ensemble de toutes les durées de bon fonctionnement de l'équipement pour la période k . Cette fonction de vraisemblance sera combinée avec la distribution a priori sur le paramètre λ pour ensuite procéder à l'inférence bayésienne.

L'estimateur du maximum de vraisemblance du paramètre λ pour la période k est donné par l'équation 3.2.

$$\hat{\lambda}_k = \frac{n_k}{\sum_{i=1}^{n_k} x_{i,k}} \quad \forall k \in \{1, \dots, K\} \quad (3.2)$$

3.2.2 Modélisation des avis d'experts

Nous supposons qu'au début de chaque période, plusieurs experts en fiabilité donnent leurs avis à propos du taux de défaillance de l'équipement λ . Il existe plusieurs façons pour éliciter les avis d'experts. Par exemple, les experts peuvent être questionnés sur la moyenne et l'écart-type du taux de défaillance. Cependant, il est préférable de leur poser des questions tangibles puisqu'ils peuvent rencontrer des difficultés en traitant des quantités statistiques (moyenne, écart-type, quantiles...). Ainsi, les experts seront questionnés sur les bornes minimale et maximale du paramètre λ . Cela nous semble moins contraignant.

Supposons que nous avons N experts, $N \geq 1$, chaque expert j donne une estimation des bornes minimale et maximale du taux de défaillance, respectivement $\lambda_{\min,j}$ et $\lambda_{\max,j}$. À partir des deux informations obtenues de chaque expert j , nous construisons la distribution a priori $\Pi_j(\lambda)$ sur le paramètre λ . Nous avons choisi de modéliser cette distribution par une loi bêta(a_j, b_j). Le choix de loi bêta pourrait être expliqué par le fait que celle-ci est fréquemment utilisée si le paramètre à modéliser est borné (Bacha *et al.* (1998)). Cette utilisation est adaptée à notre cas d'étude puisque le taux de défaillance de l'équipement peut être considéré comme borné. En effet, l'équipement est soumis au principe de la durabilité économique. En effet, si le taux de défaillance augmente au delà d'un certain seuil, il sera plus avantageux et rentable de remplacer l'équipement au lieu de procéder à sa réparation. La détermination des paramètres a_j et b_j est faite en solvant le système à deux équations non linéaires 3.3 pour chaque expert j . F_j désigne la fonction de répartition de la loi bêta(a_j, b_j). La probabilité p est un choix de modélisation, nous pouvons considérer $p = 0,05$.

$$\begin{cases} F_j(\lambda_{\min,j}, a_j, b_j) = p \\ F_j(\lambda_{\max,j}, a_j, b_j) = 1 - p \end{cases} \quad \forall j \in \{1, \dots, N\} \quad (3.3)$$

La résolution du système 3.3 se fait numériquement avec *Matlab*. En définitive, la distribution a priori pour chaque expert j est donnée par la densité de la loi bêta(a_j, b_j). L'équation 3.4 donne l'expression de cette distribution $\Pi_j(\lambda)$ pour chaque expert j . Γ désigne la fonction Gamma.

$$\Pi_j(\lambda) = \frac{\Gamma(a_j + b_j)}{\Gamma(a_j)\Gamma(b_j)} \lambda^{a_j-1} (1 - \lambda)^{b_j-1} \mathbb{1}_{[0,1]}(\lambda) \quad \forall j \in \{1, \dots, N\} \quad (3.4)$$

Nous devons, une fois les avis d'experts formulés, combiner ces avis pour obtenir une distribution a priori unique sur λ . Pour ce faire, nous introduisons les deux variables aléatoires suivantes :

- Λ : Une variable aléatoire continue désignant le taux de défaillance λ de l'équipement.
- J : Une variable aléatoire discrète désignant l'expert en fiabilité le plus performant. La performance de l'expert est mesurée par la justesse de son avis sur le paramètre λ .

La fonction de masse de la variable J est donnée par l'équation 3.5.

$$P(J = j) = \frac{1}{N} \quad \forall j \in \{1, \dots, N\} \quad (3.5)$$

Nous avons choisi d'attribuer des probabilités égales à $\frac{1}{N}$ pour chacun des experts interrogés. Ce choix se justifie par le fait que nous n'avons aucune idée sur la pertinence des avis d'experts. En prenant des probabilités égales, nous ne défavorisons aucun expert. C'est la vraisemblance qui va juger quant à la justesse des avis de ces experts. En faisant ce choix d'introduire la nouvelle variable J , la distribution a priori sur λ pour chaque expert j sera exprimée par la fonction de densité conditionnelle de Λ , étant donné que $J = j$. Cette distribution a priori s'écrit comme le montre l'équation 3.6. La distribution a priori sur λ sera exprimée par l'équation 3.7 en appliquant l'axiome des probabilités totales.

$$\Pi_{\Lambda|J}(\lambda|j) = \frac{\Gamma(a_j + b_j)}{\Gamma(a_j)\Gamma(b_j)} \lambda^{a_j-1} (1 - \lambda)^{b_j-1} \mathbb{1}_{[0,1]}(\lambda) \quad \forall j \in \{1, \dots, N\} \quad (3.6)$$

$$\begin{aligned} \Pi_{\Lambda}(\lambda) &= \sum_{j=1}^N P(J = j) \Pi_{\Lambda|J}(\lambda|j) \\ &= \frac{1}{N} \sum_{j=1}^N \frac{\Gamma(a_j + b_j)}{\Gamma(a_j)\Gamma(b_j)} \lambda^{a_j-1} (1 - \lambda)^{b_j-1} \mathbb{1}_{[0,1]}(\lambda) \end{aligned} \quad (3.7)$$

Nous notons par Θ le vecteur aléatoire mixte ($\Theta = (\Lambda, J)$). Nous construisons la distribution a posteriori sur le paramètre $\theta = (\lambda, j)$. C'est l'inférence bayésienne qui déterminera cette distribution a posteriori de θ . Cette distribution a posteriori sera ainsi évaluée pour un λ donné et un expert j donné. Étant donné n durées de bon fonctionnement de l'équipement, notées x_i ($i = 1, \dots, n$), le théorème de Bayes permet d'écrire la distribution a posteriori du paramètre θ comme suit :

$$\begin{aligned}
\Pi_{\Theta|X}(\theta|x) &= \frac{\Pi_{X|\Theta}(x|\theta)\Pi_{\Theta}(\theta)}{\Pi_X(x)} \\
&= \frac{\Pi_{X|\Theta}(x|\theta) \cdot \Pi_{\Theta}(\theta)}{\int_0^{+\infty} \Pi_{X|\Theta}(x|\theta) \cdot \Pi_{\Theta}(\theta) d\theta}
\end{aligned} \tag{3.8}$$

Le dénominateur de l'équation 3.8 est une constante de normalisation qui est indépendante de θ . Donc, la distribution a postérieure sera proportionnelle au numérateur de cette équation. $\Pi_{X|\Theta}(x|\theta)$ est la fonction de vraisemblance des données x sachant le paramètre θ et $\Pi_{\Theta}(\theta)$ est la distribution a priori sur le paramètre θ . En définitive, l'équation 3.9 donne l'expression de la distribution a postérieure sur le paramètre θ . Le symbole $\propto$ désigne proportionnel.

$$\begin{aligned}
\Pi_{\Theta|X}(\theta|x) &\propto \Pi_{X|\Theta}(x|\theta)\Pi_{\Theta}(\theta) \\
&\propto \lambda^n \exp\left(-\lambda \sum_{i=1}^n x_i\right) \cdot \Pi_{\Lambda,J}(\lambda, j) \\
&\propto \lambda^n \exp\left(-\lambda \sum_{i=1}^n x_i\right) \cdot P(J = j) \cdot \Pi_{\Lambda|J}(\lambda|j)
\end{aligned} \tag{3.9}$$

Dans le cas où la distribution a priori et la fonction de vraisemblance sont conjuguées, la forme analytique de la distribution a postérieure sera connue. Cependant, dans notre cas d'étude, il va falloir utiliser les méthodes de simulation *MCMC* pour calculer cette distribution a postérieure puisque la distribution a priori (loi bêta) et la fonction de vraisemblance (loi exponentielle) ne sont pas conjuguées. La démarche de résolution sera traitée en détail dans la section suivante.

3.2.3 Algorithme de résolution

Comme la forme analytique de la distribution a postérieure sur le paramètre θ est inconnue, nous utilisons la méthode de simulation *MCMC*, en particulier l'algorithme de Metropolis-Hasting. Pour notre modèle 3.9, la procédure de cet algorithme s'écrit comme suit :

1. initialiser $[\lambda^{(1)}, j^{(1)}]$
2. Pour $l = 1$ à R
 - Générer $u \sim U_{[0,1]}$
 - Générer $\theta^* \sim q(\theta^*|\theta^{(l)})$
 - Si $u < A$ alors $\theta^{(l+1)} = \theta^*$ sinon $\theta^{(l+1)} = \theta^{(l)}$

R est le nombre d'échantillons et $U_{[0,1]}$ est la densité de la loi uniforme. A est appelé le ratio d'acceptation, son expression est donnée par l'équation 3.10 :

$$A = \min \left(1, \frac{\Pi_{\Theta|X}(\theta^*|x)q(\theta^{(l)}|\theta^*)}{\Pi_{\Theta|X}(\theta^{(l)}|x)q(\theta^*|\theta^{(l)})} \right) \quad (3.10)$$

$q(\theta^*|\theta^{(l)})$ désigne la loi instrumentale. Cette loi permet de générer un nouveau candidat θ^* à partir de l'ancien point $\theta^{(l)}$. En effet, à chaque itération l de l'algorithme, nous devons générer un λ^* et un j^* . Nous avons choisi de générer λ^* à partir d'une loi normale $N(\lambda^{(l)}, \sigma^2)$. L'écart-type σ est un choix à faire. Le choix de la loi normale est justifié car cette loi est symétrique et par conséquent le calcul du ratio d'acceptation sera simplifié (Martinez et Martinez (2001)). Un bon choix de σ permettra à l'algorithme de converger rapidement vers la distribution a posteriori de λ . En ce qui concerne j^* , un expert est choisi aléatoirement à partir des N experts disponibles. La probabilité de choisir un expert quelconque est égale à $\frac{1}{N}$. L'équation 3.11 donne l'expression de la loi instrumentale pour le modèle 3.9.

$$\begin{aligned} q(\theta^*|\theta^{(l)}) &= q(\lambda^*, j^*|\lambda^{(l)}, j^{(l)}) \\ &= q(\lambda^*|\lambda^{(l)}, j^{(l)}) \cdot q(j^*|\lambda^{(l)}, j^{(l)}) \\ &= \frac{1}{N\sqrt{2\pi}\sigma} \exp \left(-\frac{(\lambda - \lambda^{(l)})^2}{2\sigma^2} \right) \end{aligned} \quad (3.11)$$

Nous remarquons bien que $q(\theta^{(l)}|\theta^*) = q(\theta^*|\theta^{(l)})$. Ceci nous permettra de simplifier l'équation du ratio d'acceptation A . L'équation 3.12 donne son expression résultante.

$$\begin{aligned}
A &= \min \left(1, \frac{\Pi_{\Theta|X}(\theta^*|x)}{\Pi_{\Theta|X}(\theta^{(l)}|x)} \right) \\
&= \min \left(1, \frac{(\lambda^*)^n \exp \left(-\lambda^* \sum_{i=1}^n x_i \right) \overline{P(J=j^*)} \Pi_{\Lambda|J}(\lambda^*|j^*)}{(\lambda^{(l)})^n \exp \left(-\lambda^{(l)} \sum_{i=1}^n x_i \right) \overline{P(J=j^{(l)})} \Pi_{\Lambda|J}(\lambda^{(l)}|j^{(l)})} \right) \\
&= \min \left(1, \frac{(\lambda^*)^n \exp \left(-\lambda^* \sum_{i=1}^n x_i \right) \Pi_{\Lambda|J}(\lambda^*|j^*)}{(\lambda^{(l)})^n \exp \left(-\lambda^{(l)} \sum_{i=1}^n x_i \right) \Pi_{\Lambda|J}(\lambda^{(l)}|j^{(l)})} \right) \\
&= \min \left(1, \frac{(\lambda^*)^n \exp \left(-\lambda^* \sum_{i=1}^n x_i \right) \Pi_{\Lambda|J}(\lambda^*|j^*)}{(\lambda^{(l)})^n \exp \left(-\lambda^{(l)} \sum_{i=1}^n x_i \right) \Pi_{\Lambda|J}(\lambda^{(l)}|j^{(l)})} \right) \\
&= \min \left(1, \left(\frac{\lambda^*}{\lambda^{(l)}} \right)^n \exp \left(-(\lambda^* - \lambda^{(l)}) \sum_{i=1}^n x_i \right) \frac{\Pi_{\Lambda|J}(\lambda^*|j^*)}{\Pi_{\Lambda|J}(\lambda^{(l)}|j^{(l)})} \right)
\end{aligned} \tag{3.12}$$

Ainsi, nous pouvons déterminer la moyenne $\bar{\lambda}_{post}$ et l'écart-type σ_{post} de la distribution a postérieure de λ . Nous pouvons également expliciter la fonction de masse a postérieure de la variable aléatoire J . Les équations 3.13 et 3.14 donnent respectivement les statistiques de la distribution a postérieure de λ et la fonction de masse de la variable J .

$$\begin{cases} \bar{\lambda}_{post} = \frac{1}{R} \sum_{l=1}^R \lambda^{(l)} \\ \sigma_{post} = \sqrt{\frac{1}{R-1} \sum_{l=1}^R (\lambda^{(l)} - \bar{\lambda}_{post})^2} \end{cases} \tag{3.13}$$

$$P_{post}(J = j) = \frac{1}{R} \sum_{l=1}^R \mathbb{1}_{\{j^{(l)}=j\}} \quad \forall j \in \{1, \dots, N\} \tag{3.14}$$

Nous avons développé un programme avec *Matlab* qui permet de combiner les avis d'experts avec le modèle précédent (Annexe A). Ce modèle permet de calculer les statistiques de la distribution a postérieure sur le paramètre λ . Il permet également de classer les experts suivant la fonction de masse de la variable aléatoire J . Ce programme prend comme entrée les avis d'experts donnés sous forme de bornes minimales et maximales sur λ et les durées de bon fonctionnement de l'équipement.

3.2.4 Illustration du modèle d'agrégation

Supposons que nous avons 2 experts qui donnent leurs avis à propos du taux de défaillance λ . Ces deux avis sont considérés indépendants. Nous avons aussi n données simulées de bon fonctionnement de l'équipement notées par x_i (tableau 3.1). L'estimateur de maximum de vraisemblance du paramètre λ , noté $\hat{\lambda}$, est égal à $\frac{n}{\sum_{i=1}^n x_i}$ soit $\hat{\lambda} = 0,4121$ défaillance/semaine.

Tableau 3.1 Durées de bon fonctionnement exprimées en semaine

0,9456	3,2482	0,9082	1,3518	1,1053	0,1397	4,2559	1,1054	2,8813
5,6272	4,6231	11,9487	0,5374	0,6710	3,8548	1,5933	3,6966	3,9795
6,1936	0,7986	0,8894	2,4179	1,9098	0,5552	0,1473	0,2458	0,1514
5,0576	3,7092	0,3226	0,8216	0,9702	1,4566	0,4827	1,0310	0,1072
0,1075	12,1880	2,2952	2,5573	2,6122				

Le tableau 3.2 résume les avis des 2 experts.

Tableau 3.2 Avis des 2 experts et paramètres des distributions a priori

Expert (j)	$\lambda_{\min,j}$	$\lambda_{\max,j}$	a_j	b_j
1	0,33	0,47	52,4421	78,9513
2	0,46	0,77	15,8663	9,7055

La détermination des paramètres a_j et b_j pour chaque expert j est faite en solvant le système d'équations non linéaires 3.3 à l'aide de *Matlab*. Nous allons dans la suite étudier 3 modèles :

- Modèle 1 : Ce modèle considère l'expert 1 comme la seule source d'information pour formuler la distribution a priori sur le paramètre λ .
- Modèle 2 : Pour ce modèle, seul l'expert 2 sera pris en compte dans la distribution a priori sur λ .
- Modèle 3 : Ce modèle tient compte simultanément des deux avis d'experts. Le modèle d'agrégation des deux avis est celui proposé dans la section précédente.

L'étude des 3 modèles nous permettra de quantifier la qualité de la distribution a priori de chacun. Pour ce faire, nous utilisons deux quantités statistiques qui vont mesurer la qualité d'ajustement de chaque modèle. Il s'agit du DIC (Deviance Information Criterion) et le p-value bayésien (section 2.3.3). Le calcul de chaque quantité se fait à l'aide de l'échantillon de λ que nous avons obtenu en appliquant l'algorithme de Metropolis-Hasting (section 2.3.1).

Dans ce cas, le calcul du DIC est donné par les équations suivantes.

$$\begin{cases} DIC = 2\bar{D} - D(\bar{\lambda}_{post}) \\ \bar{D} = \frac{1}{R} \sum_{l=1}^R -2 \log \left((\lambda^{(l)})^n \exp \left(-\lambda^{(l)} \sum_{i=1}^n x_i \right) \right) \\ D(\bar{\lambda}_{post}) = -2 \log \left((\bar{\lambda}_{post})^n \exp \left(-\bar{\lambda}_{post} \sum_{i=1}^n x_i \right) \right) \end{cases} \quad (3.15)$$

$$DIC = -4 \frac{n}{R} \sum_{l=1}^R \log(\lambda^{(l)}) + 2n \log(\bar{\lambda}_{post}) + 2\bar{\lambda}_{post} \sum_{i=1}^n x_i \quad (3.16)$$

Rappelons que le DIC permet de classer plusieurs modèles. En effet, plus le DIC est petit, plus la qualité d'ajustement du modèle est meilleure et par conséquent la distribution a priori sur le paramètre λ est plus vraisemblable.

La deuxième quantité statistique est le p-value bayésien. Le calcul de cette quantité se fait en utilisant l'échantillon obtenu de l'algorithme de Metropolis-Hasting. Pour un modèle donné, plus le p-value bayésien est proche de 0,5, plus la qualité du modèle est bonne. Dans notre cas, la procédure de calcul du p-value bayésien est la suivante :

1. Pour $l = 1$ à R

- Nous simulons un échantillon $x^{rep,l}$ de taille n de la loi exponentielle de paramètre $\lambda^{(l)}$.
 - Nous calculons les deux quantités statistiques $Dev_{obs}(x, \lambda^{(l)})$ et $Dev_{rep}(x^{rep,l}, \lambda^{(l)})$.
- Les expressions des deux quantités sont présentées dans le système d'équations 3.17.

$$\begin{cases} Dev_{obs}(x, \lambda^{(l)}) = -2 \sum_{i=1}^n \text{Log}(\lambda^{(l)} \exp(-\lambda^{(l)} x_i)) \\ Dev_{rep}(x^{rep,l}, \lambda^{(l)}) = -2 \sum_{i=1}^n \text{Log}(\lambda^{(l)} \exp(-\lambda^{(l)} x_i^{rep,l})) \end{cases} \quad (3.17)$$

2. Nous calculons le p-value bayésien à l'aide de l'équation 3.18 :

$$p\text{-value} = \frac{1}{R} \sum_{l=1}^R \mathbb{1}_{\{Dev_{rep}(x^{rep,l}, \lambda^{(l)}) \geq Dev_{obs}(x, \lambda^{(l)})\}} \quad (3.18)$$

Nous avons procédé au calcul des statistiques des distributions a postérieur pour chaque modèle. Les résultats sont présentés dans le tableau 3.3.

Le tableau 3.4 donne la fonction de masse a postérieur de la variable aléatoire J . Ces résultats montrent les probabilités $P_{post}(J = j)$ obtenues pour le modèle 3 représentant ainsi l'agrégation des avis d'experts et qui est le seul modèle à faire intervenir les deux avis d'experts.

Tableau 3.3 Statistiques des distributions a post riori des 3 mod les

Mod�le	$\bar{\lambda}_{post}$	σ_{post}	DIC	$p-value$
1	0,4041	0,0353	155,3442	0,5276
2	0,4855	0,0598	157,1206	0,1993
3	0,4137	0,0486	155,7683	0,4647

Tableau 3.4 Fonction de masse a post riori de la variable I

Expert (j)	$P_{post}(J = j)$
1	0,8673
2	0,1327

La figure 3.1 donne les distributions a priori de chaque expert et la distribution a post riori de λ pour le mod le 3. Elle donne aussi le d roulement de l' chantillonnage pour ce mod le.

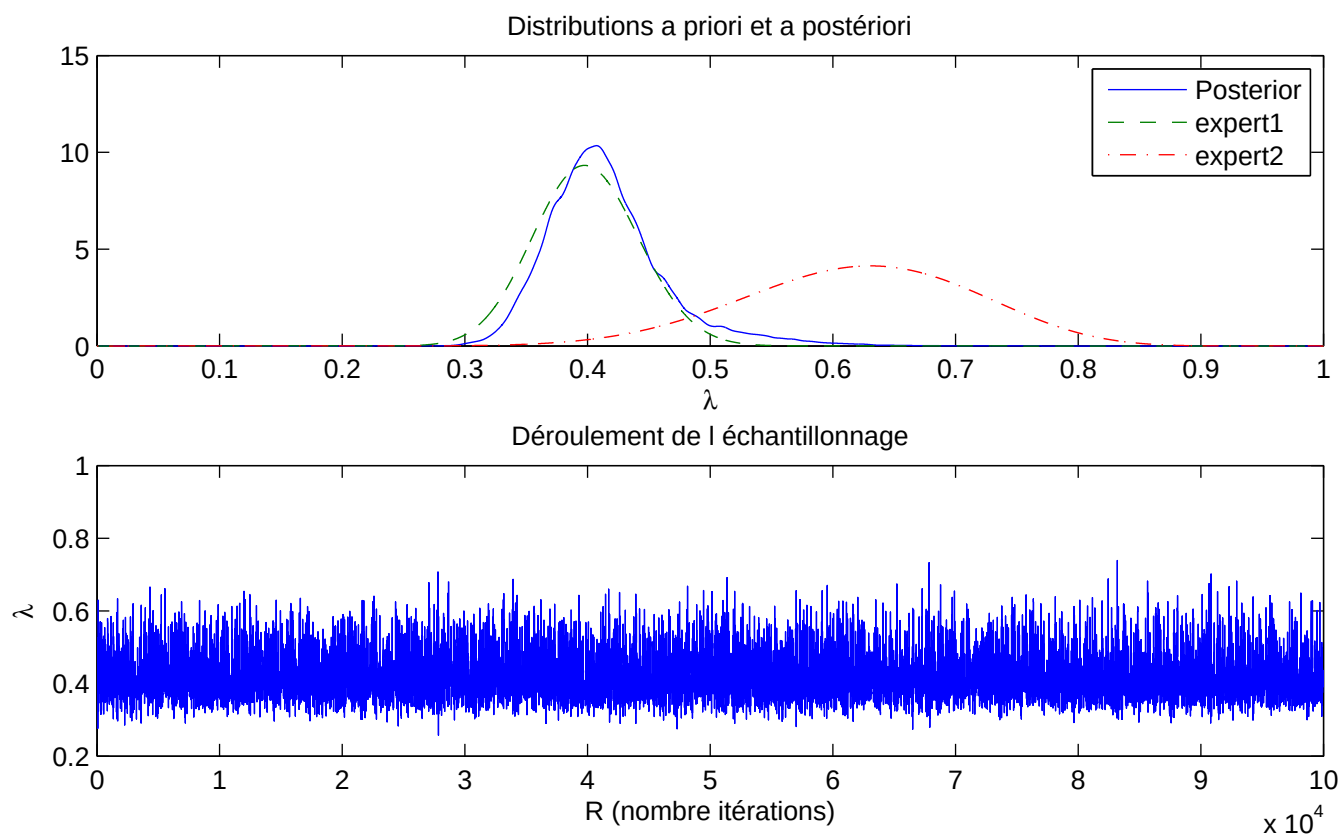


Figure 3.1 Distributions a priori des deux experts et distribution a post riori du mod le 3

Nous pouvons proc der   l'analyse des r sultats. Nous remarquons que le mod le 1 est le meilleur en terme de qualit  d'ajustement. En effet, si nous trions les mod les suivant le

critère du *DIC*, nous aurons le modèle 1 suivi du modèle 3 et puis finalement le modèle 2. Le même classement s'applique pour le cas du *p-value* bayésien. En fait, le modèle 1 demeure le meilleur vu qu'il possède le *p-value* le plus proche de 0,5. Il est tout à fait logique de trouver que le modèle 3 est moyen pour la qualité d'ajustement puisque ce modèle prend en compte les deux avis. Selon le tableau 3.4, l'expert 1 est le plus performant car il possède la probabilité a postériori la plus élevée soit 0,87. Cependant, celle de l'expert 2 est à 0,13. Ce résultat nous donne une idée bien claire sur la performance de chaque expert. Le modèle 3 permet donc de trier les distributions a priori provenant des deux experts en fonction de ces probabilités a postériori. Ce tri concorde bien avec celui du *DIC* et du *p-value*. Nous pouvons remarquer que le modèle 3 permet bien de classer les avis des experts puisqu'il en tient compte sans avoir à procéder à l'inférence bayésienne pour chaque expert seul.

3.3 Comparaison des modèles de combinaison d'avis d'experts

Dans cette section, nous allons comparer le modèle d'agrégation bayésienne avec le modèle du mélange et celui de la moyenne. Ces deux modèles de combinaison des avis d'experts ont été présentés dans la littérature (section 2.4). En se basant sur les notations de la section 3.2.2, la distribution a priori du modèle du mélange est donnée par l'équation 3.19. Pour le modèle de la moyenne, la distribution a priori est donnée par l'équation 3.20. Nous utilisons une seule période pour procéder à la comparaison des 3 modèles. L'estimateur du maximum de vraisemblance de λ de la période concernée est égal à 0,4121. Cet estimateur servira de référence pour placer les avis d'experts à droite ou à gauche de cette valeur. En effet, le but de la comparaison que nous faisons est de voir si chaque modèle de combinaison permet de détecter le bon avis d'expert. Par conséquent, nous devons être capable de placer chaque avis d'expert à droite, à gauche ou de part et d'autre de la vraisemblance.

$$\Pi_{\Lambda}(\lambda) = \frac{1}{N} \sum_{i=1}^N \Pi_i(\lambda) \quad (3.19)$$

$$\left\{ \begin{array}{l} \Pi_{\Lambda}(\lambda) = \frac{\Gamma(a+b)}{\Gamma(a)\Gamma(b)} \lambda^{a-1} (1-\lambda)^{b-1} \mathbb{1}_{[0,1]}(\lambda) \\ a = \frac{1}{N} \sum_{i=1}^N a_i \\ b = \frac{1}{N} \sum_{i=1}^N b_i \end{array} \right. \quad (3.20)$$

Pour illustrer ces propos, nous allons prendre en considération deux avis d'experts dont l'un est qualifié de 'bon' et l'autre est qualifié de 'mauvais'. Nous allons procéder à l'inférence

bayésienne pour les 3 modèles de combinaisons des avis d'experts. Le tableau 3.5 donne les avis des deux experts.

Tableau 3.5 Modélisation des 2 avis d'experts

Expert (j)	$\lambda_{\min,j}$	$\lambda_{\max,j}$	a_j	b_j
1	0,33	0,47	52,4421	78,9513
2	0,1	0,3	7,7191	32,6473

Les tableaux 3.6 et 3.7 donnent respectivement les résultats a posteriori des 3 modèles ainsi que les probabilités a posteriori du modèle d'agrégation bayésienne.

Tableau 3.6 Résultats a posteriori de la comparaison

Modèle	$\bar{\lambda}_{post}$	σ_{post}	DIC	$p-value$
Modèle du mélange	0,3812	0,1020	161,4052	0,6696
Modèle d'agrégation bayésienne	0,4014	0,0380	155,4879	0,5410
Modèle de la moyenne	0,3770	0,0387	155,8847	0,5864

Tableau 3.7 Probabilités a posteriori du modèle d'agrégation bayésienne

Expert (j)	$P_{post}(J = j)$
1	0,9677
2	0,0323

Selon le tableau 3.6, le modèle d'agrégation bayésienne est le meilleur puisqu'il possède le DIC le moins élevé et le $p-value$ le plus proche de 0,5. Ce modèle parvient à détecter que l'expert 1 est plus performant. En effet, il accorde la probabilité a posteriori la plus élevée à cet expert comme le montre le tableau 3.7 soit 0,9677. Pour les modèles du mélange et de la moyenne, ils sont influencés par l'avis de l'expert 2. Le modèle du mélange ne parvient à détecter que l'expert 1 est meilleur que l'expert 2. En effet, l'écart-type sur λ est élevé et le DIC est plus grand que celui du modèle d'agrégation bayésienne. La figure 3.2 donne les distributions a posteriori des 3 modèles (mélange, moyenne et agrégation bayésienne) ainsi que les distributions a priori pour les deux experts.

3.4 Validation du modèle d'agrégation bayésienne avec la tendance des données

Dans cette section, nous allons valider le modèle d'agrégation bayésienne proposé à l'aide du test non paramétrique de tendance des données. Nous considérons k périodes pour les-

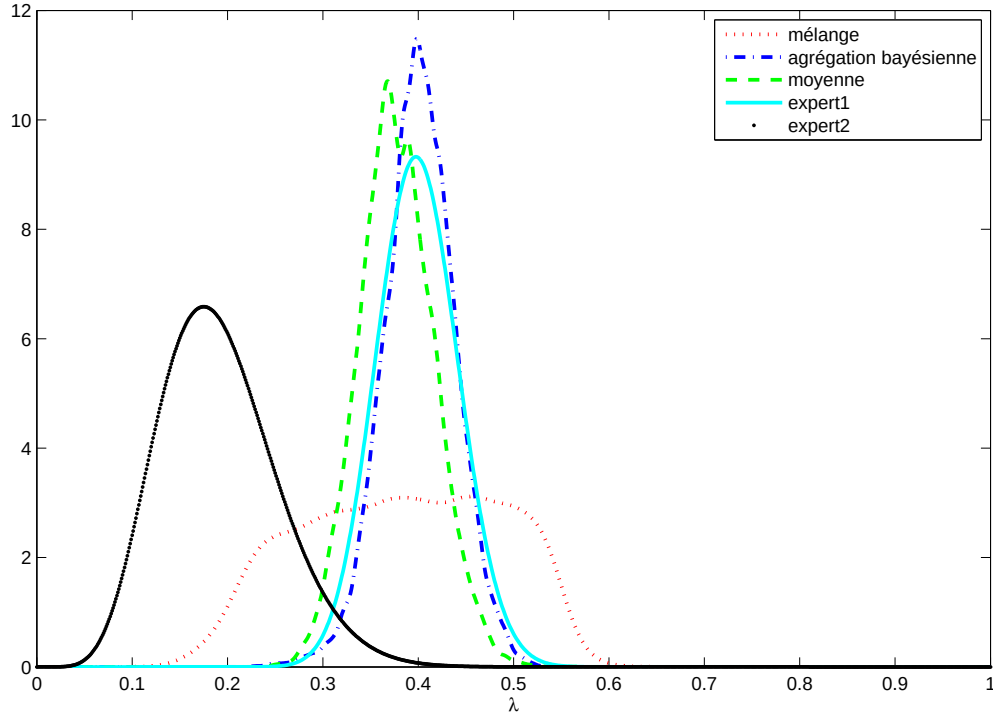


Figure 3.2 Distributions a priori des deux experts et distributions a postérieure des 3 modèles

quelles nous avons enregistré les temps de bon fonctionnement de l'équipement. Nous supposons que pour chaque période k , les durées de bon fonctionnement suivent une loi exponentielle de paramètre λ . Pour chaque période k , nous notons par :

- $x_{i,k}$: les durées de bon fonctionnement de l'équipement.
- n_k : le nombre de défaillances enregistrées durant la période k .
- $X_k = \sum_{i=1}^{n_k} x_{i,k}$: la durée totale de bon fonctionnement durant la période k .

Nous allons étudier la tendance du taux du défaut de l'équipement entre chaque deux périodes consécutives. Le test non paramétrique de *Fisher* va nous permettre de juger si le taux de défaut augmente, diminue ou reste constant entre deux périodes successives. Dans ce qui suit, nous présentons le test de *Fisher* puis nous allons considérer les avis d'experts afin de valider notre modèle d'agrégation.

3.4.1 Analyse de la tendance des données de vraisemblance

Le test de *Fisher* est proposé pour analyser la tendance du taux de défaut entre deux périodes consécutives (Lehmann et D'Abrera (2006)). L'hypothèse nulle H_0 de ce test est que

le taux de défaillance reste constant en passant d'une période k vers une période $k + 1$. La procédure du test est la suivante :

1. Nous trions les durées de bon fonctionnement de l'équipement dans l'ordre croissant pour les périodes k et $k + 1$ et nous notons par y^k et y^{k+1} les deux échantillons triés.
2. Nous calculons la quantité statistique Q selon l'équation 3.21 :

$$Q(n_k, n_{k+1}) = \frac{\frac{1}{n_k} \sum_{i=1}^{n_k} (y_{i+1}^k - y_i^k)}{\frac{1}{n_{k+1}} \sum_{i=1}^{n_{k+1}} (y_{i+1}^{k+1} - y_i^{k+1})} \quad (3.21)$$

3. Nous décidons si le taux de défaillance reste constant, augmente ou diminue. Sous H_0 , $Q(n_k, n_{k+1})$ suit une loi de *Fisher* à $2n_k$ et $2n_{k+1}$ degrés de liberté. Le critère de décision pour un seuil de signification α donné est le suivant :
 - Si $Q(n_k, n_{k+1}) > f_{(1-\frac{\alpha}{2})}$, nous rejetons H_0 en faveur de l'alternative '*le taux de défaillance augmente*'.
 - Si $Q(n_k, n_{k+1}) < f_{(\frac{\alpha}{2})}$, nous rejetons H_0 en faveur de l'alternative '*le taux de défaillance diminue*'.
 - Si $f_{(\frac{\alpha}{2})} < Q(n_k, n_{k+1}) < f_{(1-\frac{\alpha}{2})}$, nous ne rejetons pas H_0 .

Exemple :

Le tableau 3.8 donne les résultats du test de *Fisher* pour un exemple prenant en considération 6 périodes de fonctionnement de l'équipement. Les données de durées de vie période par période sont obtenues par simulation.

Tableau 3.8 Résultats du test de *Fisher* pour un seuil de signification $\alpha = 0,05$

période k - période $k + 1$	$Q(n_k, n_{k+1})$	$f_{0,025}$	$f_{0,975}$	Tendance du taux de défaillance λ
1-2	1,2831	0,6493	1,5367	reste constant
2-3	1,2315	0,6508	1,5401	reste constant
3-4	1,6442	0,6633	1,4875	augmente
4-5	1,0979	0,6957	1,4333	reste constant
5-6	0,5910	0,6810	1,4892	diminue

3.4.2 Analyse des avis d'experts selon la tendance des données

Nous allons valider le modèle d'agrégation d'avis d'experts à l'aide des résultats obtenus avec le test de *Fisher* décrit dans la section précédente. Nous considérons 3 experts indépendants qui vont exprimer leurs avis sur le taux de défaillance. Au début de chaque période k ,

chaque expert se réfère au taux de défaillance de l'année précédente pour donner son avis. Nous notons par $\hat{\lambda}_{k-1}$ le taux de référence par rapport auquel les experts se basent pour formuler leurs avis pour la période k . Les avis d'experts sont résumés comme suit :

- Expert 1 : prévoit que le taux de défaillance restera constant à la période k par rapport à la période $k - 1$.
- Expert 2 : prévoit une diminution du taux de défaillance au cours de la période k par rapport à la période $k - 1$.
- Expert 3 : prévoit une augmentation du taux de défaillance au cours de la période k par rapport à la période $k - 1$.

Les avis des 3 experts seront maintenus les mêmes pour les 6 périodes considérées (voir exemple de la section 3.4.1). La modélisation des avis des 3 experts est faite en fixant à chaque année les bornes minimale et maximale de leurs estimations du taux de défaillance. Nous avons choisi de modéliser ces avis comme le montre le tableau 3.9. La détermination des paramètres des distributions a priori pour chaque expert est faite en solvant le système d'équations 3.3.

Tableau 3.9 Modélisation des 3 avis d'experts

Expert (j)	$\lambda_{\min,j}$	$\lambda_{\max,j}$
1	$\frac{5}{6}\hat{\lambda}_{k-1}$	$\frac{7}{6}\hat{\lambda}_{k-1}$
2	$\frac{1}{2}\hat{\lambda}_{k-1}$	$\hat{\lambda}_{k-1}$
3	$\hat{\lambda}_{k-1}$	$\frac{3}{2}\hat{\lambda}_{k-1}$

Nous avons procédé à l'inférence bayésienne pour chacune des 6 périodes en utilisant le modèle d'agrégation bayésienne. Le tableau 3.10 résume les probabilités a postériori que chaque expert soit le plus performant pour chaque période. La figure 3.3 visualise les résultats du tableau 3.10.

Tableau 3.10 Validation de la tendance des données avec les probabilités a postériori

période k - période $k + 1$	Tendance du taux de défaillance λ	$P_{post}(J = j)$		
		$J = 1$	$J = 2$	$J = 3$
1-2	reste constant	0,5490	0,1456	0,3055
2-3	reste constant	0,5328	0,2645	0,2028
3-4	augmente	0,1632	0,0190	0,8178
4-5	reste constant	0,5103	0,0965	0,3932
5-6	diminue	0,1529	0,8337	0,0134

Le tableau 3.10 nous permet de conclure que la méthode d'agrégation bayésienne permet de privilégier l'expert dont l'avis concorde mieux avec la tendance des données (validée à l'aide du test de *Fisher*). À titre d'exemple, la tendance du taux de panne entre les périodes 3 et 4 est à l'augmentation, la méthode privilège l'expert 3 qui prévoit une augmentation du taux de défaillance. En effet, cet expert possède la probabilité a postériori (0.8178) la plus élevée.

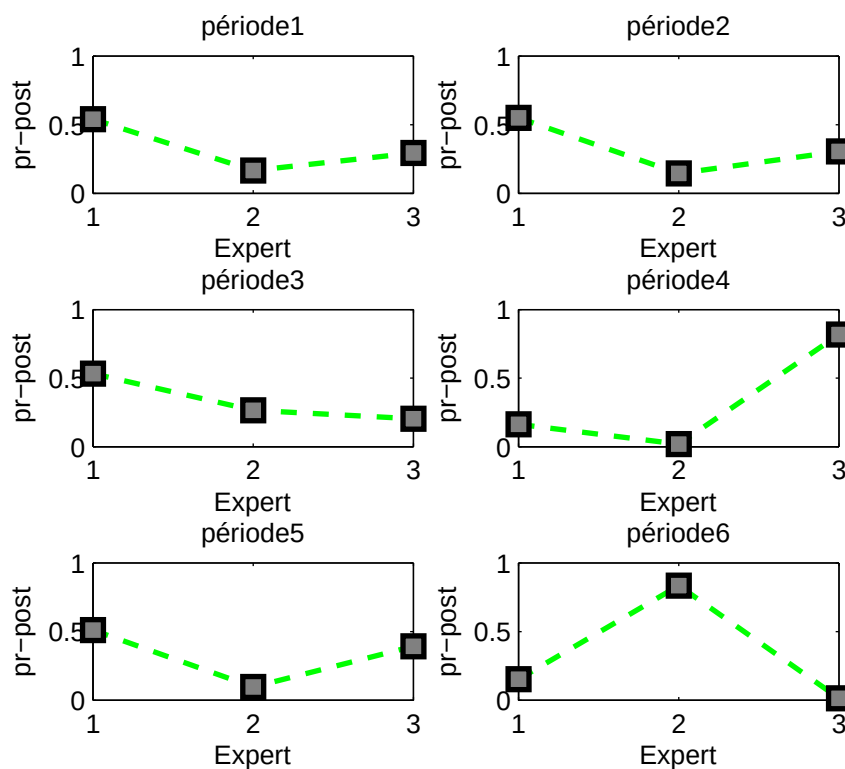


Figure 3.3 Probabilités a postériori des 3 experts

CHAPITRE 4

OPTIMISATION DE LA DISPONIBILITÉ D'UN ÉQUIPEMENT

Dans ce chapitre, nous allons présenter un modèle qui permet d'optimiser la disponibilité d'un équipement. Nous commençons, dans un premier temps, par présenter la stratégie de maintenance que nous allons utiliser telle que traitée dans la littérature. Ensuite, nous allons adapter cette stratégie à l'optimisation de la disponibilité. Ensuite, nous allons proposer un modèle d'agrégation d'avis d'experts sur les paramètres inconnus de la stratégie de maintenance : le taux de défaillance et le taux de réparation. Ce modèle servira pour calculer la distribution a posteriori de chaque paramètre et implémenter la stratégie de maintenance proposée.

4.1 Stratégie d'optimisation de la disponibilité

La stratégie de maintenance choisie est celle qui maximise la disponibilité d'un équipement. Nous expliquons dans un premier temps le principe de la stratégie tel que présentée dans la littérature puis nous allons l'adapter à notre cas d'étude afin qu'elle puisse prendre en compte l'actualisation du taux de défaillance et de réparation selon le modèle d'agrégation proposé.

4.1.1 Description de la stratégie

Nous considérons un équipement soumis à une réparation minimale à chaque défaillance et à une rénovation majeure à chaque période t_p . La stratégie optimale consiste à trouver l'intervalle t_p^* qui maximise la disponibilité de l'équipement ou, d'une manière équivalente, minimise son indisponibilité. En effet, plus le nombre de rénovations majeures augmente, plus le temps d'arrêt de l'équipement augmente suite à ces rénovations. En contre partie, le nombre d'arrêts suite à des réparations minimales de l'équipement diminue si la fréquence de rénovations majeures augmente. Nous supposons que les réparations minimales remettent l'équipement dans un état opérationnel sans toutefois augmenter sa performance. Cela suppose qu'il est remis à son état juste avant sa défaillance.

Nous allons adopter les notations suivantes :

1. T_f : Temps moyen nécessaire pour faire une réparation minimale de l'équipement.
2. T_p : Temps moyen nécessaire pour faire une rénovation majeure de l'équipement.

3. $\lambda(t)$: Le taux de défaillance de l'équipement.

Le but de la stratégie est de déterminer l'intervalle optimal entre deux rénovations majeures permettant de maximiser la disponibilité de l'équipement. La figure 4.1 illustre cette stratégie.

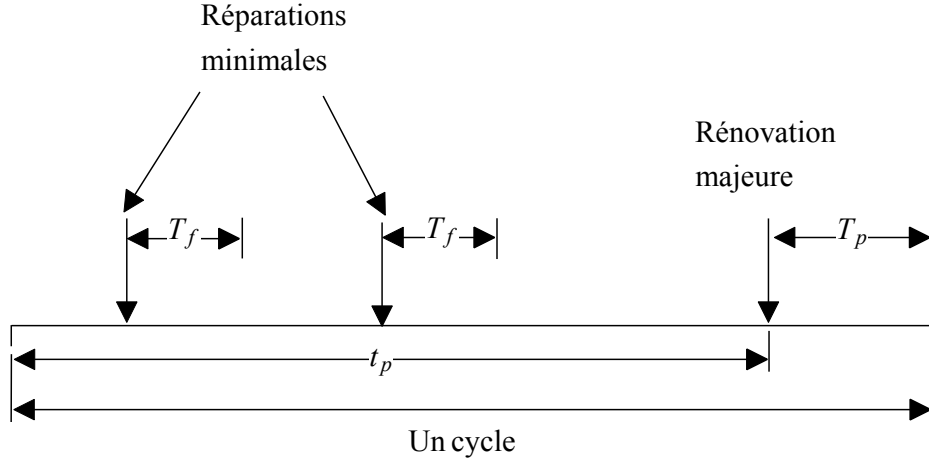


Figure 4.1 Maximisation de la disponibilité

Les seules différences avec la formulation de la stratégie minimisant les coûts (Barlow et Proschan (1996)) sont de :

- Remplacer le coût de réparation minimale C_{mr} par le temps alloué à cette intervention.
- Remplacer le coût de remplacement préventif C_p par la durée de rénovation majeure.

Le temps d'arrêt total par unité de temps, pour un intervalle de longueur t_p entre deux rénovations majeures, est noté $D(t_p)$. $D(t_p)$ représente aussi l'indisponibilité instantanée par unité de temps. Elle est donnée par l'équation 4.1.

$$D(t_p) = \frac{H(t_p) \cdot T_f + T_p}{t_p + T_p}. \quad (4.1)$$

$H(t_p) \cdot T_f$ désigne le temps d'arrêt de l'équipement suite aux défaillances. Il est égal au produit du nombre de défaillances dans l'intervalle $[0, t_p]$ par le temps moyen nécessaire pour faire une réparation minimale. Le calcul du nombre de défaillances dans l'intervalle $[0, t_p]$ est donné par l'équation 4.2.

$$H(t_p) = \int_0^{t_p} \lambda(t) dt. \quad (4.2)$$

Nous cherchons à déterminer t_p^* qui minimise la fonction $D(t_p)$. Nous commençons par dériver cette fonction :

$$D'(t_p) = \frac{T_f \cdot [\lambda(t_p)(t_p + T_p) - H(t_p)] - T_p}{(t_p + T_p)^2}$$

Nous cherchons par la suite à annuler cette dérivée. Soit la fonction $\Phi(t_p)$:

$$\Phi(t_p) = \lambda(t_p)(t_p + T_p) - H(t_p)$$

Nous avons donc : $D'(t_p^*) = 0 \implies \Phi(t_p^*) = \frac{T_p}{T_f}$. Comme $\Phi(0) < \frac{T_p}{T_f}$, il suffit que la fonction Φ soit strictement croissante pour que nous puissions trouver un unique t_p^* tel que $\Phi(t_p^*) = \frac{T_p}{T_f}$. Nous avons par conséquent :

$$\Phi'(t_p) = \lambda'(t_p)(t_p + T_p) > 0 \implies \lambda'(t_p) > 0$$

Une condition suffisante de trouver un minimum unique de la fonction $D(t_p)$ est d'avoir un taux de défaillance strictement croissant. Nous devons ensuite vérifier qu'il s'agit bien d'un minimum, il faut calculer la dérivée seconde de la fonction $D(t_p)$.

$$D''(t_p) = \frac{T_f \cdot (t_p + T_p)^2 \cdot \Phi'(t_p) - 2(t_p + T_p)(T_f \cdot \Phi(t_p) - T_p)}{(t_p + T_p)^4}$$

Comme $D''(t_p^*) = \frac{T_f \cdot (t_p^* + T_p)^2 \cdot \Phi'(t_p^*)}{(t_p^* + T_p)^4} > 0$, alors nous pouvons conclure que t_p^* est un minimum de la fonction $D(t_p)$.

4.1.2 Stratégie de disponibilité prenant en compte les avis d'experts

Dans cette partie, nous allons adapter la stratégie de maintenance présentée dans la section précédente à notre cas d'étude. Nous supposons que nous avons un équipement neuf qui est mis en fonctionnement à l'instant $t = 0$. Le fournisseur de l'équipement nous fournit un taux de défaillance de base de l'équipement $\lambda_0 = 0,4$ pannes/semaine et un taux de réparation de base $\mu_0 = 10$ réparations/semaine. Nous supposons également que les durées de bon fonctionnement de l'équipement ainsi que les durées de réparation suivent respectivement des lois exponentielles de paramètres λ et μ et que ses durées sont indépendantes.

Au début de chaque année (52 semaines), des experts en fiabilité et en maintenance donnent respectivement leurs avis sur les paramètres inconnus λ et μ . L'objectif est de trouver une stratégie optimale de maintenance permettant de maximiser la disponibilité de l'équi-

pement. Cette stratégie consiste à déterminer le temps optimal à partir duquel il faut faire une rénovation majeure de l'équipement. La stratégie décrite dans la section précédente sera modifiée pour tenir compte des avis d'experts. En effet, le modèle de l'équation 4.1 suppose que la forme du taux de défaillance est connue.

Dans notre cas d'étude, la seule information que nous possédons sur le taux de défaillance est sa moyenne a posteriori que nous pouvons déterminer en combinant les avis d'experts avec la vraisemblance de chaque année. Il va falloir donc ajuster un modèle afin de prédire l'allure du taux de défaillance à partir de ces estimations a posteriori du taux de défaillance. Il est à noter que le taux de défaillance est constant par période d'observation (processus de Poisson homogène). Par ailleurs, ce taux de défaillance peut varier d'une période d'observation à une autre (processus de Poisson non homogène). La figure 4.2 explique la stratégie proposée d'optimisation de la disponibilité.

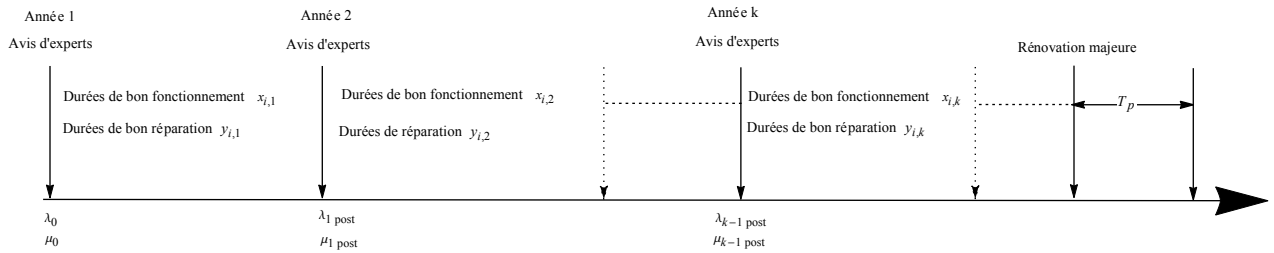


Figure 4.2 Stratégie de disponibilité et introduction d'avis d'experts

Nous allons adopter les notations suivantes dans la suite de la modélisation :

- k : Indice désignant l'année de l'étude ($k \in \{1, \dots, K\}$).
- n_k : Le nombre de défaillances durant l'année k .
- $x_{i,k}$: Les durées de bon fonctionnement durant l'année k ($i \in \{1, \dots, n_k\}$).
- $y_{i,k}$: Les durées de réparation durant l'année k ($i \in \{1, \dots, n_k\}$).
- h : La longueur de l'année ($h=52$ semaines). Nous avons $h \simeq \sum_{i=1}^{n_k} x_{i,k} + y_{i,k} \quad \forall k \in \{1, \dots, K\}$.
- $\bar{\lambda}_{post,k}$: la moyenne a posteriori du taux de défaillance de l'année k .
- $\bar{\mu}_{post,k}$: la moyenne a posteriori du taux de réparation de l'année k .
- T_p : La durée nécessaire pour faire un rénovation majeure de l'équipement. Nous supposons que cette durée est égale à 8 semaines.

L'optimisation de la disponibilité de l'équipement se fait au début de chaque année. En effet, l'information sur le taux de défaillance de l'équipement doit être mise à jour au début de chaque année pour procéder à la maximisation de la disponibilité. Cette mise à jour du

taux de défaillance est due aux avis d'experts. La forme du taux de défaillance (constant, linéaire, non linéaire) prise en compte dans l'optimisation de la disponibilité dépend de l'année d'optimisation. La démarche d'optimisation de la disponibilité de l'équipement est la suivante :

Début de l'année 1

Nous allons prendre les avis d'experts en fiabilité et en maintenance sur les paramètres λ et μ . Nous supposons que le taux de défaillance serait constant et égal à λ_0 pour le reste de la vie de l'équipement ($\lambda_1(t) = \lambda_0$). En effet, au début de la première année, c'est la seule information que nous possédons sur l'allure du taux de défaillance. Nous allons ensuite calculer la durée optimale $t_{p,1}^*$ à partir de laquelle il faut faire une rénovation majeure de l'équipement en utilisant le modèle de la section précédente (équation 4.1). Le temps d'arrêt total de l'équipement par unité de temps sur une longueur infinie est construit comme suit :

$$\begin{aligned} D_1(t_p) &= \frac{\frac{1}{\mu_0} \cdot \int_0^{t_p} \lambda_1(t) dt + T_p}{t_p + T_p} \\ &= \frac{\frac{\lambda_0}{\mu_0} \cdot t_p + T_p}{t_p + T_p} \end{aligned}$$

Nous remarquons que la fonction $D_1(t_p)$ est strictement décroissante donc $t_{p,1}^* = +\infty$. Ceci veut dire qu'il n'est pas rentable de faire une rénovation majeure de l'équipement (aucune amélioration de la disponibilité). La disponibilité optimale sera égale à la limite de $1 - D_1(t_p)$ lorsque t_p tend vers $+\infty$ donc égale à $1 - \frac{\lambda_0}{\mu_0}$.

Début de l'année 2

Nous prenons les avis d'experts sur les paramètres λ et μ . Nous calculons $\bar{\lambda}_{post,1}$ et $\bar{\mu}_{post,1}$ en utilisant les données de bons fonctionnement, de réparations et les avis d'experts pris au début de la première année. La procédure de calcul des statistiques a posteriori sur λ et μ sera traitée en détail dans la section qui suit. Nous supposons que le taux de défaillance sera linéaire à partir de la deuxième année jusqu'à la rénovation majeure de l'équipement ($\lambda_2(t) = At + B$). La détermination des constantes A et B se fait utilisant les deux conditions suivantes :

$$\begin{cases} \lambda(t=0) = \lambda_0 \\ \lambda(t=h) = \bar{\lambda}_{post,1} \end{cases}$$

L'expression du taux de défaillance sera donc donnée par l'équation suivante :

$$\lambda_2(t) = \lambda_0 + \frac{\bar{\lambda}_{post,1} - \lambda_0}{h} \cdot t \quad \forall t > h$$

Le temps total d'arrêt de l'équipement par unité de temps sera donné par :

$$D_2(t_p) = \frac{\frac{1}{\mu_0} \cdot \int_0^h \lambda_1(t) dt + \frac{1}{\mu_1} \cdot \int_h^{t_p} \lambda_2(t) dt + T_p}{t_p + T_p}$$

Nous cherchons ensuite à minimiser le temps d'arrêt total pour trouver $t_{p,2}^*$ ainsi que la disponibilité optimale correspondante.

Début de l'année k ($k \geq 3$)

Nous prenons les avis d'experts sur les paramètres λ et μ . Nous évaluons ensuite $\bar{\lambda}_{post,k-1}$ et $\bar{\mu}_{post,k-1}$ en utilisant les données de bons fonctionnement, de réparations et les avis d'experts formulés au début de l'année $k-1$. Nous supposons que le taux de défaillance aura la forme $\lambda_k(t) = a_k + b_k \cdot t^{c_k}$ à partir de l'année k ($k \geq 3$). L'ajustement du taux de défaillance avec la forme précédente est fait en utilisant les conditions suivantes :

$$\left\{ \begin{array}{l} \lambda(t=0) = \lambda_0 \\ \lambda(t=h) = \bar{\lambda}_{post,1} \\ \lambda(t=2h) = \bar{\lambda}_{post,2} \\ \vdots \\ \lambda(t=(k-1)h) = \bar{\lambda}_{post,k-1} \end{array} \right.$$

Le temps total d'arrêt de l'équipement par unité de temps sera sous la forme suivante :

$$D_k(t_p) = \frac{\sum_{s=1}^{k-1} \left(\frac{1}{\mu_{s-1}} \cdot \int_{(s-1) \cdot h}^{s \cdot h} \lambda_s(t) dt \right) + \frac{1}{\mu_{k-1}} \cdot \int_{(k-1) \cdot h}^{t_p} \lambda_k(t) dt + T_p}{t_p + T_p}$$

Nous minimisons ensuite la fonction $D_k(t_p)$ et nous cherchons $t_{p,k}^*$ ainsi que la disponibilité correspondante.

Dans la stratégie présentée précédemment, le calcul du temps optimal de rénovation majeure est effectué au début de chaque année. En effet, l'information sur les taux de défaillance

et de réparation est actualisée au début de chaque année ce qui nous amène à procéder de nouveau au calcul du temps optimal de la rénovation. Il reste maintenant à présenter la façon avec laquelle seront calculées les statistiques a posteriori sur les paramètres λ et μ . Nous allons nous baser sur les données de vraisemblance caractérisant la disponibilité et les avis d'experts formulés sur les taux de défaillance et de réparation.

4.2 Modèle de vraisemblance associé à la disponibilité

Nous supposons dans cette section que pendant une année donnée, l'équipement en exploitation est sujet à des défaillances. Après chaque défaillance, l'équipement est réparé minimalement et est remis en service. Après chaque réparation minimale, l'équipement est aussi bon que vieux, son taux de défaillance reste inchangé. Ainsi, pour une année donnée, nous aurons un certain nombre de défaillances et un nombre équivalent de réparations minimales. La figure 4.3 donne un cycle de bon fonctionnement suivi d'une réparation minimale suite à une défaillance ainsi que les variables aléatoires X et Y caractérisant respectivement la durée de bon fonctionnement et la durée de réparation minimale.

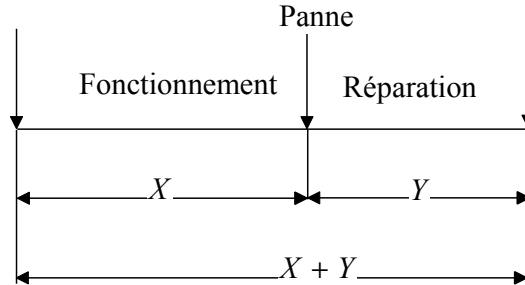


Figure 4.3 Cycle de fonctionnement de l'équipement

Nous supposons que, pendant une période donnée, nous avons n durées de bon fonctionnement x_i ($i \in \{1, \dots, n\}$) ainsi que les durées de réparation y_i ($i \in \{1, \dots, n\}$) correspondantes. Nous pouvons donc calculer la disponibilité d_i de l'équipement pour chaque cycle constitué d'un fonctionnement suivi d'une réparation. Les données d_i serviront pour constituer la fonction de vraisemblance sur laquelle nous allons effectuer une inférence bayésienne. Le calcul des données de disponibilité est donné par l'équation 4.3 .

$$d_i = \frac{x_i}{x_i + y_i} \quad \forall i \in \{1, \dots, n\} \quad (4.3)$$

Dans ce qui suit, nous allons introduire les 3 variables aléatoires continues suivantes :

- X : désigne le temps de bon fonctionnement de l'équipement. Nous supposons qu'elle suit une loi exponentielle de paramètre λ ($X \sim \text{Exp}(\lambda)$).
- Y : désigne le temps de réparation de l'équipement. Nous supposons qu'elle suit une loi exponentielle de paramètre μ ($Y \sim \text{Exp}(\mu)$).
- Z : désigne la disponibilité de l'équipement. Z est définie comme la proportion du temps de bon fonctionnement par la durée totale d'un cycle complet ($Z = \frac{X}{X+Y}$).

Les fonctions de densité des variables X et Y sont données respectivement par le système d'équations 4.4. Nous supposons que ces deux variables sont **indépendantes**.

$$\begin{cases} f_X(x) = \lambda \exp(-\lambda x) \\ f_Y(y) = \mu \exp(-\mu y) \end{cases} \quad (4.4)$$

Nous devons déterminer la fonction de densité de la variable Z en fonction des paramètres λ et μ . Nous commençons par déterminer la fonction de répartition de la variable Z :

- Si $z \leq 0$: $F_Z(z) = P(Z \leq z) = 0$
- Si $z \geq 1$: $F_Z(z) = P(Z \leq z) = 1$
- Si $0 < z < 1$:

$$\begin{aligned} F_Z(z) &= P(Z \leq z) \\ &= P\left(\frac{X}{X+Y} \leq z\right) \\ &= P\left(\frac{1-z}{z}X \leq Y\right) \\ &= \int_0^{+\infty} \int_{\frac{(1-z)x}{z}}^{+\infty} \lambda \mu \exp(-\lambda x - \mu y) dy dx \\ &= \lambda \mu \int_0^{+\infty} \left[\frac{-1}{\mu} \exp(-\lambda x - \mu y) \right]_{\frac{(1-z)x}{z}}^{+\infty} dx \\ &= \lambda \int_0^{+\infty} \exp\left[-\left(\lambda + \mu \frac{1-z}{z}\right)x\right] dx \\ &= \lambda \left[\frac{1}{-\left(\lambda + \mu \frac{1-z}{z}\right)} \exp\left[-\left(\lambda + \mu \frac{1-z}{z}\right)x\right] \right]_0^{+\infty} \\ &= \frac{\lambda z}{\lambda z + \mu - \mu z} \end{aligned}$$

La fonction de répartition de la variable Z s'écrit alors comme suit :

$$F_Z(z) = \begin{cases} 0 & \text{si } z \leq 0 \\ \frac{\lambda z}{\lambda z + \mu - \mu z} & \text{si } 0 < z < 1 \\ 1 & \text{si } z \geq 1 \end{cases} \quad (4.5)$$

Pour déterminer la fonction de densité de la variable Z , il suffit de dériver sa fonction de répartition :

- Si $z \leq 0$ ou $z \geq 1$: $f_Z(z) = 0$.
- Si $0 < z < 1$:

$$\begin{aligned} f_Z(z) &= \frac{dF_Z(z)}{dz} \\ &= \frac{\lambda\mu}{(\lambda z + \mu - \mu z)^2} \end{aligned}$$

En définitive, la fonction de densité de la variable Z est donnée par l'équation 4.6. Cette densité va servir dans la suite pour construire la fonction de vraisemblance sur les données de disponibilité.

$$f_Z(z) = \frac{\lambda\mu}{(\lambda z + \mu - \mu z)^2} \quad \text{si } 0 < z < 1 \quad (4.6)$$

4.3 Distributions a priori sur λ et μ

Au début de l'année d'intérêt, deux groupes d'experts seront consultés pour exprimer leurs avis sur les taux de panne λ et de réparation μ :

- N_1 experts en fiabilité donneront leurs avis sur λ ($N_1 \geq 1$). Chaque expert j en fiabilité donnera son estimation des bornes minimale $\lambda_{min,j}$ et maximale $\lambda_{max,j}$ du taux de panne λ . La détermination de la distribution a priori sur λ sera faite comme expliquée dans le chapitre précédent. Cette distribution a priori sera modélisée par une loi bêta(a_j, b_j) dont les paramètres seront déterminés en solvant le système d'équations 3.3 .
- N_2 experts en maintenance donneront leurs avis sur μ ($N_2 \geq 1$). Chaque expert v en maintenance donnera son estimation des bornes minimale $\mu_{min,v}$ et maximale $\mu_{max,v}$ du taux de panne μ . La distribution a priori sur μ pour chaque expert v sera modélisée par une loi normale ($N(m_v, \sigma_v^2)$). La détermination des paramètres m_v et σ_v^2 est faite en solvant le système à deux équations non linéaires 4.7.

$$\begin{cases} \frac{1}{2} \left(1 + \operatorname{erf} \left(1 + \left(\frac{\mu_{\min,v} - m_v}{\sqrt{2}\sigma_v} \right) \right) \right) = 0.05 \\ \frac{1}{2} \left(1 + \operatorname{erf} \left(1 + \left(\frac{\mu_{\max,v} - m_v}{\sqrt{2}\sigma_v} \right) \right) \right) = 0.95 \end{cases} \quad \forall v \in \{1..N_2\} \quad (4.7)$$

erf désigne la fonction erreur et son expression est donnée par l'équation suivante :

$$\operatorname{erf}(x) = \frac{2}{\sqrt{\pi}} \int_0^x \exp(-t^2) dt \quad (4.8)$$

Nous supposons que les deux groupes d'experts sont indépendants et que les experts au sein de chaque groupe sont aussi indépendants. Nous allons introduire 4 variables aléatoires qui vont servir pour formuler les distributions a priori sur les paramètres inconnus λ et μ de la disponibilité :

- Λ : variable continue désignant le taux de défaillance de l'équipement.
- M : variable continue désignant le taux de réparation de l'équipement.
- J : variable discrète désignant l'expert en fiabilité le plus performant.
- V : variable discrète désignant l'expert en maintenance le plus performant.

Les fonctions de masse des variables J et V sont données respectivement par les équations 4.9 et 4.10 .

$$P(J = j) = \frac{1}{N_1} \quad \forall j \in \{1, \dots, N_1\} \quad (4.9)$$

$$P(V = v) = \frac{1}{N_2} \quad \forall v \in \{1, \dots, N_2\} \quad (4.10)$$

Le même argument utilisé dans le chapitre précédent justifie le fait d'attribuer des probabilités égales à chaque expert d'un même groupe d'être le plus performant. En effet, la performance d'un expert est mesurée par la justesse de son a priori. C'est la vraisemblance, une fois connue, qui va permettre de mesurer cette justesse. La distribution a priori sur λ pour chaque expert j en fiabilité sera exprimée par la fonction de densité conditionnelle de Λ étant donné que $J = j$. L'expression de cette distribution est donnée par l'équation 4.11. La distribution a priori pour chaque expert v en maintenance sera exprimée par la fonction de densité conditionnelle de M étant donné que $V = v$. Son expression est donnée par l'équation 4.12.

$$\Pi_{\Lambda|J}(\lambda|j) = \frac{\Gamma(a_j + b_j)}{\Gamma(a_j)\Gamma(b_j)} \lambda^{a_j-1} (1 - \lambda)^{b_j-1} \mathbb{1}_{[0,1]}(\lambda) \quad \forall j \in \{1, \dots, N_1\} \quad (4.11)$$

$$\Pi_{M|V}(\mu|v) = \frac{1}{\sigma_v \sqrt{2\pi}} \exp \left(-\frac{1}{2} \left(\frac{\mu - m_v}{\sigma_v} \right)^2 \right) \quad \forall v \in \{1, \dots, N_2\} \quad (4.12)$$

Les distributions a priori sur les paramètres λ et μ sont déterminées en appliquant l'axiome des probabilités totales. Leurs expressions sont données par le système 4.13.

$$\begin{cases} \Pi_{\Lambda}(\lambda) = \frac{1}{N_1} \sum_{j=1}^{N_1} \frac{\Gamma(a_j + b_j)}{\Gamma(a_j)\Gamma(b_j)} \lambda^{a_j-1} (1-\lambda)^{b_j-1} \mathbb{1}_{[0,1]}(\lambda) \\ \Pi_M(\mu) = \frac{1}{N_2} \sum_{v=1}^{N_2} \frac{1}{\sigma_v \sqrt{2\pi}} \exp \left(-\frac{1}{2} \left(\frac{\mu - m_v}{\sigma_v} \right)^2 \right) \end{cases} \quad (4.13)$$

4.4 Distributions a postérieur sur λ et μ

Nous notons par Θ le vecteur aléatoire mixte ($\Theta = (\Lambda, M, J, V)$). Nous allons construire la distribution a postérieur sur le paramètre $\theta = (\lambda, \mu, j, v)$. L'inférence bayésienne sera appliquée sur les n données de disponibilités que nous notons dans la suite par d . Le théorème de Bayes permet d'écrire la distribution a postérieur sur le paramètre θ comme suit :

$$\begin{aligned} \Pi_{\Theta|D}(\theta|d) &= \frac{\Pi_{D|\Theta}(d|\theta)\Pi_{\Theta}(\theta)}{\Pi_D(d)} \\ &= \frac{\Pi_{D|\Theta}(d|\theta)\Pi_{\Theta}(\theta)}{\int_0^{+\infty} \Pi_{D|\Theta}(d|\theta)\Pi_{\Theta}(\theta)d\theta} \end{aligned} \quad (4.14)$$

Le dénominateur de l'équation 4.14 est une constante de normalisation et est indépendante de θ . La distribution a postérieur sera proportionnelle alors au numérateur de cette équation. L'équation 4.15 donne la densité a postérieur de θ .

$$\begin{aligned} \Pi_{\Theta|D}(\theta|d) &\propto \Pi_{D|\Theta}(d|\theta)\Pi_{\Theta}(\theta) \\ &\propto \prod_{i=1}^n \left(\frac{\lambda\mu}{(\lambda d_i + \mu - \mu d_i)^2} \right) \Pi_{\Lambda,J}(\lambda, j) \cdot \Pi_{M,V}(\mu, v) \\ &\propto \prod_{i=1}^n \left(\frac{\lambda\mu}{(\lambda d_i + \mu - \mu d_i)^2} \right) P(J = j) \cdot \Pi_{\Lambda|J}(\lambda|j) \cdot P(V = v) \cdot \Pi_{M|V}(\mu|v) \end{aligned} \quad (4.15)$$

L'application de l'algorithme *MCMC* de Metropolis-Hasting va nous permettre de déter-

miner les distributions a posteriori des paramètres λ et μ ainsi que les fonctions de masse des variables J et V . La procédure de cet algorithme appliquée au modèle 4.15 est la suivante :

1. initialiser $[\lambda^{(1)}, \mu^{(1)}, j^{(1)}, v^{(1)}]$
2. Pour $l = 1$ à R
 - Générer $u \sim U_{[0,1]}$
 - Générer $\theta^* \sim q(\theta^*|\theta^{(l)})$
 - Si $u < A$ alors $\theta^{(l+1)} = \theta^*$ sinon $\theta^{(l+1)} = \theta^{(l)}$

A est le ratio d'acceptation, son expression est la suivante :

$$A = \min \left(1, \frac{\Pi_{\Theta|D}(\theta^*|d)q(\theta^{(l)}|\theta^*)}{\Pi_{\Theta|D}(\theta^{(l)}|d)q(\theta^*|\theta^{(l)})} \right) \quad (4.16)$$

La loi instrumentale $q(\theta^*|\theta^{(l)})$ est construite d'une façon analogue au chapitre précédent. En effet, à chaque itération l , nous générons λ^* et μ^* respectivement à partir des lois normales $N(\lambda^{(l)}, \sigma_1^2)$ et $N(\mu^{(l)}, \sigma_2^2)$. De plus, un expert en fiabilité sera choisi aléatoirement parmi les N_1 experts avec une probabilité $\frac{1}{N_1}$. De la même manière, un expert en maintenance sera sélectionné. La probabilité de choisir un expert en maintenance v parmi les N_2 experts disponibles est égale à $\frac{1}{N_2}$. En définitive, la loi instrumentale s'écrit comme le montre l'équation 4.17.

$$\begin{aligned} q(\theta^*|\theta^{(l)}) &= q(\lambda^*, \mu^*, j^*, v^*|\lambda^{(l)}, \mu^{(l)}, j^{(l)}, v^{(l)}) \\ &= q(\lambda^*|\lambda^{(l)}, \mu^{(l)}, j^{(l)}, v^{(l)}) \times \\ &\quad q(\mu^*|\lambda^{(l)}, \mu^{(l)}, j^{(l)}, v^{(l)}) \times \\ &\quad q(j^*|\lambda^{(l)}, \mu^{(l)}, j^{(l)}, v^{(l)}) \times \\ &\quad q(v^*|\lambda^{(l)}, \mu^{(l)}, j^{(l)}, v^{(l)}) \\ &= \frac{1}{2\pi N_1 N_2 \sigma_1 \sigma_2} \exp \left(-\frac{(\lambda - \lambda^{(l)})^2}{2\sigma_1^2} \right) \exp \left(-\frac{(\mu - \mu^{(l)})^2}{2\sigma_2^2} \right). \end{aligned} \quad (4.17)$$

Il est simple à remarquer que $q(\theta^{(l)}|\theta^*) = q(\theta^*|\theta^{(l)})$. Donc, il est possible de simplifier l'équation 4.16 du ration d'acceptation A . Son expression simplifiée est donnée par l'équation 4.18.

$$\begin{aligned}
A &= \min \left(1, \frac{\Pi_{\Theta|D}(\theta^*|d)}{\Pi_{\Theta|D}(\theta^{(l)}|d)} \right) \\
&= \min \left(1, \frac{\prod_{i=1}^n \left(\frac{\lambda^* \mu^*}{(\lambda^* d_i + \mu^* - \mu^* d_i)^2} \right) \cdot \Pi_{\Lambda|J}(\lambda^*|j^*) \cdot \Pi_{M|V}(\mu^*|v^*)}{\prod_{i=1}^n \left(\frac{\lambda^{(l)} \mu^{(l)}}{(\lambda^{(l)} d_i + \mu^{(l)} - \mu^{(l)} d_i)^2} \right) \cdot \Pi_{\Lambda|J}(\lambda^{(l)}|j^{(l)}) \cdot \Pi_{M|V}(\mu^{(l)}|v^{(l)})} \right)
\end{aligned} \tag{4.18}$$

Les statistiques a posteriori des paramètres λ et μ sont données par le système suivant :

$$\left\{ \begin{array}{l} \bar{\lambda}_{post} = \frac{1}{R} \sum_{l=1}^R \lambda^{(l)}. \\ \sigma_{\lambda,post} = \sqrt{\frac{1}{R-1} \sum_{l=1}^R (\lambda^{(l)} - \bar{\lambda}_{post})^2}. \\ \bar{\mu}_{post} = \frac{1}{R} \sum_{l=1}^R \mu^{(l)}. \\ \sigma_{\mu,post} = \sqrt{\frac{1}{R-1} \sum_{l=1}^R (\mu^{(l)} - \bar{\mu}_{post})^2}. \end{array} \right. \tag{4.19}$$

Les fonctions de masse a posteriori des variables J et V sont données par le système 4.20.

$$\left\{ \begin{array}{l} P_{post}(J = j) = \frac{1}{R} \sum_{l=1}^R \mathbb{1}_{\{j^{(l)}=j\}} \quad \forall j \in \{1, \dots, N_1\}. \\ P_{post}(V = v) = \frac{1}{R} \sum_{l=1}^R \mathbb{1}_{\{v^{(l)}=v\}} \quad \forall v \in \{1, \dots, N_2\}. \end{array} \right. \tag{4.20}$$

Les statistiques a posteriori sur les paramètres λ et μ serviront pour développer la stratégie d'optimisation de la disponibilité traitée dans la section 4.1.2.

Nous avons implémenté un programme à l'aide de *Matlab* permettant d'optimiser la disponibilité (Annexe B). Ce programme permet de lire les données de bon fonctionnement et de réparation ainsi que les avis d'experts exprimés sous forme de bornes minimales et maximales sur les paramètres λ et μ pour chacune des années. Il permet aussi d'estimer les distributions a posteriori sur ces paramètres et calcule la durée optimale pour procéder à la rénovation majeure de l'équipement dépendamment de l'année d'évaluation.

4.5 Exemple d'optimisation de la disponibilité

Afin d'illustrer le modèle d'optimisation de la disponibilité proposé, nous avons simulé des durées de bon fonctionnement et de réparation d'un équipement pour 10 périodes (années).

Ensuite, nous avons appliqué la stratégie de maintenance (section 4.1.2) pour maximiser la disponibilité de l'équipement. Des avis d'experts en fiabilité et en maintenance ont été introduits au début de chaque année pour pouvoir appliquer l'inférence bayésienne sur les taux de défaillance et de réparation. Nous avons supposé que le taux de défaillance de base du constructeur de l'équipement est égal à 0,4 panne/semaine et celui de réparation est égal à 10 réparations/semaine. Nous considérons dans cet exemple 3 experts en fiabilité et 2 experts en maintenance. Les étapes de l'optimisation de la disponibilité sont les suivantes :

Étape 1 :

Au début de la première année, les 3 experts en fiabilité et les 2 experts en maintenance donnent leurs avis respectivement sur les taux de défaillance et de réparation. Leurs avis ainsi que les paramètres des distributions a priori sont résumés dans les tableaux 4.1 et 4.2. Ces avis seront combinés à la fin de la première année (début de la deuxième année) avec les données de disponibilité pour procéder à l'inférence bayésienne.

Tableau 4.1 Avis des experts en fiabilité et paramètres des distributions a priori

Expert (j)	$\lambda_{\min,j}$	$\lambda_{\max,j}$	a_j	b_j
1	0,4	0,44	690,83	954,25
2	0,41	0,47	321,31	409,08
3	0,46	0,51	519,76	551,94

Tableau 4.2 Avis des experts en maintenance et paramètres des distributions a priori

Expert (v)	$\mu_{\min,v}$	$\mu_{\max,v}$	m_v	σ_v
1	8	11	9,50	0,91
2	9	13,5	11,25	1,37

Nous supposons, qu'au début de la première année, le taux de défaillance de l'équipement va être constant et égal au taux de base du constructeur (0,4 panne/semaine). En prenant ainsi un taux de défaillance constant, nous avons montré dans la section 4.1.2 qu'il n'est pas rentable de faire une rénovation majeure de l'équipement et que la disponibilité optimale est égale à $1 - \frac{\lambda_0}{\mu_0}$ soit 0,96.

Étape 2 :

Au début de la deuxième année (fin de la première année), une vraisemblance sur les données de bon fonctionnement et de réparation est constituée. Nous pouvons ainsi construire

les données de disponibilité en utilisant l'équation 4.3. Le tableau 4.3 donne un échantillon des données de vraisemblance de la première année. Les données de disponibilité seront combinées avec les avis d'experts pris au début de la première année (tableaux 4.1 et 4.2) pour procéder à l'inférence bayésienne et déterminer les distributions a posteriori sur les paramètres λ et μ . Au début de la deuxième année, des avis d'experts en fiabilité et en maintenance seront également pris sur les paramètres λ et μ .

Tableau 4.3 Données de vraisemblance

Durées de bons fonctionnement (semaine)	Durées de réparation (semaine)	Disponibilités
0,6045	0,00570	0,9907
1,3086	0,00217	0,9983
.		
.		
0,0721	0,05276	0,5776
6,3209	0,03705	0,9942

Les statistiques a posteriori sur les paramètres λ et μ sont résumées dans le tableau 4.4. Le calcul de ces statistiques est obtenu en appliquant la procédure de la section 4.4. La figure 4.4 illustre le déroulement de l'échantillonnage pour les deux paramètres inconnus.

Tableau 4.4 Statistiques a posteriori de la première année sur les paramètres λ et μ

$\bar{\lambda}_{post,1}$	$\sigma_{\lambda,post}$	$\bar{\mu}_{post,1}$	$\sigma_{\mu,post}$
0,41	0,03	10,6	1,27

Un classement des avis d'experts en fiabilité et maintenance peut être également établi. Ce classement est donné par les fonctions de masse a posteriori des variables aléatoires J et V et est résumé dans les tableaux 4.5 et 4.6.

Tableau 4.5 Fonction de masse a posteriori de la variable J

Expert (j)	$P_{post}(J = j)$
1	0,3594
2	0,3546
3	0,2860

L'optimisation de la disponibilité au début de la deuxième année est obtenue en appliquant la procédure de la section 4.1.2. En effet, nous supposons que le taux de défaillance est linéaire à partir de la deuxième année. L'équation de ce taux est déterminée en utilisant l'information à posteriori de la première année sur le paramètre λ .

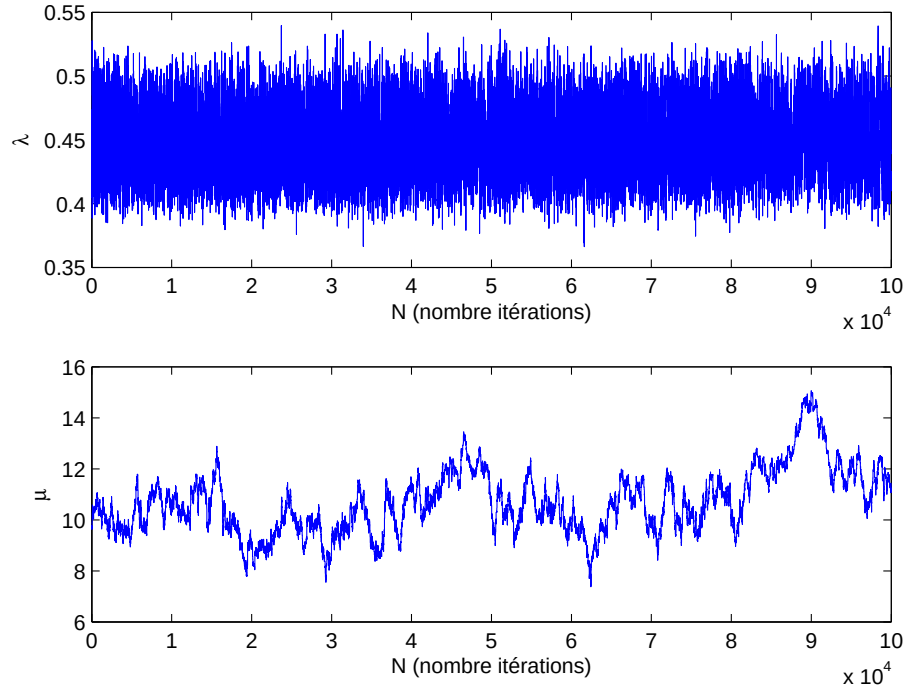


Figure 4.4 Déroulement de l'échantillonnage pour les paramètres de la première année

Tableau 4.6 Fonction de masse a postériori de la variable V

Expert (v)	$P_{post}(V = v)$
1	0,4239
2	0,5761

Étape k ($k \geq 3$) :

Au début de l'année k ($k \geq 3$), des experts en maintenance et en fiabilité donnent leurs avis sur les taux de défaillance et de réparation. Les données de vraisemblance de l'année $k-1$ sont combinées avec les avis d'experts pris au début de l'année $k-1$ pour permettre de calculer la distribution a postériori sur les paramètres inconnus λ et μ . Le taux de défaillance sera modélisé par la forme non linéaire $\lambda_k(t) = a_k + b_k \cdot t^{c_k}$ et la maximisation de la disponibilité sera faite en utilisant la procédure de la section 4.1.2.

Nous pouvons définir la durée résiduelle de l'équipement entre l'année d'évaluation et le moment de la rénovation majeure donné par $t_{p,k}^*$ avec l'équation suivante :

$$D_{résiduelle} = t_{p,k}^* - (k-1) \cdot h$$

Avec :

- $t_{p,k}^*$: Intervalle optimal entre deux rénovations majeures déterminé au début de chaque année.
- k : Désigne l'année à partir de laquelle se fait l'optimisation de la disponibilité.
- h : La longueur de l'année ($h=52$ semaines).

Le tableau 4.7 résume les statistiques a posteriori sur les paramètres inconnus pour les 10 années d'étude. Il donne aussi les intervalles optimaux et les disponibilités optimales déterminés au début de chaque année. La procédure de calcul des statistiques a posteriori ainsi que l'optimisation de la disponibilité sont déterminées en se référant aux étapes décrites précédemment.

Tableau 4.7 Maximisation de la disponibilité de l'équipement

Année (k)	$\bar{\lambda}_{post,k}$	$\bar{\mu}_{post,k}$	$t_{p,k}^*$	$1 - D_k(t_{p,k}^*)$	$D_{résiduelle}$
1	0,41	10,6	$+\infty$	0,960000	$+\infty$
2	0,43	11,2	918,81	0,945595	866,81
3	0,45	12,6	568,52	0,940053	464,51
4	0,48	11,2	644,90	0,944628	488,90
5	0,5	12,3	494,85	0,938749	286,85
6	0,53	11,9	598,06	0,942494	338,06
7	0,58	12,5	580,88	0,941778	268,88
8	0,61	13	554,13	0,941697	190,12
9	0,64	12,1	583,44	0,942318	167,44
10	0,69	11,4	545,91	0,941795	77,90

La figure 4.5 montre l'allure du taux de défaillance sur l'horizon de 10 années. Nous remarquons que la tendance générale du taux de défaillance est l'augmentation. Les cassures du taux de défaillance au début de chaque année résultent des avis d'experts qui changent la tendance du taux de défaillance à travers l'estimation a posteriori de λ .

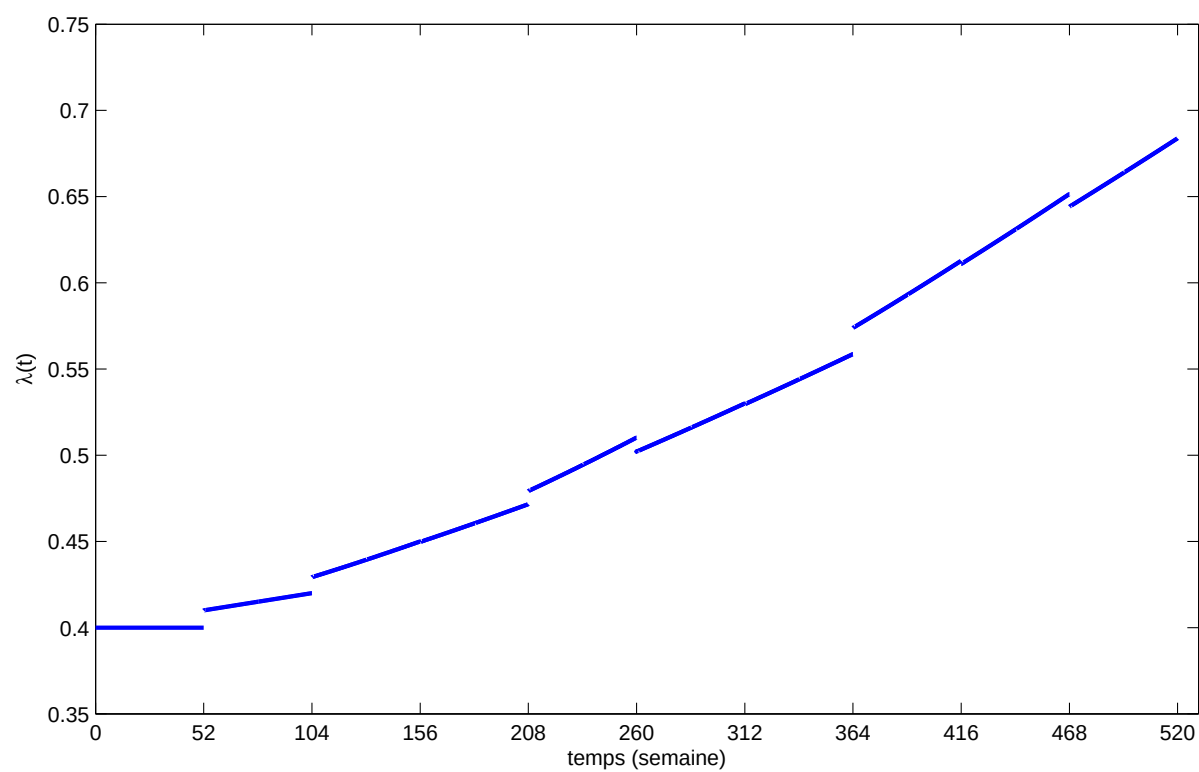


Figure 4.5 Allure du taux de défaillance

CHAPITRE 5

CONCLUSION ET PERSPECTIVES

Ce mémoire propose un modèle bayésien d'agrégation des avis d'experts en exploitation d'équipements. Le modèle permet de prendre en compte plusieurs avis d'experts provenant d'un même domaine d'expertise ou de domaines différents. L'inférence bayésienne est au cœur du développement de ce modèle. Elle permet l'inférence des données de défaillance (vraisemblance) avec les avis d'experts, exprimés sous forme de distributions a priori sur un ou plusieurs paramètres inconnus du modèle de défaillance escompté, et ce pour en déduire les distributions a posteriori sur les mêmes paramètres. Pour modéliser et résoudre le modèle d'inférence bayésienne, nous avons exploité les méthodes de simulation de chaînes de Markov de type MCMC car les distributions a posteriori sur les paramètres d'intérêt ne sont pas connues. L'algorithme de Metropolis-Hasting a été adapté pour permettre la modélisation et la résolution du modèle d'inférence avec deux variables aléatoires simultanément : le taux de défaillance de l'équipement et la variable désignant l'expert le plus performant.

Le modèle d'agrégation proposé est validé à l'aide de deux critères statistiques : le DIC et p-value bayésien. Ces critères ont permis une classification des avis d'experts formulés sur un même paramètre au cours d'une seule période d'inférence. Les résultats obtenus sont comparés aux résultats de deux modèles de combinaison des avis d'experts de la littérature : le modèle du mélange et le modèle de la moyenne des avis d'experts. Les résultats montrent que le modèle d'agrégation bayésien proposé permet non seulement de déterminer automatiquement les distributions a posteriori des paramètres d'intérêt, mais aussi de déterminer la contribution (ou le poids) de chaque expert dans l'évaluation des distributions a posteriori. De plus, le modèle d'agrégation est validé à l'aide d'un test non paramétrique de tendance des données de Fisher afin d'analyser la consistance du modèle d'agrégation dans la détection des experts les plus performants, ceux qui arrivent à prédire correctement la tendance du taux de défaillance. Là aussi les résultats sont logiques et concluants.

Un deuxième modèle d'agrégation des avis d'experts a été proposé pour actualiser deux paramètres inconnus : le taux de défaillance et le taux de réparation d'un équipement, en présence de deux catégories d'avis d'experts : en fiabilité et en maintenance. Ces paramètres font partis du modèle d'optimisation de la disponibilité d'un équipement. Ainsi, ce deuxième modèle permet d'estimer simultanément les distributions a posteriori de quatre paramètres. La mise en exploitation dans le but d'optimiser de la disponibilité est tout à fait nouvelle.

Les modèles de modélisation de la disponibilité recensés dans la littérature considèrent des distributions parfaitement connues sur les variables aléatoires de durées de vie et de durées de réparation (variables aléatoires continus dans le temps). Une nouvelle modélisation est proposée afin de déterminer la périodicité optimale des réparations majeures d'un équipement assujéti à une réparation mineure en cas de défaillance. Cette modélisation rend compte de l'actualisation du modèle d'optimisation à chaque période d'évaluation via l'actualisation de ses paramètres. Ainsi, le modèle proposé est un modèle discrétisé au cours du temps ou encore un modèle défini par morceau. Une durée résiduelle entre le moment d'évaluation de la stratégie et la période optimale est estimée afin d'aider le gestionnaire de l'équipement à maximiser la disponibilité de l'équipement.

En résumé, le travail présenté dans ce mémoire nous a permis de traiter les points suivants :

- Appliquer les méthodes statistiques bayésiennes pour inférer des données de vraisemblance avec des avis d'experts en fiabilité et en maintenance exprimés sur des paramètres inconnus du modèle statistique.
- Utiliser avec efficience les méthodes de simulation MCMC pour modéliser et résoudre des modèles statistiques complexes.
- Modéliser une stratégie prenant en compte un taux de défaillance discret bien qu'initialement elle est applicable à des taux de défaillance continus.
- Optimiser la disponibilité d'un équipement en utilisant la simulation et en prenant en compte des avis d'experts.
- Implémenter un programme Matlab permettant de combiner les avis d'experts sur les taux de défaillance et de réparation. Ce programme permet également l'optimisation de la disponibilité d'un équipement assujéti à la réparation minimale en cas de défaillance et à une rénovation majeure périodiquement.

En définitive, les deux modèles d'agrégation des avis d'experts présentent un avantage majeur. Plus précisément, ils permettent d'actualiser l'information décrivant la performance des experts une fois la vraisemblance connue. Cet avantage n'est pas accessible avec les autres méthodes traitées dans la littérature. Cependant, les modèles proposés ne permettent pas de traiter les cas où il existe une corrélation entre les avis d'experts. Ces modèles se basent sur l'interdépendance des avis d'experts. Cela constitue une limitation. En ce qui concerne le modèle discrétisé proposé pour l'optimisation de la disponibilité, ce modèle présente un avantage important du fait qu'il reflète ou incorpore les décisions prises par les experts en fiabilité et en maintenance tout au long du cycle de vie d'un équipement. Le modèle d'optimisation est tout à fait personnalisable, il permet de gérer équipement par équipement et ainsi refléter les décisions prises pour chacun au cours de son exploitation. Cependant, ce modèle présente une

limitation, celle de considérer un horizon infini pour établir la périodicité optimale. Un modèle d'optimisation de la disponibilité sur un horizon fini pourrait être développé. En pratique, ce modèle sera plus proche de la gestion réelle du cycle de vie des équipements industriels ; une gestion basée sur le principe de durabilité économique des équipements. De nombreux aspects de ce projet peuvent être améliorés. Tout d'abord, une étude de sensibilité peut être traitée pour permettre trouver le lien entre l'incertitude des experts et les distributions a posteriori sur les paramètres inconnus. Cela pourrait-être utile pour mesurer l'erreur de prédiction des distributions a posteriori par rapport aux incertitudes des experts. Aussi, les modèles d'agrégation peuvent être modifiés pour prendre en compte des corrélations possibles entre les avis d'experts.

RÉFÉRENCES

- ALBERT, I., DONNET, S., GUIHENNEUC-JOUYAUX, C., LOW-CHOY, S., Mengersen, K. et ROUSSEAU, J. (2012). Combining expert opinions in prior elicitation. *Bayesian Analysis*, 7, 503–532.
- BACHA, M., CELEUX, G., IDÉE, E., LANNOY, A. et VASSEUR, D. (1998). *Estimation de modèles de durées de vie fortement censurées*. Editions Eyrolles.
- BARLOW, R. E. et PROSCHAN, F. (1972). Availability theory for multicomponent systems. Rapport technique, DTIC Document.
- BARLOW, R. E. et PROSCHAN, F. (1996). *Mathematical theory of reliability*, vol. 17. Siam.
- BERGER, J. O. (1985). *Statistical decision theory and Bayesian analysis*. Springer.
- BURGMAN, M. A., MCBRIDE, M., ASHTON, R., SPEIRS-BRIDGE, A., FLANDER, L., WINTLE, B., FIDLER, F., RUMPF, L. et TWARDY, C. (2011). Expert status and performance. *PLoS One*, 6, e22998.
- CARLIN, B. P. et LOUIS, T. A. (2000). *Bayes and empirical Bayes methods for data analysis*. CRC Press.
- CASELLA, G. et ROBERT, C. P. (1999). Monte carlo statistical methods. Springer-Verlag, New York.
- CHAN, J.-K. et SHAW, L. (1993). Modeling repairable systems with failure rates that depend on age and maintenance. *Reliability, IEEE Transactions on*, 42, 566–571.
- CHIB, S. et GREENBERG, E. (1995). Understanding the metropolis-hastings algorithm. *The American Statistician*, 49, 327–335.
- COOKE, R. M. et GOOSSENS, L. L. (2008). Tu delft expert judgment data base. *Reliability Engineering & System Safety*, 93, 657–674.
- ELSAIED, E. A. (2012). *Reliability engineering*. Wiley Publishing.
- FRENCH, S. (2011). Aggregating expert judgement. *Revista de la Real Academia de Ciencias Exactas, Físicas y Naturales. Serie A. Matemáticas*, 105, 181–206.
- GELFAND, A. E., HILLS, S. E., RACINE-POON, A. et SMITH, A. F. (1990). Illustration of bayesian inference in normal data models using gibbs sampling. *Journal of the American Statistical Association*, 85, 972–985.
- GELFAND, A. E. et SMITH, A. F. (1990). Sampling-based approaches to calculating marginal densities. *Journal of the American statistical association*, 85, 398–409.

- GELMAN, A., MENG, X.-L. et STERN, H. (1996). Posterior predictive assessment of model fitness via realized discrepancies. *Statistica sinica*, 6, 733–760.
- GEMAN, S. et GEMAN, D. (1984). Stochastic relaxation, gibbs distributions, and the bayesian restoration of images. *Pattern Analysis and Machine Intelligence, IEEE Transactions on*, 721–741.
- GENEST, C. et ZIDEK, J. V. (1986). Combining probability distributions : A critique and an annotated bibliography. *Statistical Science*, 1, 114–135.
- GILKS, W. R., RICHARDSON, S. et SPIEGELHALTER, D. J. (1996). Introducing markov chain monte carlo. *Markov chain Monte Carlo in practice*, Springer. 1–19.
- HASTINGS, W. K. (1970). Monte carlo sampling methods using markov chains and their applications. *Biometrika*, 57, 97–109.
- HERD, G. R. (1960). Estimation of reliability from incomplete data. *Proceedings of the 6th national symposium on reliability and quality control*. 202–17.
- JARDINE, A. K. et TSANG, A. H. (2013). *Maintenance, replacement, and reliability : theory and applications*. CRC press.
- KABAK, I. W. (1969). System availability and some design implications. *Operations Research*, 17, 827–837.
- KAPLAN, E. L. et MEIER, P. (1958). Nonparametric estimation from incomplete observations. *Journal of the American statistical association*, 53, 457–481.
- KECECIOGLU, D. (2003). *Maintainability, availability, and operational readiness engineering handbook*, vol. 1. DEStech Publications, Inc.
- LEHMANN, E. L. et D’ABRERA, H. J. (2006). *Nonparametrics : statistical methods based on ranks*. Springer New York.
- LEJEUNE, M. (2011). *Statistique. La théorie et ses applications : La thêorie et ses applications*. Springer.
- LIPSCOMB, J., PARMIGIANI, G. et HASSELBLAD, V. (1998). Combining expert judgment by hierarchical modeling : An application to physician staffing. *Management Science*, 44, 149–161.
- MARTINEZ, W. L. et MARTINEZ, A. R. (2001). *Computational statistics handbook with MATLAB*. CRC press.
- METROPOLIS, N., ROSENBLUTH, A. W., ROSENBLUTH, M. N., TELLER, A. H. et TELLER, E. (1953). Equation of state calculations by fast computing machines. *The journal of chemical physics*, 21, 1087–1092.

- O'HAGAN, A., BUCK, C. E., DANESHKHAH, A., EISER, J. R., GARTHWAITE, P. H., JENKINSON, D. J., OAKLEY, J. E. et RAKOW, T. (2006). *Uncertain judgements : eliciting experts' probabilities*. John Wiley & Sons.
- OUALI, M.-S. (2013). Maintenance et sécurité industrielle. Notes de cours.
- PROCACCIA, H., FERTON, E. et PROCACCIA, M. (2011). *Fiabilité et maintenance des matériels industriels réparables et non réparables*. Lavoisier.
- ROBERTS, G. O. (1996). Markov chain concepts related to sampling algorithms. *Markov chain Monte Carlo in practice*, Springer. 45–57.
- SINHARAY, S., JOHNSON, M. S. et STERN, H. S. (2006). Posterior predictive assessment of item response theory models. *Applied Psychological Measurement*, 30, 298–321.
- SINHARAY, S. et STERN, H. S. (2003). Posterior predictive model checking in hierarchical models. *Journal of Statistical Planning and Inference*, 111, 209–221.
- SPIEGELHALTER, D. J., BEST, N., CARLIN, B. P. et VAN DER LINDE, A. (1998). Bayesian deviance, the effective number of parameters, and the comparison of arbitrarily complex models. Rapport technique, Research Report, 98-009.
- SPIEGELHALTER, D. J., BEST, N. G., CARLIN, B. P. et VAN DER LINDE, A. (2002). Bayesian measures of model complexity and fit. *Journal of the Royal Statistical Society : Series B (Statistical Methodology)*, 64, 583–639.
- WALLACH, M. A., KOGAN, N. et BEM, D. J. (1962). Group influence on individual risk taking. *Journal of abnormal and social psychology*, 65, 75–86.
- WANG, H. et PHAM, H. (1999). Some maintenance models and availability with imperfect maintenance in production systems. *Annals of Operations Research*, 91, 305–318.
- WANG, H. et PHAM, H. (2006a). Availability and maintenance of series systems subject to imperfect repair and correlated failure and repair. *European Journal of Operational Research*, 174, 1706–1722.
- WANG, H. et PHAM, H. (2006b). *Reliability and optimal maintenance*. Springer.
- WINKLER, R. L. (1968). The consensus of subjective probability distributions. *Management Science*, 15, B-61–B-75.
- ZHAO, M. (1994). Availability for repairable components and series systems. *Reliability, IEEE Transactions on*, 43, 329–334.

ANNEXE A

PROGRAMME MATLAB 1 : AGRÉGATION DES AVIS D'EXPERTS

```

% importation des donnees
n_beta=xlsread('donnees',1,'R19');
n_normal=xlsread('donnees',1,'S19');
n=n_normal+n_beta;
p=xlsread('donnees',1,'C:C');
min_exp=xlsread('donnees',1,'D:D');
max_exp=xlsread('donnees',1,'E:E');

% Avis des experts
e=zeros(n,2);
for i=1:n_beta
e(i,:)=fsolve(@(z) val(min_exp(i),max_exp(i),z),[1 1]);
end

for i=(n_beta+1):n
e(i,:)=fsolve(@(z) valn(min_exp(i),max_exp(i),z),[0.5 0.1]);
end

% Generation des observations.
load('A1');
t_obs =A1;
t=sum(t_obs);
m=length(t_obs);

% Generation de N echantillonnage.

% Les parametres
N =50000;
sigma = 0.2;

```

```

lambda = zeros(1,N);
par_exp=zeros(1,N);

lambda(1) =m/t;
par_exp(1)=1;
vect_saut=(1/n)*ones(1,n);
alea=zeros(1,N);

matrice_alea=mnrnd(1,vect_saut,N);

for j=1:N
    alea(j)=find(matrice_alea(j,:)==1);
end

Moy_lambda=zeros(n,1);
Et_lambda=zeros(n,1);
DIC=zeros(n,1);
p_value=zeros(n,1);

% Debut de l'echantillonnage

for i=2:N

    % Generation d'un point

    theta = lambda(i-1) + sigma.*randn(1);

    % ratio d'acceptation

    alpha = min([1,(theta/lambda(i-1))^m*exp(t*(lambda(i-1)-theta))...
    *(posteriori_exp(theta,e(alea(i),1),e(alea(i),2),alea(i),n_beta) ...
    .*p(alea(i)).*vect_saut(par_exp(i-1)) ...
    ./vect_saut(alea(i))./p(par_exp(i-1))./...
    ((posteriori_exp(lambda(i-1),e(par_exp(i-1),1),e(par_exp(i-1),2),par_exp(i-1),n_beta)
    )))]);

```

% Generation d'un point de la loi uniforme

u = rand;

if u <= alpha

lambda(i) = theta;

par_exp(i)=alea(i);

else

lambda(i) = lambda(i-1);

par_exp(i)=par_exp(i-1);

end

end

% moyenne et ecart type du taux de defaillance

Moy = **mean**(lambda);

Et = **std**(lambda);

% DIC du modele d'aggregation

DIC_ag=-4*m***mean**(**log**(lambda))+2*m***log**(Moy)+2*Moy*t;

% p-value du modele d'aggregation

DEVag_obs=**zeros**(1,N);

DEVag_rep=**zeros**(1,N);

for i=1:N

DEVag_obs(i)=-2*m***log**(lambda(i))+2*m*lambda(i)***mean**(t_obs);

t_rep=exprnd(1/lambda(i),1,m);

DEVag_rep(i)=-2*m***log**(lambda(i))+2*m*lambda(i)***mean**(t_rep);

end

p_value_ag=**length**(**find**(DEVag_rep>=DEVag_obs))/N;

SF=[Moy Et DIC_ag p_value_ag];

% probabilite a posteriori de chaque expert

pr=**zeros**(n,1);

for i=1:n

pr(i,1)=**length**(**find**(par_exp==i))/N;

end

```

for j=1:n
lambda = zeros(1,N);
lambda(1) = m/t;

for i=2:N

    % Generation d'un point

    theta = lambda(i-1) + sigma.*randn(1);

    % ratio d'acceptation

    alpha = min([1,(theta/lambda(i-1))^m*exp(t*(lambda(i-1)-theta))...
        *(posteriori_exp(theta,e(j,1),e(j,2),j,n_beta) ...
        ./((posteriori_exp(lambda(i-1),e(j,1),e(j,2),j,n_beta))))]);

    % Generation d'un point de la loi uniforme

    u = rand;

    if u <= alpha
    lambda(i) = theta;
    else
    lambda(i) = lambda(i-1);
    end

end

Moy_lambda(j) = mean(lambda);
Et_lambda(j) = std(lambda);

DEV_obs=zeros(1,N);
DEV_rep=zeros(1,N);

for i=1:N

```

```

DEV_obs(i)=-2*m*log(lambda(i))+2*m*lambda(i)*mean(t_obs);
t_rep=exprnd(1/lambda(i),1,m);
DEV_rep(i)=-2*m*log(lambda(i))+2*m*lambda(i)*mean(t_rep);
end
p_value(j)=length(find(DEV_rep>=DEV_obs))/N;

DIC(j)=-4*m*mean(log(lambda))+2*m*log(Moy_lambda(j))+2*Moy_lambda(j)*t;
Pd(j)=DIC(j)-2*Moy_lambda(j)*t+2*m*mean(log(lambda));
end

figure;
subplot(2,1,1)
x=transpose(lambda);
x_values = 0:0.0001:1;
z = makedist('Beta','a',e(1,1),'b',e(1,2));
k = makedist('beta','a',e(2,1),'b',e(2,2));
zz = pdf(z,x_values);
kk = pdf(k,x_values);
h = ksdensity(x,x_values);
plot(x_values,h,'-', x_values,zz,'--',x_values,kk,'-.');
xlabel('\lambda');
legend('Posterior','expert1','expert2')
title('Distributions_a_priori_et_a_posteriori')
subplot(2,1,2)
yy=1:1:N;
plot(yy,lambda);
xlabel('R_(nombre_itrations)');
ylabel('\lambda');
title('Droulement_de_l_chantillonnage')

```

ANNEXE B

PROGRAMME MATLAB 2 : OPTIMISATION DE LA DISPONIBILITÉ

```

% initialisation
h=52;
T_op=[];
D_op=[];
DIC_post=zeros(10,1);
temps=[0 h 2*h 3*h 4*h 5*h 6*h 7*h 8*h 9*h 10*h];
load('Donnees_vr.mat');
lambda_post=[];
mu_post=[];
lambda_post(1)=0.45;
mu_post(1)=10;
T_op(1)=inf;
D_op(1)=1-lambda_post(1)/mu_post(1);
Tp=8;
a=zeros(1,8);
b=zeros(1,8);
c=zeros(1,8);

% importation des donnees
for k=2:11

T=[];
D=[];
p_lambda=xlsread('maitrise',k-1,'F4:F20');
p_mu=xlsread('maitrise',k-1,'J4:J20');
min_exp_lambda=xlsread('maitrise',k-1,'G4:G20');
max_exp_lambda=xlsread('maitrise',k-1,'H4:H20');
min_exp_mu=xlsread('maitrise',k-1,'K4:K20');
max_exp_mu=xlsread('maitrise',k-1,'L4:L20');
n_lambda=xlsread('maitrise',k-1,'O2');
n_mu=xlsread('maitrise',k-1,'P2');

```



```

T=eval(['T' int2str(k-1)]);
D=eval(['D' int2str(k-1)]);
t=T./(D+T);
    actualisation_disponibilite ;
% rsultats a posteriori
DIC_post(k-1)=DIC;
lambda_post(k)=Moy_lambda;
mu_post(k)=Moy_mu;
Et_lambda_post(k)=Et_lambda;
Et_mu_post(k)=Et_mu;
S{k-1}=pr_lambda;
R{k-1}=pr_mu;

% optimisation de la strategie de maintenance
if k>2
x= wkeep(temps,k,'l');
f = fit(x , lambda_post', 'a*x^b+c', 'Lower',[0,0,0], 'StartPoint', [0, 0, 0], 'upper'
    ,[10,10,10]) ;
r=coeffvalues(f);
a(k-2)=r(1);
b(k-2)=r(2);
c(k-2)=r(3);
end
v = @(T) panne7(lambda_post,k,h,T,Tp,mu_post,temps,a,b,c);
T_op(k)= fminbnd(v, temps(k), 10000);
D_op(k)=1-panne7(lambda_post,k,h,T_op(k),Tp,mu_post,temps,a,b,c);
end

% allure du taux de defaillance , allure du temps optimal de remplacement et celle de la
    disponibilite
x=temps;
y=lambda_post;
subplot(3,2,[3 4])
plot([x(1) x(2) ],[ y(1) y(1) ], 'linewidth',2);
hold on;
z=x(2):.01:x(3);

```

```

r=@(t) ((y(2)-y(1))/h)*t+y(1);
plot(z,r(z), 'linewidth',2);
hold on;
for i=3:10
z=x(i):.01:x(i+1);
v = wkeep(y,i,'l')';
u= wkeep(x,i,'l')';
f = fit( u, v, 'a*x^b+c', 'Lower', [0,0,0], 'StartPoint', [0, 0, 0], 'upper', [10,10,10]) ;
plot(z,f(z), 'linewidth',2);
hold on;
end
set(gca,'Xtick',0:h:10*h);
% set(gca,'XtickLabel',{'0','h','2h','3h','4h','5h','6h','7h','8h','9h','10h'});
set(gca,'XtickLabel',{'0','52','104','156','208','260','312','364','416','468','520'});
axis([0 530 0.45 0.9]) ;
xlabel('temps_(semaine)')
ylabel('\lambda');

subplot(3,2,1);
plot(temps,T_op,'--gs',...
'LineWidth',2,...
'MarkerSize',10,...
'MarkerEdgeColor','k',...
'MarkerFaceColor',[0.5,0.5,0.5])
set(gca,'Xtick',0:h:10*h);
% set(gca,'XtickLabel',{'0','h','2h','3h','4h','5h','6h','7h','8h','9h','10h'});
set(gca,'XtickLabel',{'0','52','104','156','208','260','312','364','416','468','520'});
xlabel('temps_(semaine)');
ylabel('T_optimal');

subplot(3,2,2)
plot(temps,D_op,'--gs',...
'LineWidth',2,...
'MarkerSize',10,...
'MarkerEdgeColor','k',...
'MarkerFaceColor',[0.5,0.5,0.5])

```

```

set(gca,'Xtick',0:h:10*h);
% set(gca,'XtickLabel',{'0','52','104','156','208','260','312','364','416','468','520'});
xlabel('temps_(semaine)');
ylabel('D_optimal');

subplot(3,2,[5 6])
plot(temps,mu_post,'--gs',...
     'LineWidth',2,...
     'MarkerSize',10,...
     'MarkerEdgeColor','k',...
     'MarkerFaceColor',[0.5,0.5,0.5])
set(gca,'Xtick',0:h:10*h);
% set(gca,'XtickLabel',{'0','h','2h','3h','4h','5h','6h','7h','8h','9h','10h'});
set(gca,'XtickLabel',{'0','52','104','156','208','260','312','364','416','468','520'});
xlabel('temps_(semaine)');
ylabel('\mu');

% Ecriture des resultats
filename = 'resultats.xlsx';
xlswrite(filename,T_op,1,'B2');
xlswrite(filename,D_op,1,'C2');
xlswrite(filename,lambda_post,1,'D2');
xlswrite(filename,mu_post,1,'E2');
xlswrite(filename,S{1},2,'A3');
xlswrite(filename,S{2},2,'C3');
xlswrite(filename,S{3},2,'E3');
xlswrite(filename,S{4},2,'G3');
xlswrite(filename,S{5},2,'I3');
xlswrite(filename,S{6},2,'K3');
xlswrite(filename,S{7},2,'M3');
xlswrite(filename,S{8},2,'O3');
xlswrite(filename,S{9},2,'Q3');
xlswrite(filename,S{10},2,'S3');
xlswrite(filename,R{1},2,'B3');
xlswrite(filename,R{2},2,'D3');
xlswrite(filename,R{3},2,'F3');

```

```
xlswrite(filename,R{4},2,'H3');  
xlswrite(filename,R{5},2,'J3');  
xlswrite(filename,R{6},2,'L3');  
xlswrite(filename,R{7},2,'N3');  
xlswrite(filename,R{8},2,'P3');  
xlswrite(filename,R{9},2,'R3');  
xlswrite(filename,R{10},2,'T3');
```