



AVERTISSEMENT

Ce document est le fruit d'un long travail approuvé par le jury de soutenance et mis à disposition de l'ensemble de la communauté universitaire élargie.

Il est soumis à la propriété intellectuelle de l'auteur. Ceci implique une obligation de citation et de référencement lors de l'utilisation de ce document.

D'autre part, toute contrefaçon, plagiat, reproduction illicite encourt une poursuite pénale.

Contact : ddoc-theses-contact@univ-lorraine.fr

LIENS

Code de la Propriété Intellectuelle. articles L 122. 4

Code de la Propriété Intellectuelle. articles L 335.2- L 335.10

http://www.cfcopies.com/V2/leg/leg_droi.php

<http://www.culture.gouv.fr/culture/infos-pratiques/droits/protection.htm>

THÈSE

Présentée à

L'Université Paul Verlaine – METZ

UFR Mathématiques, Informatique, Mécanique

Pour obtenir le titre de

DOCTEUR

en

Automatique, Traitement du Signal et des Images, Génie Informatique

Par

Jérémie SCHUTZ

**Contribution à l'optimisation des plans d'exploitation et de maintenance,
selon une approche basée sur le pronostic : Application au domaine naval**

Soutenue le jeudi 3 décembre 2009, devant le jury composé de :

Rapporteurs

Daoud AÏT-KADI

Professeur à l'Université de Laval, Canada

Bernard GRABOT

Professeur à l'École Nationale d'Ingénieurs de Tarbes

Examineurs

El Houssaine AGHEZZAF

Professeur à l'Université de Gent, Belgique

Benoît IUNG

Professeur à l'Université Henri Poincaré Nancy I

Jean-Baptiste LÉGER

Président de la société PREDICT

Directeur de thèse

Nidhal REZG

Professeur à l'Université Paul Verlaine – Metz

À mes très chers parents,

*J'ai longtemps cherché les mots les plus justes
pour vous remercier d'avoir été, et d'être toujours présents,
de m'avoir épaulé, soutenu lorsque je voulais tout abandonner.
Mais après mûre réflexion, je ne vois rien de plus fort que...
Je vous aime.*

À la mémoire de Théodore et Louis

Remerciements

Le travail de recherche présenté dans ce mémoire a été accompli au sein de l'équipe Fiabilité & Maintenance du Laboratoire de Génie Industriel et de Production de Metz. Ce travail s'intègre dans le cadre du projet CosTeam soutenu par l'INRIA Nancy – Grand Est dont le responsable scientifique est le Pr. Nidhal REZG, directeur du LGIPM.

Je ne pouvais pas commencer ces remerciements sans évoquer mon directeur de thèse, Nidhal REZG, Professeur à l'Université Paul Verlaine – Metz. En plus de sa confiance et de son aide scientifique qu'il m'a accordées, son amabilité et sa bonne humeur ont rendu ces trois années de Doctorat fort agréables.

Je tiens également à adresser mes plus vifs remerciements aux membres du jury, qui ont accepté d'évaluer mon travail de thèse.

Merci à MM. Daoud AÏT-KADI, Professeur à l'Université de Laval et Bernard GRABOT, Professeur à l'École Nationale d'Ingénieurs de Tarbes, d'avoir accepté d'être les rapporteurs de ce manuscrit. Leurs remarques et suggestions me permettront d'apporter des améliorations à la qualité de ce dernier.

Merci à MM. El Houssaine AGHEZZAF, Professeur à l'Université de Gent, Benoît IUNG, Professeur à l'Université Henri Poincaré Nancy I, Jean-Baptiste LÉGER, Président de la société PREDICT et Docteur de l'Université Henri Poincaré Nancy I, pour avoir accepté d'examiner mon travail de recherche et de faire partie de mon jury de thèse.

C'est avec un plaisir non dissimulé que j'exprime plus particulièrement ma reconnaissance amicale à Mohammed DAHANE et Abdelhakim KHATAB, Maîtres de Conférences à l'École Nationale d'Ingénieurs de Metz, qui par leur disponibilité, leur aide et leurs grandes compétences scientifiques, ont largement contribué à la réalisation et à la qualité de ce travail.

Je tiens également à remercier vivement l'ensemble des chercheurs et secrétaires du laboratoire pour l'ambiance amicale dans laquelle s'est déroulé ce travail et, en particulier, Z. ACHOUR, S. DELLAGI, L. BENYOUCEF, N. SAUER, T. MONTEIRO, sans oublier C. FOUSSE et C. WIEMERT pour leur gentillesse et leur dévouement. Enfin, une pensée pour F. MUNERATO avec qui j'ai eu le plaisir de partager quelques points de vues avant qu'il ne nous quitte.

Je remercie aussi l'ensemble des doctorants sans qui tout ce travail n'aurait pas été aussi enrichissant qu'agréable. Je pense particulièrement à Simón, Norly, Foued, Hatem, Zied et Zhouhang.

Ma reconnaissance et ma profonde gratitude s'adressent aux Dr. Anis CHELBI et Dr. Mehdi RADHOU pour les quelques moments partagés en leur compagnie ainsi que pour leurs avis éclairés qui s'avèrent d'une aide précieuse.

J'associe aussi à ces remerciements l'ensemble du corps enseignants qui m'a permis d'atteindre ce niveau et, en particulier, Mme. CHÉRY et M. REMY-VINCENT.

Enfin, j'adresse tous mes remerciements à ma famille et notamment à mes parents pour leur amour et leur soutien inconditionnel durant toutes ces années.

Tentative

Entre

Ce que je pense,

Ce que je veux dire,

Ce que je crois dire,

Ce que je dis,

Ce que vous avez envie d'entendre,

Ce que vous croyez entendre,

Ce que vous entendez,

Ce que vous avez envie de comprendre,

Ce que vous croyez comprendre,

Ce que vous comprenez,

Il y a dix possibilités qu'on ait des difficultés à communiquer.

Mais essayons quand même...

Encyclopédie du Savoir Relatif et Absolu

Bernard WERBER

Sommaire

Introduction générale	13
Chapitre I : Position du problème et mise en situation	15
I. Introduction	16
II. Les types de maintenances	16
III. Les politiques de maintenances	19
IV. L'horizon de temps fini	21
V. Les conditions de fonctionnement	23
VI. Ordonnancement des tâches de production	24
VII. Couplage de la maintenance et de la production	24
VIII. Description du problème	27
VIII.1. Notations	28
VIII.2. Définitions	28
VIII.3. Hypothèses de travail	29
IX. Conclusion	30
Chapitre II : Modélisation de la loi de défaillance	31
I. Introduction	32
II. Les différentes approches de pronostics	32
II.1. Le pronostic basé sur un modèle physique	33
II.2. Le pronostic guidé par les données	33
II.3. Le pronostic basé sur l'expérience	33
III. Modèle de vie accélérée	34
IV. Modèle à risques proportionnels	35
V. Formalisation de la fonction de risque	36
VI. Structure d'un raisonnement par logique floue	37
VI.1. Fuzzification	38
VI.2. Raisonnement flou	39
VI.3. Défuzzification	40
VII. Application à notre problématique	41
VII.1. Structure des variables linguistiques	42
VII.2. Définition du moteur d'inférence flou	45
VII.3. Détermination de la valeur de sortie	46
VII.4. Exemple illustratif	46
VIII. Conclusion	50

Chapitre III : Études et analyses de politiques intégrées	51
I. Introduction	52
II. Les différentes stratégies d'optimisation	52
II.1. La stratégie séquentielle	52
II.2. La stratégie intégrée	54
II.2.a. Recherche de solutions efficaces	54
II.2.b. Principe de fonctionnement des algorithmes génétiques	55
II.2.c. Application de l'algorithme génétique à la détermination des plans d'exploitation	57
II.2.c.i Le codage	57
II.2.c.ii La stratégie de reproduction	57
II.2.c.iii L'opérateur de croisement	57
II.2.c.iv L'opérateur de mutation	59
III. La politique minimaliste	59
III.1. L'âge fonctionnel	60
III.1.a. Exemple illustratif	62
III.2. Modélisation des maintenances préventives	63
IV. La politique systématique	64
IV.1. Recherche du facteur d'amélioration optimal	66
IV.1.a. Modélisation du coût et de la durée des actions de maintenances préventives	66
IV.2. Recherche des facteurs d'amélioration optimaux	71
IV.2.a. Méta-heuristique pour la recherche des facteurs d'amélioration optimaux	73
IV.2.a.i Le codage	74
IV.2.a.ii La stratégie de reproduction	74
IV.2.a.iii L'opérateur de croisement	74
IV.2.a.iv L'opérateur de mutation	75
V. La politique sporadique	75
V.1. Recherche du facteur d'amélioration optimal	78
V.2. Recherche des facteurs d'amélioration optimaux	79
V.2.a. Méta-heuristique pour la recherche des facteurs d'amélioration optimaux	80
V.2.a.i Le codage	80
V.2.a.ii Définition du voisinage	81
VI. La politique périodique	81
VI.1. Recherche du facteur d'amélioration optimal	87
VI.2. Recherche des facteurs d'amélioration optimaux	88
VI.2.a. Méta-heuristique pour la recherche des facteurs d'amélioration optimaux	89
VII. La politique séquentielle	90
VII.1. Méta-heuristique pour l'obtention de solutions efficaces	92
VII.1.a. Le codage	92
VII.1.b. L'opérateur de croisement	93
VII.1.c. L'opérateur de mutation	94
VII.2. Recherche du facteur d'amélioration optimal	94
VII.3. Recherche des facteurs d'amélioration optimaux	95
VII.3.a. Méta-heuristique pour la recherche des variables de décision	95

VIII. Conclusion	96
Chapitre IV : Exemple illustratif	97
I. Introduction	98
II. Données numériques	98
II.1. Les missions et l'horizon de temps	98
II.2. Les données relatives à la fiabilité du système	99
II.3. Les données relatives aux méta-heuristiques	100
III. Première étude : la stratégie séquentielle	100
III.1. Détermination du plan d'exploitation optimal	100
III.2. Les politiques de maintenances	101
IV. Seconde étude : la stratégie intégrée	102
IV.1. La politique systématique	102
IV.2. La politique sporadique	104
IV.3. La politique périodique	106
IV.4. La politique séquentielle	108
V. Synthèse au sujet des stratégies séquentielle et intégrée	109
Conclusion	111
Bibliographie	113

Introduction générale

Habituellement, dans le cas des systèmes manufacturiers, il est supposé que la production de biens est réalisée sous des conditions de fonctionnement prédéterminées et, généralement, constantes au cours du temps. De ce fait, la planification des inspections et/ou des actions préventives, utilisée pour maintenir l'équipement et réduire l'occurrence des défaillances, ne prend pas en considération ces conditions. Cependant, dès lors que l'on s'éloigne de ces systèmes de production et que l'on s'intéresse aux domaines aéronautiques ou maritimes, cette hypothèse ne peut plus être vérifiée. Prenons le cas de l'aéronautique : les conditions de navigation et les conditions météorologiques tiennent un rôle très important. Par exemple, les sondes Pitot, situées sur le nez des appareils ou sous les ailes, servent à déterminer, à partir des pressions statique et totale, la vitesse des avions. Malheureusement, lorsque les conditions météorologiques deviennent inclementes, à cause de la pluie ou de la neige, ces sondes peuvent givrer, engendrant des conséquences dramatiques. Il est donc primordial d'intégrer ces conditions de fonctionnement dans les politiques d'inspection. D'un point de vue analogue, dans le domaine maritime, les risques de défaillances peuvent être liés aux conditions opérationnelles et environnementales. Un navire peut être amené à réaliser des missions pour lesquelles les conditions opérationnelles et environnementales sont différentes. Un autre exemple concerne un navire qui peut être amené à réaliser des missions d'ordre militaire, diplomatique, humanitaire, etc. Alors que des missions océanographiques nécessitent un régime de fonctionnement constant, des missions de démonstration de forces requièrent des conditions opérationnelles bien moins idéales, liées à des variations importantes au niveau du régime de fonctionnement. En plus des conditions opérationnelles liées au type de missions, les conditions environnementales impactent également le niveau de dégradation du système. En effet, la température à la surface de l'eau varie principalement en fonction de la latitude. Ainsi, les mers situées à des latitudes élevées peuvent atteindre des températures de l'ordre de -2 degrés Celsius, alors que la température à des latitudes faibles peut atteindre les 36 degrés Celsius. À l'instar de la température, d'autres paramètres peuvent influencer tels que le taux de salinité, les ondes électromagnétiques. La connaissance de ces conditions de fonctionnement permet donc, connaissant les différentes missions à réaliser, de définir des politiques de maintenances appropriées pour minimiser les coûts de maintenances. Par ailleurs, lorsque le navire quitte le quai, il s'engage en mer pour une durée conséquente. Pendant cette période, aucun retour à quai n'est prévu pour la réalisation d'actions de maintenances préventives. Il convient donc de définir, au préalable, le plan de maintenance adéquat afin de disposer des ressources nécessaires, tant du point de vue humain que matériel. La détermination de ce plan de maintenance est donc liée aux conditions de fonctionnement suivant lequel le navire sera soumis, ces conditions dépendant des missions accomplies. De ce fait, avant de quitter le quai, le plan

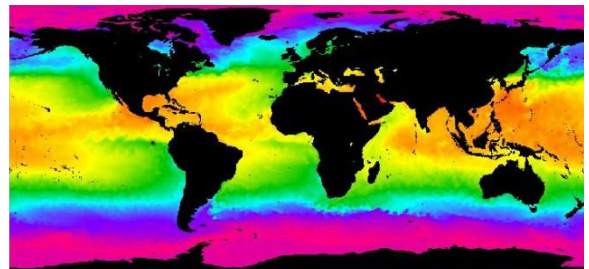


Figure 1 : Température à la surface des eaux

d'exploitation (l'ensemble ordonné des missions à réaliser) ainsi que le plan de maintenance (dates et types des actions de maintenances préventives) doivent être attribués au navire. Le travail de recherche, proposé dans ce manuscrit, consiste donc à déterminer ces deux plans dans l'objectif de maximiser les gains, sachant que les plans de maintenance dépendent des plans d'exploitation.

Ce mémoire sera organisé de la manière suivante :

Dans le premier chapitre, le contexte d'étude sera présenté et permettra la mise en situation de notre problématique par rapport aux travaux existants du point de vue de la maintenance et de la gestion de la production. Afin de souligner la spécificité de notre problématique, nous nous focaliserons sur les lacunes présentes dans la littérature, à savoir les variations des conditions de fonctionnement au cours du temps ainsi que l'application de politiques de maintenance à des horizons de temps fini.

Le deuxième chapitre sera axé sur la modélisation de la loi de défaillance. En effet, les conditions de fonctionnement étant dépendantes des missions à réaliser, l'évolution des dégradations, et donc des défaillances, n'est pas homogène au cours du temps. Par conséquent, la loi de dégradation devra s'adapter aux différentes missions afin de permettre une prévision de l'état futur du système. Parmi l'étude des deux modèles existants dans la littérature, il sera nécessaire d'utiliser un raisonnement basé sur la logique floue afin d'adapter les données historiques issues du (ou des) navires au modèle sélectionné.

Par la suite, le troisième chapitre s'intéressera à l'optimisation des plans d'exploitation et de maintenance. À ce titre, une première partie portera sur les différentes stratégies utilisables pour une optimisation conjointe des deux plans. Pour la stratégie séquentielle, il conviendra de déterminer, initialement, le plan d'exploitation optimal et d'y associer un plan de maintenance satisfaisant. Concernant la stratégie intégrée, la détermination des plans d'exploitation et de maintenance devra être réalisée simultanément. Une méta-heuristique sera utilisée afin de déterminer une solution satisfaisante. La phase d'optimisation des plans de maintenance nécessitera l'étude préalable des différentes politiques proposées. En effet, pour minimiser les coûts liés aux actions de maintenance, différentes politiques préventives seront développées. Dans un premier temps, les politiques préventives se caractériseront par la réalisation des actions préventives uniquement à la fin des missions. La politique systématique consiste à réaliser une action préventive après chaque mission accomplie alors que pour la politique sporadique, il conviendra de choisir les missions suivies de maintenances préventives. Par la suite, cette contrainte sera relaxée pour étudier les politiques de types périodique et séquentiel. La politique périodique se caractérise par des intervalles de temps réguliers, dépendant du nombre d'actions préventives, alors que pour la politique séquentielle, les durées de chaque intervalle de maintenances préventives peuvent être différentes. Afin d'améliorer ces différentes politiques de maintenances préventives, nous chercherons à optimiser le facteur d'amélioration duquel dépend l'efficacité de la maintenance.

Enfin, le dernier chapitre sera consacré à un exemple numérique. Il permettra d'illustrer les différentes politiques développées au cours de ce mémoire et de percevoir l'impact des différentes stratégies et politiques de maintenances sur l'évolution des bénéfices.

Chapitre I

Position du problème et mise en situation

DEPUIS de nombreuses années, l'intérêt de la fonction maintenance et son intégration au sein d'un système de production ne sont plus à démontrer. Aussi, dans ce premier chapitre, un état de l'art concernant le domaine de la maintenance sera mené et permettra de mettre en évidence des lacunes encore présentes dans la littérature. Une partie de ce chapitre sera consacrée au domaine de la production et plus particulièrement à l'ordonnancement des tâches de production ainsi qu'au couplage entre la gestion de la production et la gestion de la maintenance. Finalement, ce chapitre s'achèvera par le positionnement du problème traité dans cette thèse et sa mise en situation.

Sommaire

I. Introduction	16
II. Les types de maintenances	16
III. Les politiques de maintenances	19
IV. L'horizon de temps fini	21
V. Les conditions de fonctionnement	23
VI. Ordonnancement des tâches de production	24
VII. Couplage de la maintenance et de la production	24
VIII. Description du problème	27
IX. Conclusion	30

I. Introduction

Dans un système industriel, chacun est conscient que des défaillances peuvent s'avérer très coûteuses mais également très dangereuses tant au niveau matériel, environnemental et humain. Des exemples tels que le naufrage du Bourbon Dolphin ayant provoqué trois morts et la disparition de cinq personnes en mer, la perte de plusieurs dizaines de tonnes de fuel lourd sur les côtes Ouest de l'Europe par le pétrolier Prestige sont là pour illustrer les conséquences dramatiques qui ont pu se produire.

Dès les années 1960, Barlow et Hunter [Barlow, 1960] furent les premiers à s'intéresser au domaine de la maintenance. À l'époque, l'objectif recherché était principalement une réduction du coût des actions de maintenances correctives. De ce fait, pour pallier l'accroissement de ces coûts liés aux pannes, la théorie de renouvellement fut proposée. Cette théorie aborde la problématique du remplacement des systèmes. Parmi les différentes politiques de remplacement existantes, la plus connue est certainement celle proposée par Barlow et Proschan [Barlow, 1965], qui est basée sur l'âge. Cette dernière se caractérise par le remplacement de l'unité après une durée prédéterminée ou dès l'apparition d'une défaillance, la première des deux situations atteintes. Bien évidemment, les coûts liés au remplacement planifié sont inférieurs à ceux liés au remplacement dû à une défaillance. Bien que d'autres politiques de remplacements existent, la stratégie de type bloc, proposée par Nakagawa [Nakagawa, 1979] est couramment employée. Dans cette stratégie, l'unité est remplacée à une périodicité fixe et lors de chaque défaillance (sans tenir compte de l'âge du composant). Les coûts liés à la stratégie bloc respectent les caractéristiques des coûts liés à la stratégie de type âge. Cette politique de remplacement a permis la détermination d'intervalles de temps optimaux dans le cadre de la minimisation des coûts de maintenance ou de la maximisation de la disponibilité de composants. Dans ces deux stratégies, les remplacements planifiés correspondent en définitive à des actions préventives, c'est-à-dire « exécutées à des intervalles de temps prédéterminés ou selon des critères prescrits, et destinées à réduire la probabilité de défaillance ou la dégradation du fonctionnement d'un bien » [AFNOR, 2001]. Parallèlement, les remplacements suites aux défaillances correspondent en réalité à des maintenances correctives, que la norme [AFNOR, 2001] définit comme « l'ensemble des actions exécutées après détection d'une panne et destinées à remettre un bien dans un état dans lequel il peut accomplir une fonction requise ».

II. Les types de maintenances

Lorsque ces actions de maintenances sont opérées, les résultats post-maintenance peuvent différer suivant le type de maintenance réalisée. Les plus couramment utilisées sont les maintenances minimales et parfaites. Après, sont apparues les maintenances imparfaites pour modéliser un état intermédiaire.

Les actions de maintenances parfaites peuvent être considérées comme des remplacements dans la mesure où le système se retrouve dans un état *As Good As New* (AGAN), littéralement traduisible par « Aussi Bon Que Neuf ». Après une maintenance ou réparation parfaite, la distribution des durées de vie et la fonction du taux de défaillance sont les mêmes que celles d'une nouvelle unité

neuve. Ce modèle est employé pour modéliser l'effet de maintenance corrective et/ou préventive. Mathématiquement, le taux de défaillance en fonction du temps t s'exprime par :

$$\lambda(t) = \lambda_N(t - T_{N_i}) \quad (\text{I. 1})$$

Avec :

- $\lambda(t)$: fonction du taux de défaillance,
- $\lambda_N(t)$: fonction nominale du taux de défaillance,
- T_{N_i} : date de la dernière défaillance.

À l'opposé des maintenances parfaites, nous retrouvons les maintenances minimales. Ces actions de maintenances ont pour but de restituer le système dans un état opérationnel, et plus précisément dans l'état qu'il était juste avant la défaillance, c'est-à-dire *As Bad As Old* (ABAO), qui signifie « Aussi Mauvais Que Vieux ». De ce fait, les maintenances minimales ne sont utilisées que pour modéliser des actions correctives. Puisque ces actions de maintenance n'améliorent pas l'état du système, le taux de défaillance au cours du temps s'exprime comme suit :

$$\lambda(t) = \lambda_N(t) \quad (\text{I. 2})$$

Néanmoins, nous pouvons constater que seules les maintenances parfaites peuvent être utilisées pour modéliser des actions préventives, actions réalisées pour améliorer l'état du système. Il aura fallu attendre approximativement deux décennies pour qu'apparaissent les premiers modèles de maintenances imparfaites. Ce type de maintenance est utilisé pour modéliser l'effet d'une maintenance entre minimal et parfait. Le premier modèle, proposé par Brown et Proschan [Brown, 1983] correspond en réalité à la réalisation des deux modèles de bases (minimal et parfait) mais suivant une probabilité. Dans leur modèle, une maintenance préventive peut être considérée comme parfaite avec une probabilité ρ et minimale avec la probabilité complémentaire, à savoir $(1 - \rho)$. Par ailleurs, ce modèle suppose que les effets des actions de maintenances ne dépendent ni des dates de réalisation, ni des maintenances réalisées précédemment. Il aura fallu attendre 1988 pour qu'apparaissent les premiers modèles de maintenances imparfaites basés sur l'âge virtuel. Cette notion, proposée par Malik [Malik, 1979], suppose qu'après la $i^{\text{ème}}$ maintenance, le système serait équivalent à un système neuf mais ayant fonctionné, sans jamais tomber en panne, durant une période A_i . A_i représente donc l'âge virtuel du système après la $i^{\text{ème}}$ maintenance. Basée sur ce concept, Kijima et al. proposent une première modélisation [Kijima, 1988] suivant laquelle l'efficacité de la maintenance est caractérisée par une réduction de l'âge virtuel, d'une quantité proportionnelle à la durée écoulée depuis la précédente maintenance. Mathématiquement, l'âge du système est donné par :

$$A_i = A_{i-1} + (1 - Z_i) \cdot X_i \quad \forall i \geq 1 \quad (\text{I. 3})$$

Avec :

- X_i : Durée entre la $(i-1)^{\text{ème}}$ maintenance et la $i^{\text{ème}}$ maintenance,
- Z_i : Facteur de réduction de la durée X_i .

Le deuxième modèle, proposé par Kijima [Kijima, 1989], repose sur la même idée. Néanmoins, la réduction d'âge n'est plus proportionnelle à la durée écoulée depuis la dernière maintenance mais à l'âge virtuel lui-même qui peut s'exprimer comme suit :

$$A_i = (1 - Z_i) \cdot (A_{i-1} + X_i) \quad \forall i \geq 1 \quad (\text{I. 4})$$

Ces modèles sont respectivement connus sous le nom de modèles de Kijima de type I et II.

Dans ces modèles, le facteur de réduction Z est compris dans l'intervalle $[0, 1]$, dont nous pouvons distinguer quelques particularités :

Z_i	Efficacité	
	Kijima type I	Kijima type II
$Z_i = 0$	Minimale (ABAO)	
$Z_i = 1$	Optimale	Parfaite (AGAN)
$Z_i \in]0, 1[$	Efficace	
$Z_i < 0$		Nuisible

Tableau 1 : Efficacité des maintenances imparfaites

Dans le même ordre d'idées, Doyen et Gaudoin [Doyen, 2004A] ont développé des modèles à réduction arithmétique de l'âge pouvant tenir compte des instants de défaillances précédents. Ces instants peuvent influencer l'intensité de défaillance du système. L'expression analytique s'exprime donc comme suit :

$$\lambda(t) = \lambda_N \left(t - \rho \cdot \sum_{j=0}^{\text{Min}(m-1, N_i-1)} (1-\rho)^j T_{N_i-j} \right) \quad (\text{I. 5})$$

Avec :

- m : nombres des derniers instants de défaillances mémorisés,
- ρ : facteur de réduction, également appelé facteur d'amélioration.

Lorsque ce modèle fait référence à tous les instants de défaillances ($m = \infty$), nous retrouvons le modèle de Kijima de type II dans le cas où l'efficacité des maintenances est déterministe et constante. Au contraire, lorsque la mémoire fait référence uniquement à la dernière défaillance ($m = 1$), ce modèle correspond à celui défini par Malik.

Une autre approche repose sur une modification du taux de défaillance. Ce modèle semble avoir été développé par Nakagawa [Nakagawa, 1988] et se caractérise par une augmentation de l'intensité de la fonction de taux de défaillance après chaque maintenance préventive. Néanmoins, après chaque action imparfaite, le taux de défaillance est ramené à la valeur zéro comme si, à cet instant précis, le système était considéré comme neuf :

$$\lambda_i(\tau) > \lambda_{i-1}(\tau) \quad \forall i \geq 1 \quad (\text{I. 6})$$

Où $\lambda_i(t)$ représente la fonction du taux de défaillance avant la $i^{\text{ème}}$ maintenance imparfaite et où τ est compris dans l'intervalle $[0, X_i]$.

Une autre approche se base sur la réduction du taux de défaillance. Ainsi, Chan et Shaw [Chan, 1993] proposent de réduire l'intensité de défaillance après une maintenance d'une quantité proportionnelle à sa valeur juste avant la défaillance. Après la $i^{\text{ème}}$ maintenance dont la durée est supposée négligeable, le taux de défaillance s'exprime par :

$$\lambda_{T_i^+} = \lambda_{T_i^-} - \rho \cdot \lambda_{T_i^-} \quad (\text{I. 7})$$

Où T_i^- représente la limite à gauche de l'instant de maintenance et T_i^+ la limite à droite. À l'instar du modèle de Kijima de type I, une maintenance réalisée avec un facteur d'amélioration égal à 1 ne retourne pas le système dans un état *As Good As New*, puisque le taux de défaillance évolue en considérant l'âge réel du système. Dans ce cas, la maintenance est considérée comme optimale.

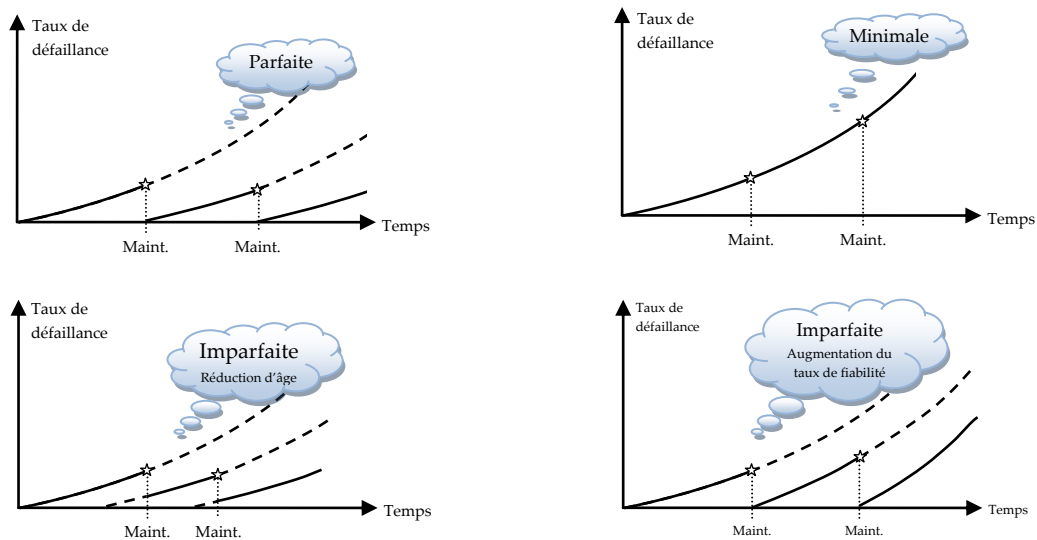


Figure 2 : Évolution de la fiabilité pour les maintenances de type parfait, minimal, et imparfait

Grâce à ces différents types de maintenances, de nombreux travaux de recherches ont porté sur la modélisation et l'étude de politiques de maintenances.

III. Les politiques de maintenances

Avant l'apparition des maintenances imparfaites, les politiques de maintenances consistaient essentiellement en des politiques de remplacement. De nos jours, de nombreux travaux portent encore essentiellement sur des politiques de maintenances constituées uniquement de maintenances minimales et parfaites. Parmi ces travaux, nous pouvons citer [Dellagi, 2006], [Hajjej, 2009]. Les maintenances imparfaites ont ainsi permis l'émergence de nouvelles politiques de maintenances. Les plus couramment employées sont les politiques périodiques et séquentielles [Wang, 2002], dont les premiers travaux ont été effectués par Nakagawa [Nakagawa, 1986].

Dans une politique de maintenance de type périodique, l'unité est maintenue préventivement à intervalles réguliers, soit toutes les $k \cdot X$ ($k = 1, 2, \dots, N - 1$) unités de temps, avant d'être remplacée au bout de $N \cdot X$ unités de temps. Les maintenances périodiques sont considérées comme imparfaites et les maintenances correctives comme minimales. Habituellement, les durées de ces actions de maintenances sont considérées comme négligeables. Pour cette politique, le principal objectif consiste à déterminer le nombre optimal d'intervalles de maintenances préventives N^* , mais également la durée optimale X^* de ces intervalles. L'objectif généralement recherché par cette politique de maintenance réside dans la minimisation des coûts de maintenance [Nakagawa, 1986]. Différentes contraintes peuvent être ajoutées, comme, par exemple, maintenir un seuil de disponibilité minimal [Schutz, 2008C]. De nos jours, cette politique de maintenance est encore communément employée pour des problèmes faisant référence aux garanties [Liu, 2008], aux équipements de location [Yeh, 2009], ou à la planification de *jobs* dans des ateliers de type *flowshop* ou *jobshop*, en vue de minimiser le *makespan* [Ji, 2007]. Cette politique de maintenance est généralement utilisée dans le cas d'horizon infini. Dans ce cas, des activités de maintenances imparfaites sont réalisées périodiquement entre les remplacements [Nagakawa, 2005].

La politique séquentielle se différencie de la politique périodique par une variation de la taille des intervalles de temps séparant deux actions de maintenances préventives consécutives. Alors que dans la politique périodique, l'unité était maintenue préventive à intervalles fixes, la politique séquentielle se caractérise par des tailles d'intervalles variables. L'objectif consiste donc à déterminer le nombre optimal de maintenances préventives à réaliser, c'est-à-dire $(N-1)^*$, avant d'effectuer le remplacement de l'unité. Cependant, il est également nécessaire de déterminer la taille optimale de ces N intervalles de temps, notée $\{X_1^*, X_2^*, \dots, X_N^*\}$. Cette politique de maintenance est plus réaliste dans la mesure où les maintenances préventives sont imparfaites. En effet, après chaque maintenance préventive, l'état du système se trouve dans un état entre AGAN et ABAO. Il convient donc de réaliser des actions de maintenances de manière plus rapprochée afin de minimiser le risque de défaillance. Comparée à la politique précédente, cette politique de maintenance est assez peu étudiée car le nombre de variables de décisions est plus important et de plus, elles sont corrélées. Le nombre de variables de décisions dépend du nombre de maintenances préventives réalisées.

Depuis peu de temps, une nouvelle politique de maintenance semble émerger. Cette politique, basée sur la fiabilité, ne se caractérise pas réellement par la recherche des intervalles de temps optimaux, mais par la recherche d'un seuil de fiabilité optimal. La détermination du seuil de fiabilité implique explicitement la connaissance des durées séparant les maintenances préventives. À ce titre, Zhou et al. [Zhou, 2007] cherchent le seuil de fiabilité optimal minimisant les coûts de maintenance en considérant des maintenances préventives imparfaites hybrides. Ces maintenances sont appelées hybrides car leur effet se caractérise par une réduction de l'âge virtuel et par une augmentation de la fonction du taux de défaillance comme défini dans [Lin, 2000]. L'approche proposée par Cheng et al. [Cheng, 2007] varie dans la mesure où un remplacement est effectué dès lors que la fiabilité atteint le seuil préalablement fixé. Dans leur premier modèle, le système reçoit des actions de maintenances préventives périodiquement, soit toutes les $k \cdot X$ ($k = 1, 2, \dots, N - 1$) unités de temps, et le système est remplacé dès lors que la fiabilité atteint ce seuil fixé. Le second modèle proposé, sous le nom de politique périodique partielle, consiste à déterminer dans un premier temps le nombre d'intervalles de maintenances préventives à réaliser ainsi que la longueur optimale de ces intervalles en relaxant le seuil de fiabilité. Par la suite, la contrainte portant sur le seuil de fiabilité est réintégrée, le nombre

d'intervalles de maintenances préventives est adapté, et la durée du dernier intervalle est réduite pour atteindre le seuil de fiabilité. Pour résumer, ce modèle sera constitué de $(k-1)$ intervalles de durées X et la durée du dernier intervalle est inférieure ou égale à X .

Malheureusement, ces différentes politiques se basent toutes sur la même hypothèse. Les différents modèles mathématiques proposés pour ces politiques de maintenance partent du principe que l'horizon de temps est considéré comme infini. Bien que cette hypothèse soit admissible dans de nombreux cas, il n'en demeure pas moins que pour d'autres situations, cette hypothèse ne peut pas être satisfaisante. Par exemple, l'obsolescence de certains équipements ou structures nécessite de prendre en considération cette échéance.

IV. L'horizon de temps fini

Tous les travaux cités précédemment ainsi que la quasi-totalité des travaux de recherches concernant l'optimisation de politiques de maintenances considèrent un horizon de temps infini. Cette hypothèse est toujours utilisée lorsque le système se situe dans sa période de maturité et que les opérations de production sont constantes au cours du temps. Habituellement, l'objectif recherché consiste en une minimisation des coûts de maintenance. L'horizon de temps étant infini, le but recherché est la détermination des intervalles de temps entre maintenances imparfaites et renouvellement en déterminant le coût moyen de maintenance par unité de temps.

Jayabalan [Jayabalan, 1992A] semble être le premier à considérer une période de planification fixée. Dans ses premiers travaux, des actions préventives imparfaites et/ou des remplacements sont réalisés dès lors que le taux de défaillance du système atteint un seuil maximal défini. L'objectif consiste alors à déterminer le nombre d'interventions préventives, à réaliser entre chaque remplacement, ainsi que le nombre de remplacements durant l'horizon de temps ; le but étant de minimiser les coûts totaux. Cependant, le temps de calcul dépendant de la taille de la période de planification, Jayabalan [Jayabalan, 1992B] propose une heuristique pour déterminer une solution satisfaisante. L'algorithme proposé se base sur l'affirmation que la durée prévisionnelle avant la prochaine défaillance d'un système maintenu à l'aide de maintenances préventives imparfaites est plus courte que pour un système neuf.

Les travaux de Dedopoulos [Dedopoulos, 1998] se basent également sur un horizon de temps fini. Comme bien souvent, l'objectif exposé dans son article consiste à déterminer le nombre et les dates de réalisation de maintenances préventives pour une politique de type séquentiel. Ces maintenances préventives sont considérées comme imparfaites et sont modélisées par une réduction de l'âge basée sur un facteur dépendant de l'intervalle de maintenance. Un profit étant généré pour chaque unité de temps, le but consiste à maximiser le bénéfice durant l'intervalle alloué. Bien évidemment, les actions de maintenances préventives nécessitent une durée pendant laquelle aucun profit n'est généré. De plus, le coût de ces actions préventives est variable ; il dépend du facteur de réduction d'âge, de la durée de l'action de maintenance ainsi que de l'âge du système. Des actions correctives minimales de durée négligeable et de coûts fixes sont appliquées aux défaillances qui interviennent entre ces actions préventives. Compte tenu de la complexité du modèle obtenu, une résolution numérique est proposée. Le principe consiste à incrémenter le nombre d'actions de maintenance préventive et pour chaque cas, déterminer les dates optimales de maintenance préventive ainsi que

les facteurs optimaux de réduction d'âge. Une extension est également proposée à ces travaux ; il s'agit de garantir une fiabilité minimale au système. Cette contrainte peut se justifier par la production de produits intermédiaires qui ne peuvent être stockés et qui se détériorent après une courte période.

Récemment, de nombreux modèles de politiques de maintenances ont été adaptés aux horizons de temps fini. Le premier modèle abordé par [Nakagawa, 2005] porte sur le remplacement basé sur l'âge. Dès 1996, Legát et *al.* [Legát, 1996] avaient proposé une approximation pour déterminer l'âge optimal de remplacement, mais uniquement valable dans le cas d'une distribution de type Weibull. Cette approximation a été établie à partir de la solution issue de l'horizon de temps infini. Dans son ouvrage, Nakagawa a adapté la formule de l'horizon de temps infini à l'horizon de temps fini sans considérer de distribution particulière, après avoir prouvé l'existence d'une solution optimale. La deuxième politique de remplacement étudiée (remplacement de type bloc) se différencie de la politique précédente par le remplacement des éléments défaillants au cours du temps. La dernière politique de remplacement se caractérise par un remplacement périodique. Dans ce cas, le système est remplacé toutes les $k \cdot T$ ($k = 1, 2, \dots, N = S/T$) unités de temps où S représente la durée totale de l'horizon de temps. Dès lors que le système tombe en panne entre ces périodes de remplacement, un coût de non fonctionnement (par unité de temps) est engendré et le système reste dans un état non opérationnel. Ces trois politiques de remplacement peuvent être résolues à l'aide de la méthode dite de « partition ». Différents modèles d'inspection sont également présentés à travers cet ouvrage. L'inspection périodique se rapproche des travaux concernant le remplacement périodique dans la mesure où des inspections sont réalisées à des intervalles constants. La seconde politique d'inspection est dite séquentielle. Elle se caractérise par des intervalles de temps variables, qui dans ce cas, diminuent au cours du temps. Nakagawa et Mizutani [Nakagawa, 2009] ont complété cette étude par l'adaptation aux horizons de temps fini de politiques de maintenances préventives imparfaites. Après ces maintenances, le taux de défaillance du système est ramené à zéro, mais il augmente plus rapidement par rapport à l'intervalle précédent. La première politique étudiée est de type périodique. Des actions préventives imparfaites sont réalisées toutes les $k \cdot T$ ($k = 1, 2, \dots, N - 1$) unités de temps et le système est remplacé au bout de $S = N \cdot T$ unité de temps. Durant cet intervalle de temps, des actions minimales sont réalisées pour palier les défaillances. L'existence d'une solution optimale étant prouvée, le nombre de maintenances préventives peut être déterminé en incrémentant la variable de décision N . Par conséquent, la durée des intervalles de maintenance préventive est donnée par $T = S / N$. La seconde politique étudiée est de type séquentiel, c'est-à-dire que les intervalles de maintenances sont variables. Afin d'apporter des solutions approximatives pour cette politique séquentielle, [Jiang, 2009] propose un algorithme simple, sans itération et applicable à l'ensemble des distributions probabilistes. Enfin, pour l'ensemble des modèles proposés précédemment, les durées d'inspections, de maintenances correctives et préventives ainsi que de remplacement sont considérées comme négligeables.

De manière générale, l'ensemble des travaux abordant l'hypothèse d'horizon de temps fini conserve encore certaines hypothèses simplificatrices (durées négligeables, taux de défaillance continu entre les maintenances préventives, etc.) qui, si elles étaient relaxées, complexifieraient considérablement la résolution de ces politiques de maintenance.

V. Les conditions de fonctionnement

Comme ce fut le cas pour les horizons de temps fini, très peu de travaux abordent la problématique de conditions de fonctionnement variables au cours du temps. Dans tous les travaux cités précédemment, les auteurs ont toujours supposé que les systèmes fonctionnaient sous des conditions de fonctionnement constantes au cours du temps.

La première contribution semble résulter des travaux de [Özekici, 1995]. Dans cet article, l'auteur s'intéresse à deux politiques de maintenance. Il définit un intervalle comme la durée pendant laquelle les conditions de fonctionnement sont constantes. Dès que ces conditions changent, un nouvel intervalle débute. Pour la première politique, au début de chaque nouvel intervalle, un choix est proposé permettant le remplacement ou non du système. Si le système est défaillant au moment de l'inspection, c'est-à-dire au début de l'intervalle, le système est remplacé par un nouveau et le coût engendré est supérieur au coût du remplacement préventif. La deuxième politique proposée permet de nuancer l'action entreprise au début de chaque nouvel intervalle. En plus du remplacement, des actions de maintenance peuvent être réalisées, ramenant le système dans un état situé entre AGAN et ABAO. Le coût de ces actions de maintenance dépend de l'effet de la réparation. Dans cet article, le processus de défaillance, et donc le taux de défaillance, dépendent des conditions de fonctionnement. Par conséquent, l'auteur introduit la notion d'âge intrinsèque plutôt que d'âge réel. Il définit l'âge intrinsèque par : *"If the real time is t , then the intrinsic clock show time $R_i(t)$ [...] In other words, it takes $R_i^{-1}(x)$ units of real time operation to age a brand new device to intrinsic age x in environment i ",* où $R_i(\cdot)$ et $R_i^{-1}(\cdot)$ représentent respectivement la fonction de fiabilité et la fonction inverse de fiabilité associées à l'intervalle i .

Dans les travaux de Martorel et al. [Martorell, 1999], l'influence des conditions de fonctionnement est introduite directement au sein du modèle de fiabilité du système. Le modèle de vie accélérée (Accelerated Life Model - ALM) et le modèle à risques proportionnels (Proportional Hazards Model - PHM) sont les deux outils considérés pour modéliser l'impact des conditions de fonctionnement. Ces deux modèles seront présentés plus en détail au cours du deuxième chapitre. Dans leurs travaux, les auteurs supposent que les variations des conditions de fonctionnement ne peuvent survenir qu'aux instants de maintenances préventives et que, par conséquent, entre deux actions préventives successives, les conditions de fonctionnement restent inchangées. Contrairement aux autres travaux présentés où l'intérêt consistait à maximiser ou minimiser une fonction objectif, leur objectif vise plutôt à comparer deux modèles de maintenances préventives imparfaites, à savoir les modèles PAR (*Proportional Age Reduction*) et PAS (*Proportional Age Setback*).

Concernant les conditions de fonctionnement, nous pouvons constater qu'au moment où des variations surviennent, des actions de maintenances sont entreprises. Aucun auteur ne semble considérer la situation où les maintenances préventives sont indépendantes des conditions de fonctionnement.

Les travaux présentés jusqu'à présent concernent uniquement la planification des activités de maintenance. Néanmoins, l'application de ces dernières suppose l'existence de tâches de production. Dans la suite de cette étude bibliographique, nous nous intéresserons à l'optimisation de l'ordonnancement des tâches de production, ainsi qu'au couplage entre ces deux grands domaines que sont la production et la maintenance.

VI. Ordonnancement des tâches de production

Avec la planification de la maintenance, la planification de la production est l'une des tâches les plus importantes dans un système manufacturier. Elle consiste en la programmation des tâches, également appelées *jobs*, au niveau des machines, et notamment la spécification de la séquence de réalisation ainsi que le temps nécessaire à l'accomplissement de ces dernières. De nombreux systèmes de production présentent une configuration particulière. Cette dernière implique que la réalisation des tâches passe toujours par les mêmes machines et suivant le même ordre. Ce type de configuration, appelé *permutation flowshop* (PFSP), a été largement étudié dans la littérature [Pinedo, 2005]. Depuis la présentation du problème de flowshop par Johnson [Johnson, 1954], de très nombreux algorithmes et techniques ont été développés, allant de simples heuristiques jusqu'à des méta-heuristiques modernes et complexes, sans oublier les méthodes exactes, telles que la programmation en nombre entier [Coleman, 1992] et les méthodes de Branch and Bound [Brucker, 1996]. Parmi les techniques les plus connues, issues de l'étude comparative de [Ruiz, 2005], nous pouvons citer l'heuristique NEH [Nawaz, 1983] considérée comme l'une des meilleures méthodes de résolution pour le PFSP, ainsi que l'heuristique améliorée de Suliman [Suliman, 2000]. Parmi les méta-heuristiques utilisées, nous pouvons notamment évoquer les travaux de [Ruiz, 2006] et [Rajendran, 2004] pour l'utilisation respective des algorithmes génétiques et des colonies de fourmis. Actuellement, les travaux concernant l'ordonnancement prennent en compte de nouvelles contraintes, telles que les durées de réglages, les précédences [Balas, 2008], et essaient d'y associer des actions de maintenance.

VII. Couplage de la maintenance et de la production

Depuis quelques années, des problèmes, liés à la planification des tâches et intégrant des actions de maintenances préventives, émergent. Parmi ces nouveaux travaux, l'un des plus connus est celui développé par Cassady et Kutanoğlu [Cassady, 2003]. Les auteurs considèrent qu'une seule machine est utilisée pour accomplir différents *jobs*. L'objectif de leur article consiste à proposer un modèle analytique intégré (production et maintenance) afin de minimiser le retard total pondéré (*total weighted tardiness*). Pour minimiser ce retard, ils utilisent des maintenances préventives périodiques, dont la durée optimale des intervalles est déterminée dans l'objectif de maximiser la disponibilité du système. La résolution étant obtenue par une énumération totale, et par conséquent le nombre de *jobs* étant restreint, cette approche est difficilement applicable. Néanmoins, l'intégration de la planification de la production et de la maintenance, génère une économie d'environ 30% du retard total pondéré, comparé à l'étude indépendante de l'ordonnancement des *jobs* et de la planification de la maintenance. Pour ce même type de problème, Sortrakul et al. [Sortrakul, 2005] proposent une résolution basée sur un algorithme évolutionnaire. Dans leur étude, ils considèrent que la politique de maintenance est constituée de maintenances préventives pouvant être réalisées avant chaque *job*. L'algorithme génétique est utilisé pour déterminer la séquence optimale des différents *jobs*, mais également pour choisir les *jobs* précédés de maintenances préventives. Les solutions sont représentées grâce à deux chromosomes. Le premier chromosome est constitué de n gènes, variables entières correspondant à l'ordonnancement des n -*jobs*. Le second chromosome comporte n variables binaires caractérisant si une maintenance préventive doit précéder le *job* associé. Ainsi, si le $i^{\text{ème}}$ gène du chromosome associé à la maintenance est égal à 1, le $i^{\text{ème}}$ *job* sera précédé d'une maintenance

préventive ; dans le cas contraire, aucune maintenance préventive ne sera réalisée entre les *jobs* $i-1$ et i .

Au niveau du couplage de la production et de la maintenance, Benbouzid-Sitayeb et *al.* [Benbouzid-Sitayeb, 2006] comparent les stratégies séquentielle et intégrée. La stratégie séquentielle consiste en la détermination du plan de production optimal, puis de considérer ce dernier comme une contrainte forte pour déterminer un plan de maintenance adapté. La stratégie séquentielle consiste, quant à elle, à déterminer simultanément le plan de production et de maintenance. Cependant, l'utilisation d'une politique périodique ne permet pas réellement la détermination conjointe car ces maintenances ne sont pas contraintes par les tâches de production, mais par leur périodicité. Pour y remédier, les auteurs utilisent un algorithme évolutionnaire (colonies de fourmis) qui permet donc d'obtenir des solutions satisfaisantes.

Dans Ruiz et *al.* [Ruiz, 2007], les auteurs considèrent la réalisation de n jobs sur m machines. Le couplage entre la planification de la production et les actions de maintenances préventives est réalisé de manière intégrée. Dans ce papier, deux politiques de maintenances sont étudiées : celle de Cassady et Kutanoğlu [Cassady, 2003] qui consiste à déterminer l'âge optimal de maintenance préventive, dans le but de maximiser la disponibilité. Concernant la seconde politique de maintenance étudiée, il s'agit de maintenir un seuil de fiabilité minimal pour le système. Les auteurs ont adapté six techniques pour résoudre ce problème intégré suivant les deux politiques de maintenances. Les méthodes utilisées sont l'heuristique NEH ; pour les méta-heuristiques classiques, ils se sont basés sur les outils de recherche tabou de [Widmer, 1989], connus sous le nom de SPIRIT, et le recuit simulé de [Osman, 1989], également dénommé SA_OP. Concernant les algorithmes évolutifs, il s'agissait de l'algorithme génétique de Ruiz et *al.* [Ruiz, 2006], appelé GARMA et de l'algorithme à colonie de fourmis, PACO [Rajendran, 2004].

En dehors des problèmes de type *flowshop*, de nombreux travaux de recherche se sont intéressés à l'intégration entre la production et la maintenance en considérant le cas d'une seule machine produisant un seul type de produit. De nombreuses études ont démontré l'efficacité de l'aspect intégré entre le domaine de la production et celui de la maintenance. Nous pouvons citer, entre autres, les travaux menés par Srinivasan et Lee [Srinivasan, 1996] qui ont analysé les effets combinés entre les politiques de production et de maintenance préventive d'un point de vue monétaire. Cheung et Hausmann [Cheung, 1997] ont cherché à optimiser simultanément un stock avec une politique de maintenance de type âge. Plus récemment, Rezg et *al.* [Rezg, 2004] ont présenté une optimisation conjointe entre une politique de maintenance préventive et la gestion d'un stock, dans le cas d'une ligne de production constituée de N machines. Kenné et Gharbi [Kenné, 2004] ont étudié l'optimisation stochastique de la gestion de production couplée à des activités de maintenances corrective et préventive. Ils ont proposé une méthode permettant de déterminer de manière optimale la périodicité des maintenances préventives et du taux de production pour un système de production constitué de plusieurs machines identiques. Dans les travaux menés par Aghezzaf et *al.* [Aghezzaf, 2007], l'objectif consistait à déterminer, conjointement, les plans de production et de maintenance de manière à minimiser les coûts totaux. Dans cette étude, l'horizon de temps considéré est fini et ce dernier est constitué de N sous-intervalles de tailles identiques où, pour chacun d'entre eux, plusieurs produits peuvent être requis. Le plan de maintenance est constitué de maintenances correctives minimales réalisées lors de l'apparition de défaillances et de remplacement réalisé périodiquement. Récemment, Hajje et *al.* [Hajje, 2009] se sont intéressés à l'intégration entre

un plan de production et un plan de maintenance associés à un système manufacturier. L'optimisation conjointe est faite en vue d'établir, de manière optimale, un plan de production et une planification de maintenance. L'originalité de ce papier réside dans l'influence du taux de production sur le degré de dégradation. En effet, le taux de production varie tout au long de l'horizon de temps puisque les demandes sont hétérogènes, c'est-à-dire que les quantités désirées varient suivant les mois. D'après le modèle proposé, la vitesse de dégradation du système est proportionnelle à la cadence de production. L'objectif de leur travail consiste à satisfaire la demande, avec un seuil minimum de service, et à minimiser les coûts de stockage, production et maintenance, en déterminant le plan de production optimal, sous entendues les cadences optimales pour chaque mois. Concernant le plan de maintenance, la politique adaptée consiste en la réalisation de maintenances préventives parfaites. Le nombre de maintenances à réaliser est déterminé en fonction du plan de production associé.

Depuis peu de temps, des travaux concernant le couplage entre la gestion de la production et la gestion de la maintenance ont intégré de nouvelles contraintes, telles que des tâches de sous-traitance, ou encore la notion de contrôle de la qualité.

Dans les travaux de Dahane [Dahane, 2007], l'objectif consiste à déterminer des politiques de maintenances intégrées (considérant la maintenance et la production) sous contrainte de sous-traitance dans le cas d'entreprises prestataires de sous-traitance. À partir d'une politique de maintenance simple, dissociant la gestion de maintenance de la gestion de la production, il a déterminé l'âge optimal de remplacement puis, à partir de ces données, le seuil optimal de stock. Par la suite, il a cherché à optimiser simultanément ces deux variables de décisions au sein d'une politique de maintenance dite intégrée. En se basant sur cette politique intégrée, il a défini analytiquement les conditions nécessaires à partir desquelles il est économiquement intéressant d'allouer une machine à des tâches de sous-traitance. La contrainte de sous-traitance impose des périodes d'indisponibilité des machines de façon périodique afin d'accomplir des tâches de sous-traitance. Une extension de cette étude a été proposée dans le cas de deux machines. Par la suite, il s'est intéressé aux aspects temporels de la sous-traitance. Par rapport à ces aspects temporels, il a défini un modèle analytique permettant de déterminer la date optimale d'allocation des ressources aux tâches de sous-traitance, tout en montrant l'impact du placement d'une tâche de sous-traitance, tant au niveau du retard en début de période de sous-traitance qu'en prolongation inattendue de la période de sous-traitance. En se basant sur les résultats obtenus, il a proposé une politique généralisée de maintenance intégrée qui consiste en une généralisation basée sur le nombre de sous-traitances à réaliser pendant un cycle. La dernière partie de leurs travaux concerne l'analyse d'un point de vue quantitatif. Cet aspect prend en considération les différentes politiques de maintenances soumises à des contraintes par rapport au nombre de tâches de sous-traitance ainsi que la manière dont la machine est allouée à la sous-traitance.

Radhoui [Radhoui, 2008] aborde la problématique du couplage production/maintenance en considérant le contrôle de la qualité pour des systèmes de production pouvant générer des unités non-conformes au-delà d'un certain degré aléatoire de production. En considérant une demande constante au cours du temps et en supposant que le passage à un état de production non-conforme est aléatoire, un modèle analytique a été proposé afin de déterminer les valeurs optimales de la taille de lot ainsi que l'âge des maintenances préventives parfaites. Dans ce modèle, le système produit à sa cadence maximale, jusqu'à atteindre un seuil d'inventaire fixé, puis la cadence de production est

ramenée à la cadence des demandes. La production de pièces non conformes apparaît de façon aléatoire et engendre automatiquement le passage en phase de restauration. La taille du lot déterminée a pour objectif de compenser les produits non conformes pendant la phase de préparation de la restauration, de durée connue, ainsi que la consommation moyenne durant la phase de restauration, dont la durée est aléatoire. Dans la suite de cette étude, un nouveau modèle est développé, permettant d'intégrer le dimensionnement du stock tampon, correspondant à la taille de lot, et du contrôle qualité, c'est-à-dire le seuil du taux de non conformité. Le stock tampon est utilisé pour satisfaire les demandes pendant les phases d'arrêts liés aux actions correctives et préventives, ainsi que pour compenser les produits non-conformes. Contrairement au modèle précédent, la production de pièces non-conformes augmente avec la dégradation du système. Suivant la valeur du taux de non-conformité, des maintenances préventives ou de révisions sont appliquées. La suite de ces travaux repose sur ce même modèle à la différence que tous les lots ne sont pas automatiquement inspectés. Une nouvelle variable de décision est ajoutée et permet de déterminer la périodicité optimale d'inspection. Enfin, le dernier modèle suppose que les actions de maintenances préventives sont considérées comme imparfaites. La qualité de ces dernières se représente par une diminution du taux de non-conformité, entre 0 (la machine est dans un état *As Good As New*) et le taux actuel.

Malheureusement, dans l'ensemble des articles traitant de l'intégration de la production et de la maintenance, les politiques de maintenances préventives restent simples, dans la mesure où ces actions sont de type parfait, restaurant le système dans un état neuf.

VIII. Description du problème

La problématique soulevée dans ce mémoire concerne un navire, qui peut être assimilé à un système de production, qui doit s'engager en mer pour une durée définie, c'est-à-dire pour un horizon de temps donné. Durant cette période, le navire doit réaliser un ensemble de missions (*jobs*) de sorte à maximiser les bénéfices engendrés par chaque mission. Cependant, la réalisation de missions engendre un processus de dégradation du système qui peut entraîner des défaillances immobilisantes dont les coûts de maintenances viennent en déduction des bénéfices générés. Par ailleurs, les missions à réaliser étant situées dans des lieux différents (hémisphère nord, sud, etc.) et les missions étant de caractéristiques différentes, elles impactent différemment le processus de dégradation du navire.

L'objectif posé par cette thèse consiste donc à déterminer, avant le départ du navire, un plan d'exploitation, c'est-à-dire un ensemble ordonné de missions à réaliser et à définir une politique de maintenance adaptée dans le but de maximiser les bénéfices minorés des coûts liés aux actions de maintenances. La loi de défaillance, associée à la politique de maintenance, devra évidemment prendre en considération les caractéristiques des différentes missions du plan d'exploitation.

En plus du plan d'exploitation, un plan de maintenance devra être déterminé conjointement de manière à maximiser le bénéfice total, c'est-à-dire la totalité des profits générés par les missions réalisées, minorées des coûts liés aux actions de maintenances tant correctives que préventives. De ce fait, différentes politiques de maintenances seront étudiées tout au long du manuscrit.

Avant de continuer, il est nécessaire de définir les principales notations qui seront utilisées au cours du manuscrit.

VIII.1. Notations

Les principales notations utilisées sont données ci-dessous :

- H : Horizon de temps alloué, durée entre le départ et le retour à quai, exprimée en unité de temps,
- $\mathcal{M} = \{1, 2, \dots, M\}$: Ensemble des M missions proposées,
- $\mathcal{X} = \{x_1, x_2, \dots, x_k\}$: Vecteur des k variables x_j qui caractérise le plan d'exploitation,
- $x_j (\in \mathcal{M})$: Variable (naturelle) précisant la $j^{\text{ème}}$ mission à réaliser dans le plan d'exploitation,
- $\dim(\cdot)$: dimension (nombre d'éléments) du vecteur concerné,
- $\Gamma_x(\cdot)$: Coût de maintenance lié la politique x , exprimé en unité monétaire,
- $\phi_x(\cdot)$: Nombre moyen de pannes pouvant survenir suivant la politique x ,
- $\rho \in]0, 1[$: Facteur d'amélioration pour la maintenance préventive imparfaite à réduction d'âge,
- M_C : Coût d'une action de maintenance corrective, exprimé en unité monétaire,
- M_P : Coût d'une action de maintenance préventive, exprimé en unité monétaire,
- μ_p : Durée d'une action de maintenance préventive, exprimée en unité de temps
- $\lambda_0(t)$: Fonction du taux nominal de défaillance,
- $R_0(t)$: Fonction de fiabilité nominale.

Une mission m est définie par un ensemble de caractéristiques :

- δ_m : durée de la mission m , exprimée en unité de temps,
- ψ_m : profit généré par la mission m , exprimé en unité monétaire,
- Z_m : Vecteur précisant les variables explicatives (conditions de fonctionnement) de la mission m ,
- z_m^{CO} : variable précisant les conditions opérationnelles de la mission m ,
- z_m^{CE} : variable précisant les conditions environnementales de la mission m ,
- $g(Z_m)$: Fonction de risque de défaillance pour la mission m ,
- $\lambda_m(t) = g(Z_m) \cdot \lambda_0(t)$: Fonction du taux de défaillance associée à la mission m ,
- $R_m(t)$: Fonction de fiabilité associée à la mission m ,
- $R_m^{-1}(t)$: Fonction inverse de fiabilité associée à la mission m .

VIII.2. Définitions

- Les composantes z_m^{CO} et z_m^{CE} , qui modélisent les conditions opérationnelles et environnementales de la mission m , sont représentées par une valeur située dans l'intervalle $[0, 10]$. Cette note caractérise l'impact de ces conditions sur la fiabilité du système. Puisque ces conditions de fonctionnement seront pondérées par la suite, le choix arbitraire de cet

intervalle n'implique aucune influence pour le reste de l'étude. D'après l'intervalle indiqué, les conditions de fonctionnement seront considérées comme nominales lorsqu'elles prendront la valeur 0 et comme extrêmes pour la valeur 10. Bien évidemment, les conditions de fonctionnement pourront prendre toutes les valeurs réelles comprises dans cet intervalle, comme l'illustre la figure suivante :



Figure 3 : Intervalle des conditions de fonctionnement

- Les conditions de fonctionnement associées aux missions sont estimées par des experts du domaine naval. Les conditions opérationnelles correspondent aux sollicitations subies par le navire c'est-à-dire, les changements entre le fonctionnement en régime constant, les accélérations, les décélérations, leurs durées, etc. Comme le laisse supposer leur nom, les conditions environnementales dépendent de l'environnement. Ainsi, le lieu de réalisation, et par conséquent la salinité de l'eau, les ondes électromagnétiques ont des influences sur la fiabilité du système.

Pour la suite de ces travaux, certaines hypothèses ont été posées.

VIII.3. Hypothèses de travail

Les hypothèses présentées ci-dessous régissent les travaux menés dans cette thèse. Au cours du manuscrit, certaines d'entre-elles seront relaxées.

- Le navire est considéré comme un système monolithique. Sa structure ne sera pas prise en considération.
- Les ressources, tant d'un point de vue humain que matériel, seront considérées comme suffisantes.
- Les durées de changement de missions seront considérées comme négligeables et n'auront, par conséquent, aucune influence sur la dégradation du système.
- Seules les défaillances entraînant une immobilisation et un arrêt de fonctionnement du navire seront considérées.
- Les maintenances correctives seront considérées comme minimales ; le seul but étant de rétablir le système dans un état opérationnel. De ce fait, les durées associées à ces actions seront considérées comme négligeable.
- Les maintenances préventives réalisées seront imparfaites, car effectuées en mer dans des conditions non optimales. Par ailleurs, elles seront de durées non négligeables.
- L'influence des missions, ou plus précisément des conditions de fonctionnement, est indépendante du temps.
- Lorsque le navire quitte le quai, son état général est considéré comme *As Good As New*.
- Chaque mission ne peut être réalisée qu'une seule fois au maximum.

IX. Conclusion

Au cours de ces dernières décennies, il a été démontré que la fonction maintenance jouait un rôle aussi important que celui de la production. L'état de l'art retrace brièvement les principaux types de maintenances qui sont couramment utilisés dans la littérature (maintenances parfaites, minimales et imparfaites) ainsi que les trois principales politiques de maintenances rencontrées (politiques périodiques, séquentielles et basées sur la fiabilité). Par ailleurs, l'étude de ces différents types et politiques de maintenance nous a permis de remarquer que la quasi-totalité des travaux menés concerne des horizons de temps infini, pour lesquels le système est considéré dans un mode de fonctionnement constant au cours du temps. Une partie a été consacrée au problème de l'ordonnancement des tâches de production dans le cas d'atelier de type *flowshop* ou *jobshop* et il a été montré que le couplage entre le domaine de la production et de la maintenance était un sujet en pleine expansion.

Dans la suite de ce chapitre, nous avons exposé la problématique et nous avons pu remarquer qu'elle permettra d'apporter une contribution par rapport aux lacunes présentes dans ce domaine de recherche, et principalement, au niveau des horizons de temps finis et des conditions de fonctionnement variables au cours du temps.

Chapitre II

Modélisation de la loi de défaillance

À travers le chapitre précédent, il a été montré qu'une des principales lacunes, au niveau de notre problématique, correspondait aux conditions de fonctionnement qui n'étaient pas constantes au cours du temps. Pour pallier cette difficulté, la loi de défaillance régissant l'évolution des dégradations au sein du système doit être modélisée en tenant compte des missions à réaliser. Cette modélisation est effectuée en intégrant une fonction de risque de défaillance à la loi de dégradation. Face au modèle sélectionné, un raisonnement par logique floue sera développé pour adapter le modèle à notre problématique.

Sommaire

I. Introduction	32
II. Les différentes approches de pronostics	32
III. Modèle de vie accélérée	34
IV. Modèle à risques proportionnels	35
V. Formalisation de la fonction de risque	36
VI. Structure d'un raisonnement par logique floue	37
VII. Application à notre problématique	41
VIII. Conclusion	50

I. Introduction

Dans les travaux portant sur l'optimisation de politiques de maintenance, les auteurs supposent généralement que le système fonctionne en régime constant sous des conditions de fonctionnement fixées. Bien que cette hypothèse puisse être acceptable dans de nombreux cas de systèmes manufacturiers, il paraît évident que dans cette étude, elle ne peut être retenue. En effet, il est indéniable que la dégradation du navire dépend du type de missions (conditions opérationnelles) ainsi que du lieu de réalisation (conditions environnementales). En d'autres termes, la loi de défaillance doit permettre de tenir compte des conditions de fonctionnement des missions futures et ainsi de pouvoir pronostiquer l'état futur du système. D'ailleurs, Lebold et Thurston [Lebold, 2001] ont proposé de définir la fonction de pronostic comme *"The primary function of the prognostic is to project the current health state of equipment into the future taking into account estimates of future usage profile"*. Pour Byington et al. [Byington, 2002], le pronostic se caractérise par *"Prognostic is the ability to predict the future condition of a machine based on the current diagnostic state of the machinery and its available operating and failure history data"*.

II. Les différentes approches de pronostics

Pour évaluer l'état futur du système, l'approche de pronostic peut être réalisée suivant différentes formes, selon la précision souhaitée et/ou suivant les modèles et dispositifs disponibles. A ce titre, des travaux concernant les différentes approches ont été proposés par Byington et al. [Byington, 2003].

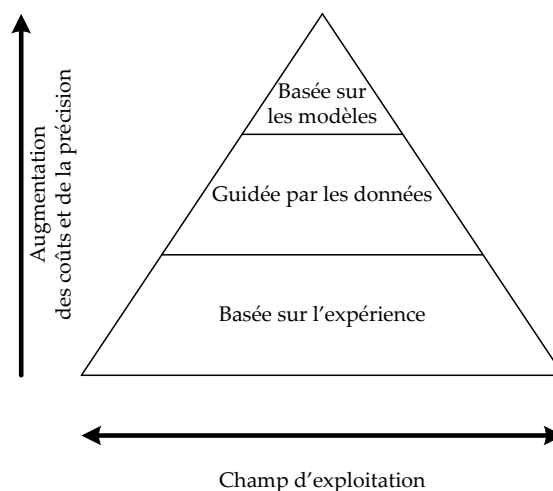


Figure 4 : Les trois approches de pronostic d'après [Byington, 2003]

Comme l'illustre la figure 4, Byington a défini le pronostic suivant trois approches. Nous nous proposons de les présenter succinctement ci-après :

II.1. Le pronostic basé sur un modèle physique

L'approche de pronostic basée sur le modèle se situe en haut de la pyramide. Elle permet donc de déterminer avec précision l'état futur du système. Néanmoins, ce type de pronostic est peu applicable en raison de la nécessité de disposer d'une représentation mathématique du processus de dégradation du système, ce qui s'avère particulièrement difficile pour des systèmes complexes. Comme le précise Muller [Muller, 2005], « la dégradation est considérée comme une variable continue, dont l'évolution est déterminée par une loi déterministe ou stochastique ». Dans cette approche de pronostic, les conditions de fonctionnement du système peuvent être directement intégrées au modèle et intervenir dans l'évolution de la dégradation.

II.2. Le pronostic guidé par les données

Dans le cas où aucun modèle physique du mécanisme de dégradation n'existe, et à condition de pouvoir évaluer de manière continue le niveau de dégradation du système, l'état futur du système peut être estimé par l'approche guidée par les données, à l'aide de progiciels tel que CASIP [Léger, 2003]. Cette approche se base notamment sur l'utilisation de méthodes statistiques comme, par exemple, l'analyse de tendance, les méthodes bayésiennes [Iung, 2005] ou encore des réseaux de neurones flous, comme proposés par Cocheteux et *al.* [Cocheteux, 2009].

II.3. Le pronostic basé sur l'expérience

La dernière approche de pronostic proposée est dite basée sur l'expérience. Elle est usuellement employée lorsqu'aucune des deux précédentes méthodes n'est applicable. Ce type de pronostic se base sur la connaissance fiabiliste du système, résultant des données de retour d'expériences ou de l'avis d'experts. Pour tenir compte des conditions de fonctionnement, les fonctions de fiabilité ou de taux de défaillances peuvent être influencées par des variables explicatives en accélérant le processus de dégradation (*Accelerated Life Model*) ou en augmentant le taux de défaillance (*Proportional Hazards Model*).

Comme précisé auparavant, l'objectif de ce travail de recherche vise à déterminer, avant que le navire ne s'engage en mer, les plans d'exploitation et de maintenance optimaux. Parmi les différentes techniques de pronostic, l'approche guidée par les données ne peut être appliquée à notre problème, car elle repose sur l'exploitation d'indicateurs de dégradation au cours du temps. Or justement, le but de ce travail est de déterminer les plans optimaux avant que le navire ne quitte le quai, c'est-à-dire avant d'obtenir les données de suivis de dégradations. Le système étudié étant très complexe, il n'existe pas de modèle physique définissant l'évolution de la dégradation, qui plus est, en faisant intervenir des facteurs extérieurs tels que les conditions environnementales, etc. L'état futur du système sera donc déterminé en se basant sur l'exploitation d'un modèle fiabiliste. Pour tenir compte des caractéristiques des missions et de leurs influences sur la fiabilité du système, des variables explicatives seront introduites dans l'expression fiabiliste pour prendre note de l'impact opérationnel et environnemental des différentes missions. Les deux principaux concepts présents dans la littérature [Bagdonavičius, 1994], à savoir les modèles ALM (*Accelerated Life Model*) et PHM

(*Proportional Hazards Model*), sont étudiés précisément dans la partie suivante en vue de sélectionner celui qui sera utilisé dans le reste de l'étude pour modéliser l'aspect pronostic.

III. Modèle de vie accélérée

Le modèle de vie accélérée [Martorell, 1999] se caractérise par l'existence d'une fonction de fiabilité nominale qui peut être influencée par un vecteur de covariables, en l'occurrence les conditions opérationnelles et environnementales. L'effet de ces variables explicatives consiste à modifier l'évolution de la fiabilité en modifiant l'âge du système. Basée sur ces variables, la fonction de fiabilité devient :

$$R_{ALM}(t, Z) = R_0(g(Z) \cdot t) \quad t \geq 0 \quad (\text{II. 1})$$

Avec :

- $R_{ALM}(t, Z)$: Fonction de fiabilité basée sur le modèle ALM
- $g(Z)$: Fonction de risque dépendant du vecteur Z
- Z : Vecteur qui contient q covariables, représentant les conditions de fonctionnement

La fonction de fiabilité nominale suppose que les q variables explicatives sont nulles. Les variables explicatives sont liées à la fonction de risque et doivent vérifier les relations suivantes :

$$\begin{cases} g(0) = 1 \\ g(Z) > 0 \quad \forall Z \end{cases} \quad (\text{II. 2})$$

Généralement, la fonction de risque est modélisée par une fonction log-linéaire de la forme $g(Z) = e^{B^T \cdot Z}$. Le vecteur transposé B^T correspond aux q coefficients associés à chaque covariable du vecteur Z . Ces coefficients sont utilisés pour nuancer l'influence de chacune des variables explicatives. Dans ce modèle de vie accélérée, les variables explicatives permettent d'accélérer l'âge du système lorsque $g(Z) > 1$. Au contraire, si $g(Z) < 1$, ce modèle permet de réduire l'âge du système. La fonction de défaillance associée à ce modèle peut se déduire à partir de l'équation (II. 1) et s'exprime comme suit :

$$\lambda_{ALM}(t, Z) = g(Z) \cdot \lambda_0(g(Z) \cdot t) \quad t \geq 0 \quad (\text{II. 3})$$

Où $\lambda_{ALM}(t, Z)$ représente la fonction du taux de défaillance basée sur le modèle ALM. Suivant ce modèle, les fonctions de fiabilité et de taux de défaillance évoluent, comme le précise la figure 5.

Le second modèle couramment employé dans la littérature correspond au modèle à risques proportionnels. Ce modèle présente une grande similitude avec le modèle de vie accélérée.

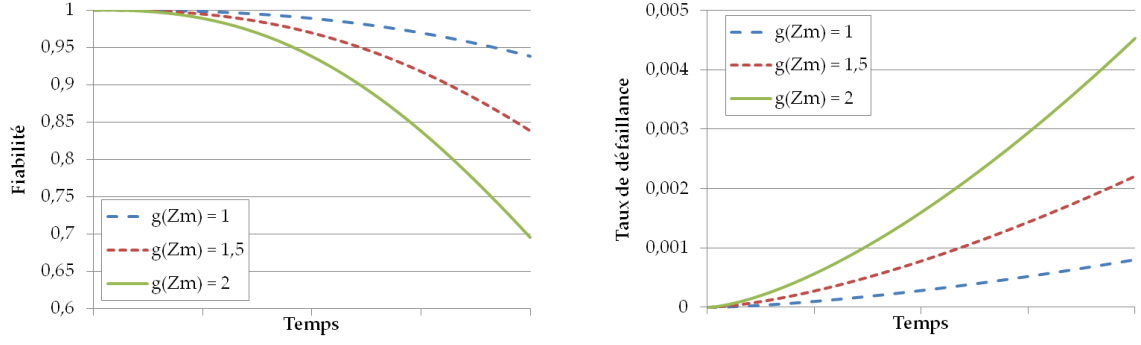


Figure 5 : Évolution de la fiabilité et du taux de défaillance pour le modèle ALM pour une distribution nominale de type Weibull, de paramètres de forme 2,5 et d'échelle 600

IV. Modèle à risques proportionnels

Bien qu'il fût initialement utilisé dans le domaine médical, d'après Lyonnet [Lyonnet, 2006], ce modèle a été utilisé pour mesurer l'effet de traitements thérapeutiques sur des populations humaines. Ce ne fut que bien plus tard qu'on l'utilisât dans le domaine industriel, et plus précisément dans le domaine de la maintenance.

Alors que le modèle de vie accélérée se caractérise par une modification du processus d'usure (variation de l'âge du système), le modèle à risques proportionnels se définit par une variation du taux de défaillance proportionnellement aux conditions de fonctionnement. En considérant ces variables explicatives, la fonction de défaillance s'exprime par :

$$\lambda_{PHM}(t, Z) = g(Z) \cdot \lambda_0(t) \quad t \geq 0 \quad (\text{II. 4})$$

Où $\lambda_{PHM}(t, Z)$ représente la fonction du taux de défaillance basée sur le modèle PHM.

Comme pour le modèle précédent, la fonction de risque est régulièrement représentée par la même fonction log-linéaire représentant des effets similaires. Ainsi, les relations exprimées par (II. 2) doivent être vérifiées. Par ailleurs, lorsque $g(Z) > 1$ le taux de défaillances augmente, et par conséquent, ce taux diminue pour $g(Z) < 1$. À partir de la fonction du taux de défaillance, la fonction de fiabilité peut être déterminée et s'exprime par :

$$R_{PHM}(t, Z) = [R_0(t)]^{g(Z)} \quad t \geq 0 \quad (\text{II. 5})$$

Où $R_{PHM}(t, Z)$ caractérise la fonction de fiabilité associée au modèle PHM

Les modèles de vie accélérée ou à risque proportionnels sont des modèles semi-paramétriques et à risques proportionnels. En effet, l'effet des covariables est indépendant du temps. La différence entre ces deux modèles se situe dans la représentation de l'effet des variables explicatives. Par rapport à notre système d'étude, le modèle à risques proportionnels semble être le plus approprié pour représenter l'effet des facteurs extérieurs sur la distribution des durées de vies. Connaissant le modèle utilisé pour la modélisation des facteurs de fonctionnement, l'étape suivante consiste à formaliser cette fonction de risque.

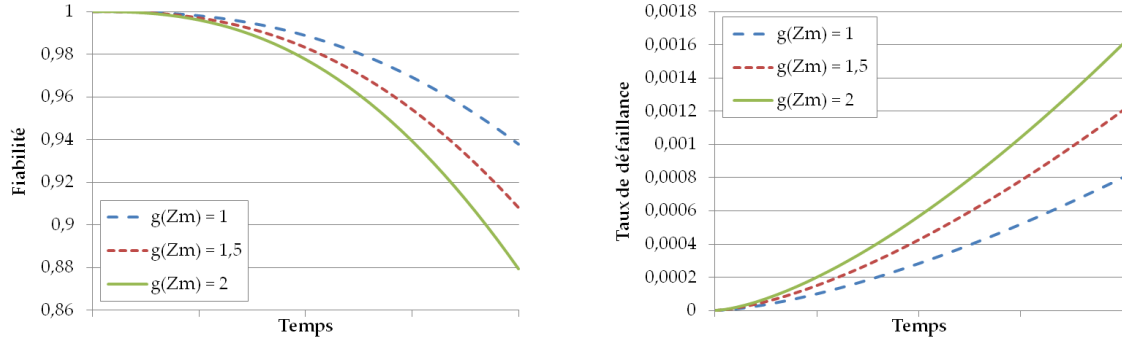


Figure 6 : Évolution de la fiabilité et du taux de défaillance pour le modèle PHM pour une distribution nominale de type Weibull, de paramètres de forme 2,5 et d'échelle 600

V. Formalisation de la fonction de risque

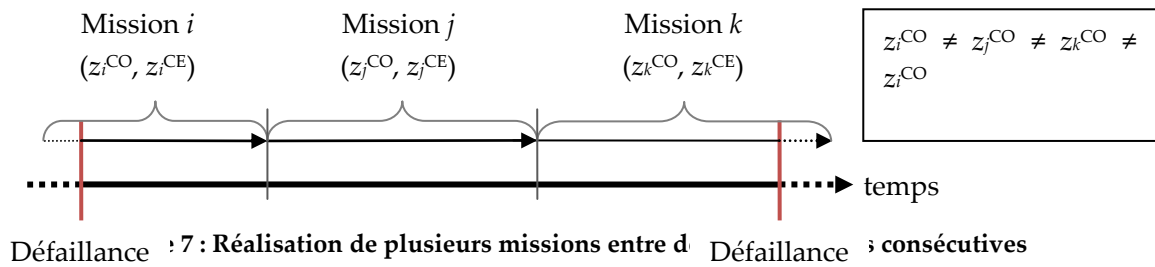
Comme précisé précédemment, les missions sont caractérisées par des conditions de fonctionnement, liées aux risques opérationnels et aux risques environnementaux. Bien évidemment, ces deux types de risques de défaillance n'ont pas le même impact sur la modification du taux de défaillance. Par conséquent, la fonction de risque $g(Z_m)$ peut être formulée comme suit :

$$\begin{aligned} g(Z_m) &= e^{B^T \cdot Z_m} \\ &= e^{\beta^{CO} \cdot z_m^{CO} + \beta^{CE} \cdot z_m^{CE}} \quad \forall m \in \mathcal{K} \end{aligned} \quad (\text{II. 6})$$

Avec :

- β^{CO} : coefficient de pondération associé à la covariable z_m^{CO}
- β^{CE} : coefficient de pondération associé à la covariable z_m^{CE}

Pour déterminer les coefficients associés aux variables explicatives dans le modèle à risques proportionnels, Cox [Cox, 1972] a proposé la méthode du maximum de vraisemblance partielle. Cette méthode est dérivée du maximum de vraisemblance classique et ne nécessite pas la connaissance de la distribution des durées de vies. Cette vraisemblance partielle se calcule comme le produit des contributions pour chaque date de défaillance suivant les variables explicatives (facteurs de risques) associées. D'après Lyonnet [Lyonnet, 2000], « la contribution V_i du composant i défaillant en t_i à la vraisemblance partielle V^* est égale à la probabilité conditionnelle, que ce soit le sujet soumis aux contraintes (facteurs de risques) Z_i qui soit défaillant en t_i , connaissant la population à risques en cet instant $n(t_i)$ » avec $i = \{1, 2, \dots, n\}$. Dans cette définition, il est clairement précisé que le composant i est soumis à des contraintes Z_i , ce qui signifie que ces conditions de fonctionnement sont constantes au cours du temps. Autrement dit, entre deux dates de défaillances consécutives, les valeurs des variables explicatives restent inchangées. Nonobstant, dans notre étude, cette hypothèse ne peut être vérifiée. En effet, plusieurs missions de caractéristiques différentes peuvent être réalisées successivement entre deux maintenances correctives consécutives, comme l'illustre la figure suivante.



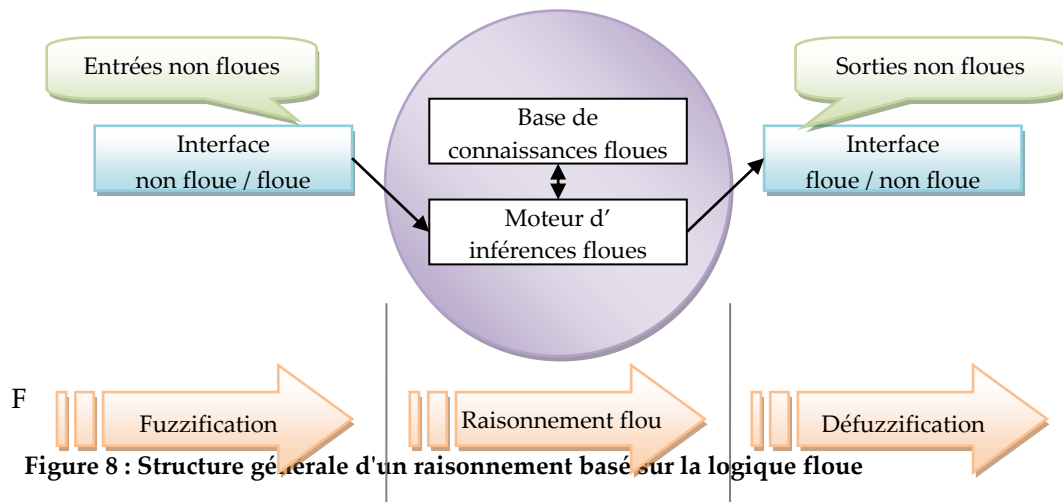
Par conséquent, il convient de définir pour chaque période de bon fonctionnement, des conditions de fonctionnement équivalentes qui représentent les missions effectuées. En d'autres termes, il convient de déterminer une mission qui, réalisée entre les deux défaillances, modéliserait l'effet des missions i , j et k (les missions i et k étant tronquées). Bien évidemment, la recherche de ces covariables équivalentes fait référence à la durée des missions réalisées. À titre d'exemple, si le système est soumis à des conditions de fonctionnement élevées pendant une durée importante et à des conditions quasi nominales pendant une courte durée, il est raisonnable de penser que les conditions équivalentes seront assez proches des conditions élevées. Pour déterminer ces conditions de fonctionnement équivalentes, nous avons choisi d'employer un raisonnement par logique floue, principalement pour tenir compte de la subjectivité, de l'imprécision et des incertitudes.

VI. Structure d'un raisonnement par logique floue

Les premiers travaux de recherches traitant le principe de logique floue sont dus à Zadeh [Zadeh, 1965]. Cette théorie du floue a rencontré un grand essor car elle se base sur un raisonnement intuitif et permet de considérer la subjectivité et l'imprécision. En effet, si une donnée n'est pas connue précisément, elle peut néanmoins être exprimée par un intervalle de confiance précis, c'est-à-dire par un ensemble de valeurs possibles.

Un raisonnement par logique floue est constitué de trois étapes principales :

- La fuzzification consiste en la transformation de données non floues en donnée floues.
- Le raisonnement flou est constitué de deux parties, une statique, l'autre dynamique. La base de connaissance floue (partie statique) permet de traduire l'avis des experts, alors que le moteur d'inférence floue (partie dynamique) permet d'appliquer les règles issues de la base de connaissances.
- La défuzzification est l'opération opposée à la fuzzification. Elle transforme les valeurs floues obtenues en valeur non floues.



À présent, nous allons détailler le principe de fonctionnement d'un raisonnement basé sur la logique floue.

VI.1. Fuzzification

Comme l'illustre la figure 8, la fuzzification est la première étape d'un contrôleur flou. Comme un système flou opère sur des valeurs qualitatives, il convient de transformer des valeurs quantitatives en valeurs floues. En respectant la syntaxe de la logique floue [Godjevac, 1999], il convient donc de définir les variables linguistiques qui correspondent à un état. On définit par état les variables d'entrées ou de sorties du contrôleur flou. Une variable linguistique est caractérisée par un ensemble de données qui sont son nom, l'ensemble de ses valeurs linguistiques ainsi que l'univers du discours. Pour illustrer ces différentes spécificités, nous allons considérer comme variable linguistique la taille d'un homme. Le terme « taille » peut être utilisé pour définir le nom de cette variable. Les valeurs prises par cette variable constituent les sous-ensembles flous. Qualitativement, on peut définir la taille d'un homme comme « très petite », « petite », « moyenne », grande » et « très grande ». L'ensemble des valeurs linguistiques prises par une variable constitue l'univers du discours. Un ensemble flou se définit comme l'ensemble des éléments qui appartiennent à une valeur linguistique, mais avec des degrés d'appartenance différents. Par exemple, en supposant qu'une taille de 1,70 m est considérée comme « moyenne » à 100%, nous pouvons présumer qu'une taille de 1,69 m ou 1,71 m est toujours une taille moyenne, mais à 90 %. Ainsi, la logique floue est justement employée pour éviter les limites posées par la logique booléenne classique. Aussi, à chaque valeur linguistique est associée une fonction d'appartenance. Elle modélise l'évolution de la valeur linguistique. Dans la littérature [Ross, 2004], de nombreuses formes de fonctions d'appartenance ont été proposées. Suivant les applications et la base de connaissances fournies par les experts, elles peuvent être de forme monotone, triangulaire, trapézoïdale ou gaussienne. Généralement, lorsqu'une variable appartient à 100% à une valeur linguistique, son niveau d'appartenance à ce sous-ensemble flou est égal à 1. Au contraire, lorsqu'un élément n'appartient pas à un sous-ensemble flou, le niveau d'appartenance à cette valeur linguistique vaut 0. Pour les autres situations, le niveau d'appartenance se situe dans l'intervalle]0, 1[.

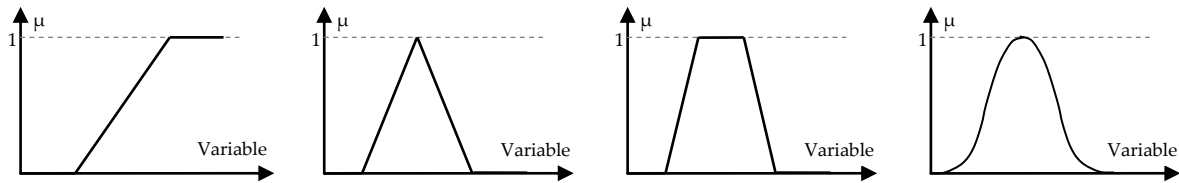


Figure 9 : Différentes formes de fonctions d'appartenance

L'étape subséquente à la fuzzification consiste à définir la base de connaissances floues ainsi que le moteur d'inférences.

VI.2. Raisonnement flou

L'étape intitulée « raisonnement flou » est formée de deux étapes. La première étape, à savoir la base de connaissances floues, requiert les informations issues d'experts. Il s'agit d'établir les règles linguistiques, notions introduites par Zadeh [Zadeh, 1973], qui régissent le moteur d'inférences floues. Ces règles sont de la forme « Si... alors... ». La partie antécédente de cette règle, encore appelée prémisse, est définie par « Si... », alors que la partie conséquente ou conclusion est exprimée par « alors... ». Les prémisses de ces règles linguistiques peuvent être constituées d'une ou plusieurs variables d'entrées. Dans ce dernier cas, les variables d'entrées sont combinées en utilisant les opérateurs *et* et/ou *ou*.

Lors de la seconde étape, le moteur d'inférences floues vérifie pour chaque règle, définie dans la base de connaissances, si cette dernière peut être appliquée au présent problème. Une règle est utilisée si la prémisse est vérifiée, c'est-à-dire si le poids d'activation est non nul. Le poids d'activation d'une règle est obtenu en fonction du niveau d'appartenance des variables d'entrées associées à la dite règle. Le calcul du poids d'activation des règles dépend du type d'inférence floue choisi. Les plus connues et utilisées dans la littérature sont les méthodes du « Max – Prod » et du « Max – Min ». Par ailleurs, le chevauchement des fonctions d'appartenances peut conduire à l'activation simultanée de deux ou plusieurs règles. La méthode d'inférence dite du « Max – Prod », présentée à la figure 10, réalise au niveau des conditions, la formation du maximum par l'opérateur *OU* alors que la formation du minimum résulte de l'opérateur *ET*. La conclusion introduite par *ALORS*, au niveau de chaque règle, se caractérise par le produit de la fonction d'appartenance par le poids d'activation, poids issu de la prémisse. L'opérateur *OU* qui lie les différentes règles est caractérisé par la formation du maximum.

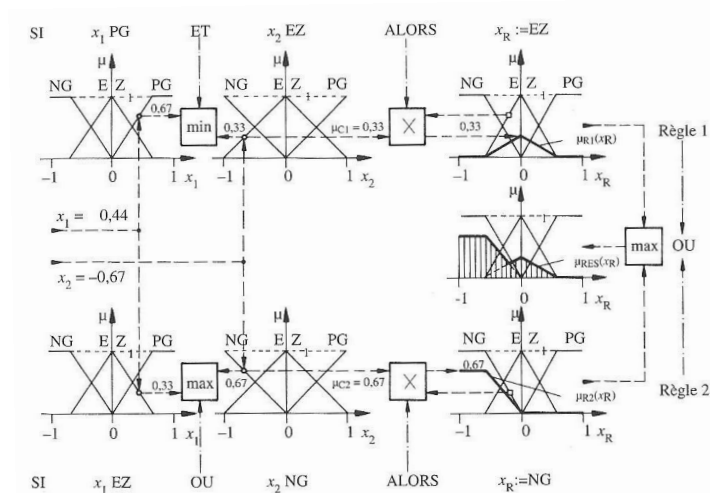


Figure 10 : Méthode d'inférence « Max - Prod » pour deux variables d'entrée et deux règles [Bühler, 1994]

Dans la méthode d'inférence du « Max – Min », les opérateurs *ET* et *OU* sont également réalisés, respectivement, par la formation du minimum et du maximum. La différence se situe au niveau de la conclusion, introduite par *ALORS* ; où la valeur de sortie est obtenue par la formation du minimum qui sert de valeur d'écrtage pour le sous-ensemble flou concerné par la variable de sortie. L'opérateur *OU*, qui relie les différentes règles, est quant à lui, réalisé par la formation du maximum.

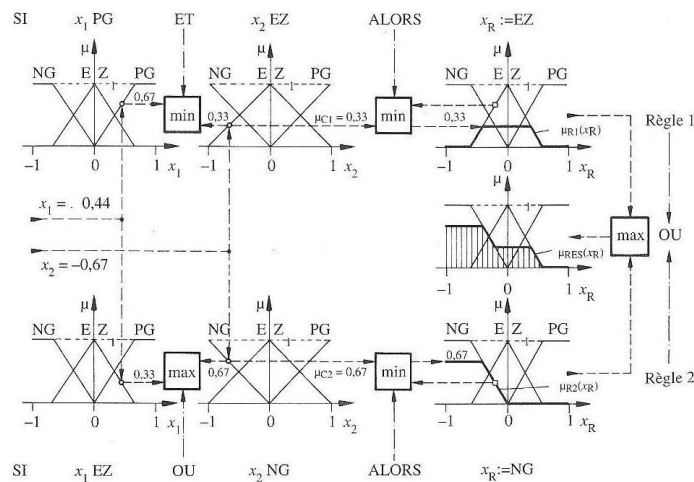


Figure 11 : Méthode d'inférence « Max - Min » pour deux variables d'entrée et deux règles [Bühler, 1994]

Lorsque la fonction d'appartenance résultante est obtenue, l'étape suivante concerne l'obtention de la valeur non-floue de la variable de sortie.

VI.3. Défuzzification

À l'opposé de la fuzzification, la défuzzification a pour rôle de transformer les valeurs floues en valeurs non floues. Comme nous avons pu le remarquer, le raisonnement flou renvoie à une fonction d'appartenance résultante pour la variable de sortie. Les deux principales méthodes que l'on rencontre dans la littérature sont la défuzzification par valeur maximum et celle par centre de

gravité. Pour la première méthode, l'obtention de la valeur de sortie est obtenue en considérant l'abscisse de la valeur maximale de la fonction d'appartenance résultante. Dans le cas où cette dernière est écrêtée (méthode du « Max – Min »), nous considérons la moyenne des abscisses de la valeur maximale de la fonction d'appartenance. Bien que cette méthode offre l'avantage d'une réponse très rapide, elle n'est pas recommandée pour un raisonnement par logique floue. En effet, elle néglige l'influence des autres valeurs sur la fonction d'appartenance résultante. Pour pallier ce problème, la méthode du centre de gravité permet d'intégrer l'influence de toutes les valeurs sur la fonction d'appartenance résultante. Le centre de gravité, c'est-à-dire la valeur de sortie, est obtenu en déterminant le rapport entre la surface (fonction d'appartenance résultante) et le moment de celle-ci. Bien que cette méthode offre une très bonne précision concernant la variable de sortie, elle exige, néanmoins, une envergure de calcul assez importante.

VII. Application à notre problématique

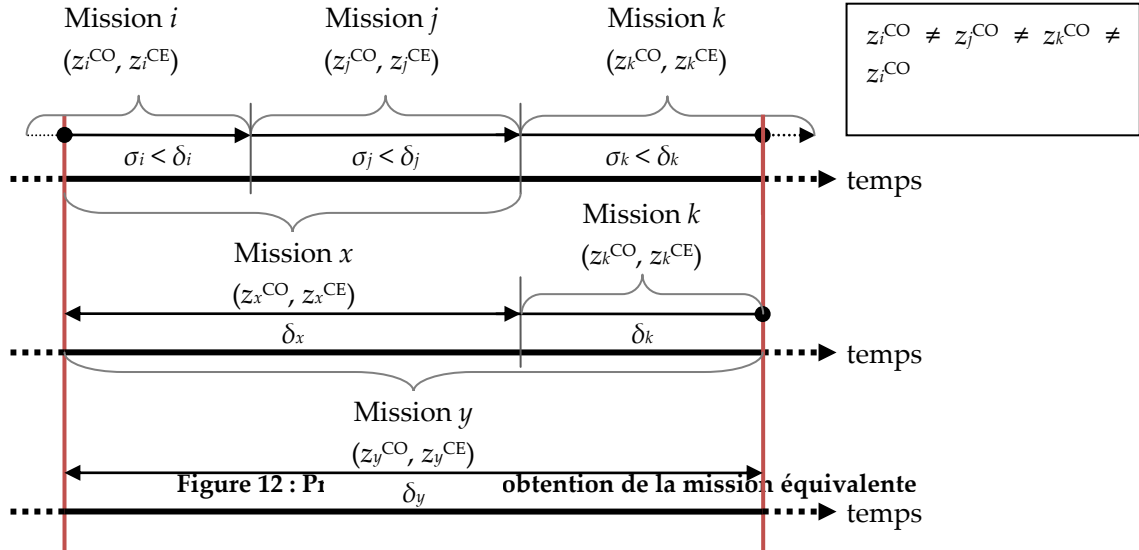
Dans ce travail, l'objectif consiste à modéliser des conditions de fonctionnement équivalentes aux niveaux opérationnels et environnementaux, représentatives de l'ensemble des missions réalisées. L'obtention de ces conditions de fonctionnement dépend, en toute logique, des conditions de fonctionnement des missions concernées, mais également de la durée durant laquelle le système a été soumis à ces conditions. Par exemple, les conditions de fonctionnement pour un système soumis à 70 % du temps à des conditions extrêmes et le reste à des conditions nominales seront, de toute évidence, plus importantes que la situation opposée.

Cependant, l'application d'un raisonnement par logique floue à cette problématique pose quelques difficultés. Comme nous venons de le voir précédemment, un raisonnement flou nécessite un nombre de variables d'entrées déterminé. Dans ce problème, le nombre de variables de sorties est fixe, à savoir la condition opérationnelle équivalente et la condition environnementale équivalente. Par contre, le nombre des variables d'entrées est variable. Il s'agit des conditions opérationnelles et environnementales ainsi que de la durée de chacune des missions appartenant à l'intervalle de bon fonctionnement. Par conséquent, le nombre de variables d'entrées est multiple du nombre de missions appartenant à l'intervalle de bon fonctionnement.

Pour minimiser la taille du problème de logique floue, et en se basant sur l'hypothèse de conditions de fonctionnement indépendantes (les conditions opérationnelles n'interviennent pas sur les conditions environnementales et inversement), la détermination de ces dites conditions peut être traitée séparément. Un raisonnement par logique floue sera opéré pour la détermination des conditions opérationnelles et, de façon similaire, un autre raisonnement par logique floue permettra de déterminer les conditions environnementales. Pour la détermination de ces conditions, les fonctions d'appartenances, les règles linguistiques et toutes les méthodes employées resteront identiques.

Pour pallier le problème du nombre dynamique des variables d'entrées, l'idée consiste à déterminer la mission équivalente, en se basant sur des missions intermédiaires constituées uniquement de deux missions. En reprenant l'exemple illustré par la figure 7, la mission équivalente sera déterminée en suivant la procédure suivante. À partir des missions i et j , une nouvelle mission intermédiaire x peut être déterminée. Avec cette nouvelle mission x et en considérant la mission

suivante, à savoir la mission k , une nouvelle mission y peut être déterminée. Cette procédure est répétée tant que toutes les missions appartenant à l'intervalle de bon fonctionnement n'auront pas participé à l'obtention de la mission équivalente finale.



Dans cet exemple, la mission y constitue donc la mission équivalente. Grâce à cette méthode, le nombre de variables d'entrées pour chaque raisonnement flou est limité à quatre, réparti comme suit :

- Condition opérationnelle (ou environnementale) de la première mission,
- Durée (durant l'intervalle de bon fonctionnement) de la première mission,
- Condition opérationnelle (ou environnementale) de la seconde mission,
- Durée (durant l'intervalle de bon fonctionnement) de la seconde mission.

Concernant la durée des missions, il ne s'agit pas de la durée totale de la mission, mais uniquement du temps passé par la mission dans l'intervalle de temps de bon fonctionnement. Par conséquent, celle-ci est inférieure ou égale à la durée totale de la mission considérée. Ainsi, on note par σ_i le temps passé par la mission i dans l'intervalle de bon fonctionnement.

Connaissant les variables d'entrées et de sorties, il est possible de définir les variables linguistiques associées. Dans la suite, nous traiterons uniquement la détermination des conditions opérationnelles équivalentes puisque les conditions environnementales seront obtenues suivant le même procédé. Il s'agira simplement de modifier l'intitulé « Conditions opérationnelles » par « Conditions environnementales ».

VII.1. Structure des variables linguistiques

Parmi les quatre variables d'entrées, nous définirons uniquement les variables linguistiques associées à la première mission, notée m . Les variables associées à la seconde mission ($m + 1$) seront strictement identiques. Au cours du premier chapitre, les conditions opérationnelles ont été définies comme une note (un nombre réel) variant entre zéro, pour les conditions nominales, et dix, pour modéliser les conditions extrêmes. Les caractéristiques associées à cette variable d'entrée sont donc exprimées par :

- z_m^{CO} : nom de la variable associée à la mission m ,

- $[0, 10]$: l'univers du discours.
- L'univers du discours est constitué de trois sous-ensembles flous dont les valeurs linguistiques sont « Nominale », « Intermédiaire » et « Extrême ».

La variable d'entrée, associée à la durée de la mission, se caractérise par :

- σ_m : nom de la variable associée à la mission m ,
- $]0, +\infty]$: l'univers du discours.
- L'univers du discours est fractionné de trois sous-ensembles flous dont les valeurs linguistiques sont « Courte », « Modérée » et « Longue ».

Dans la phase de fuzzification, il est nécessaire de définir pour chaque variable linguistique les fonctions d'appartenance. Celles-ci régissent les valeurs linguistiques en pondérant entre 0 et 1 les valeurs des variables. Pour représenter les variables linguistiques associées aux conditions de fonctionnement, des fonctions d'appartenance de formes triangulaires ont été employées. Cette forme triangulaire, composée uniquement de droites est certes simple, mais s'avère suffisante pour délimiter les ensembles flous.

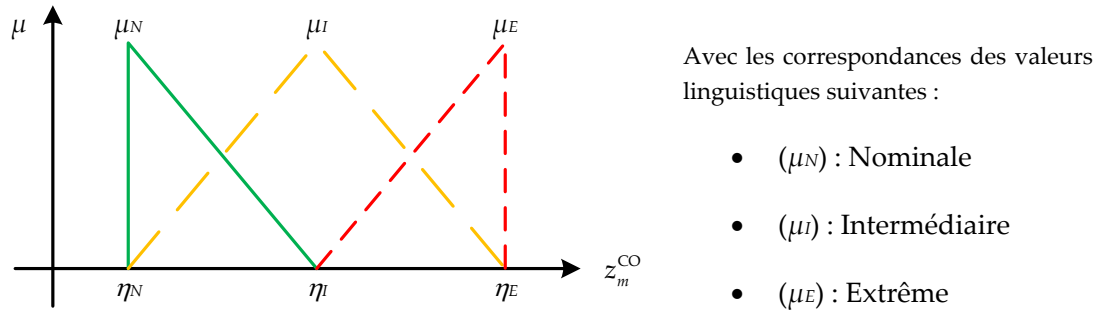


Figure 13 : Fonctions d'appartenance, avec trois ensembles, pour la variable linguistique « condition opérationnelle », associées à la mission m

De manière générale, ces fonctions d'appartenance, associées aux conditions de fonctionnement, sont régies par les expressions mathématiques suivantes :

$$\mu(z_i^{C?}) = \begin{cases} \mu_N(z_i^{C?}) = -\frac{1}{\eta_I - \eta_N} \cdot z_i^{C?} + \frac{\eta_I}{\eta_I - \eta_N} & \text{pour } \eta_N \leq z_i^{C?} \leq \eta_I \\ \mu_I(z_i^{C?}) = \begin{cases} \frac{1}{\eta_I - \eta_N} \cdot z_i^{C?} - \frac{\eta_N}{\eta_I - \eta_N} & \text{pour } \eta_N \leq z_i^{C?} \leq \eta_I \\ -\frac{1}{\eta_E - \eta_I} \cdot z_i^{C?} + \frac{\eta_E}{\eta_E - \eta_I} & \text{pour } \eta_I \leq z_i^{C?} \leq \eta_E \end{cases} \\ \mu_E(z_i^{C?}) = \frac{1}{\eta_E - \eta_I} \cdot z_i^{C?} - \frac{\eta_I}{\eta_E - \eta_I} & \text{pour } \eta_I \leq z_i^{C?} \leq \eta_E \end{cases} \quad (\text{II. 7})$$

Dans cette équation (II. 7), la variable i correspond à la mission m ou $(m + 1)$ et $z_i^{C?}$ correspond à z_i^{CO} et z_i^{CE} , respectivement pour les conditions opérationnelles et environnementales. Par ailleurs, nous avons choisi d'associer aux paramètres η_N , η_I et η_E les valeurs (ou notes) de 0, 5 et 10.

Concernant les variables d'entrées associées à la durée de la mission, les fonctions d'appartenances de formes trapézoïdales et triangulaires ont été utilisées. Les formes trapézoïdales s'expliquent par l'univers de discours qui couvre une durée pseudo-infinie. En effet, une mission a une durée strictement supérieure à 0 ut (unités de temps), mais la durée maximale d'une mission n'est pas fixée.

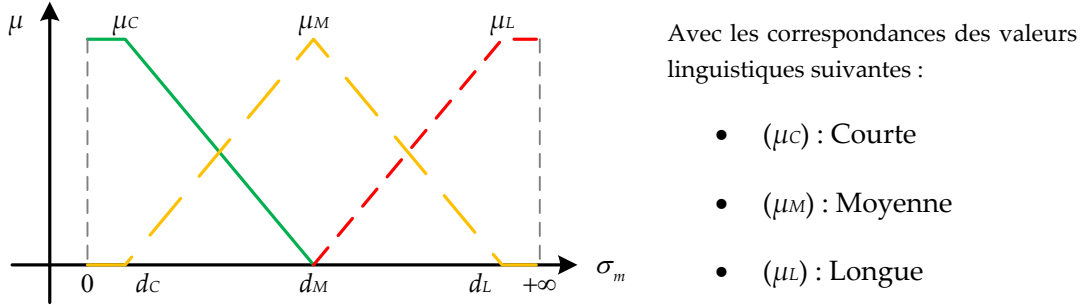


Figure 14 : Fonctions d'appartenance, avec trois ensembles, pour la variable linguistique « durée », associées à la mission m

Les fonctions d'appartenances modélisant chacun des sous-ensembles flous sont définies par les expressions mathématiques suivantes :

$$\mu(\sigma_i) = \begin{cases} \mu_C(\sigma_i) = \begin{cases} 1 & \text{pour } 0 \leq \sigma_i \leq d_C \\ -\frac{1}{d_M - d_C} \cdot \sigma_i + \frac{d_M}{d_M - d_C} & \text{pour } d_C \leq \sigma_i \leq d_M \end{cases} \\ \mu_M(\sigma_i) = \begin{cases} \frac{1}{d_M - d_C} \cdot \sigma_i - \frac{d_C}{d_M - d_C} & \text{pour } d_C \leq \sigma_i \leq d_M \\ -\frac{1}{d_L - d_M} \cdot \sigma_i + \frac{d_L}{d_L - d_M} & \text{pour } d_M \leq \sigma_i \leq d_L \end{cases} \\ \mu_L(\sigma_i) = \begin{cases} \frac{1}{d_L - d_M} \cdot \sigma_i - \frac{d_M}{d_L - d_M} & \text{pour } d_M \leq \sigma_i \leq d_L \\ 1 & \text{pour } d_L \leq \sigma_i \leq \infty \end{cases} \end{cases} \quad (\text{II. 8})$$

dans l'équation (II. 8) et de façon analogue à l'équation définie par (II. 7), la variable i correspond à la mission m ou $(m + 1)$. De même, σ_i représente la durée passée par la mission i dans l'intervalle de bon fonctionnement, délimité par les deux défaillances consécutives. Pour cette étude, nous pouvons associer aux paramètres d_C , d_M et d_L les valeurs de 50, 100 et 150 ut.

Pour rappel, l'objectif de ce raisonnement flou était l'obtention des conditions de fonctionnement équivalentes aux différentes missions réalisées. Par conséquent, la variable de sortie correspond aux conditions opérationnelles et environnementales, respectivement pour chacun des deux raisonnements flous. La modélisation de cette variable est identique aux variables d'entrées $z_i^{C?}$ avec $i = \{m, m + 1\}$ et $C? = \{CO, CE\}$, comme défini par la figure 13.

VII.2. Définition du moteur d'inférence flou

Concernant le moteur d'inférence flou, la principale étape consiste à définir l'ensemble des règles linguistiques qui contrôleront ce système flou. Ce dernier étant constitué de quatre variables linguistiques définissant les entrées floues, et compte tenu de la subdivision de ces univers en trois sous-ensembles, 81 règles peuvent être établies. Nous avons souhaité définir chaque règle afin d'obtenir une précision correcte quant à la valeur de sortie. L'ensemble des règles établies peut être représenté sous la forme matricielle à l'aide d'une matrice d'inférence ("*Fuzzy Associative Matrix*") de dimension quatre.

Conditions de fonctionnement → Durée ↓			Mission m								
			Nominale			Intermédiaire			Extrême		
			Courte	Modérée	Longue	Courte	Modérée	Longue	Courte	Modérée	Longue
$m+1$	Nominale	Courte	μ_N	μ_N	μ_N	μ_N	μ_I	μ_I	μ_I	μ_I	μ_E
		Modérée	μ_N	μ_N	μ_N	μ_N	μ_N	μ_I	μ_I	μ_I	μ_I
		Longue	μ_N	μ_N	μ_N	μ_N	μ_N	μ_N	μ_N	μ_I	μ_I
	Intermédiaire	Courte	μ_N	μ_N	μ_N	μ_I	μ_I	μ_I	μ_E	μ_E	μ_E
		Modérée	μ_I	μ_N	μ_N	μ_I	μ_I	μ_I	μ_I	μ_E	μ_E
		Longue	μ_I	μ_I	μ_N	μ_I	μ_I	μ_I	μ_I	μ_I	μ_E
	Extrême	Courte	μ_I	μ_I	μ_I	μ_E	μ_I	μ_I	μ_E	μ_E	μ_E
		Modérée	μ_I	μ_I	μ_I	μ_E	μ_E	μ_I	μ_E	μ_E	μ_E
		Longue	μ_E	μ_I	μ_I	μ_E	μ_E	μ_E	μ_E	μ_E	μ_E

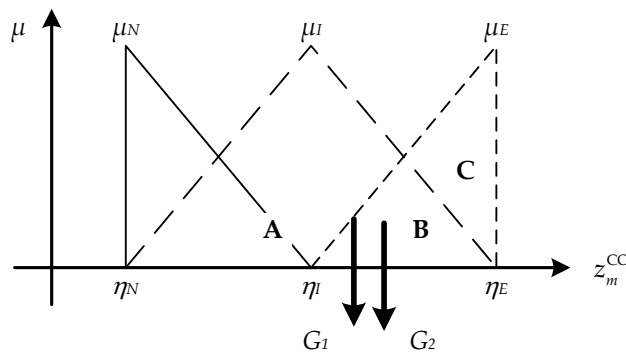
Tableau 2 : Matrice d'inférence

À travers cette matrice d'inférence, nous pouvons remarquer deux caractéristiques principales. La première spécificité se situe lorsque deux missions successives ont des conditions de fonctionnement identiques. Dans ce cas, il est parfaitement logique que la mission résultante conserve les mêmes conditions, peu importe les durées des missions. La deuxième caractéristique fait référence au modèle utilisé pour représenter l'aspect dynamique de la loi de défaillance. Le modèle à risques proportionnels étant indépendant du temps, l'ordre des missions n'a donc aucune influence. Par conséquent, la matrice d'inférence présente un aspect symétrique.

Ce raisonnement flou faisant intervenir quatre variables d'entrées, et compte tenu du chevauchement des fonctions d'appartenance de ces variables, l'obtention des conditions de fonctionnement pour chaque mission intermédiaire peut résulter, au maximum, de l'activation simultanée de seize règles. Pour traiter ces inférences, nous avons opté pour la méthode d'inférence max-min, dont la valeur de sortie est généralement obtenue par la méthode du centre de gravité.

VII.3. Détermination de la valeur de sortie

La détermination de la valeur de sortie correspond à la dernière étape du raisonnement par logique floue. Par rapport aux différentes méthodes présentées dans la partie « Chapitre IIVI.3 Défuzzification », les conditions de fonctionnement seront obtenues grâce à la méthode des centroïdes car elle permet d'obtenir une réponse précise en intégrant l'impact des différentes valeurs de sortie sur les fonctions d'appartenance. Néanmoins, dans la littérature, la réalisation de l'opérateur *OU* liant les règles s'effectue toujours par la formation du maximum. En considérant l'exemple proposé sur la figure 15, le centre de gravité obtenu par cette méthode est noté G_1 .



Le centre de gravité G_1 est obtenu en considérant une seule fois les différentes parties hachurées (**A**, **B** et **C**).

Le centre de gravité G_2 est obtenu en considérant deux fois la partie hachurée **B** et une fois les surfaces **A** et **C**.

Figure 15 : Défuzzification par la méthode du centre de gravité (méthode des centroïdes)

Dans notre problématique, l'activation de seize règles peut amener à obtenir plusieurs fois les mêmes valeurs de fonctions d'appartenance résultantes. Pour prendre en considération l'ensemble de toutes les règles actives et tenir compte de l'influence de chacune d'elles, nous avons décidé d'intégrer dans le calcul du centre de gravité l'ensemble des surfaces obtenues. Dans ce cas, nous nous apercevons que la valeur résultante est légèrement supérieure ($G_2 > G_1$ sur l'axe des abscisses) puisque la surface **B** est considérée deux fois.

Afin d'illustrer la démarche présentée jusqu'à présent, nous allons déterminer uniquement les conditions opérationnelles équivalentes pour l'exemple présenté par la figure 12.

VII.4. Exemple illustratif

Nous considérons les missions notées i , j et k définies par :

	Durée réelle ($\delta_m m=\{i, j, k\}$)	Durée tronquée ($\sigma_m m=\{i, j, k\}$)	Cond. opérationnelle	Cond. environnementale
Mission i	100	90	5	6
Mission j	130	130	3	6
Mission k	150	120	8	4

Tableau 3 : Caractéristiques des missions i, j et k

Afin d'obtenir les conditions de fonctionnement équivalentes pour la mission y , il est nécessaire de déterminer la mission intermédiaire x , ainsi que les conditions de fonctionnement associées.

À partir des équations (II. 7) et (II. 8), nous pouvons déterminer les facteurs d'appartenances pour chacun des sous-ensembles flous des variables « durée » et « conditions opérationnelles » pour les missions i et j .

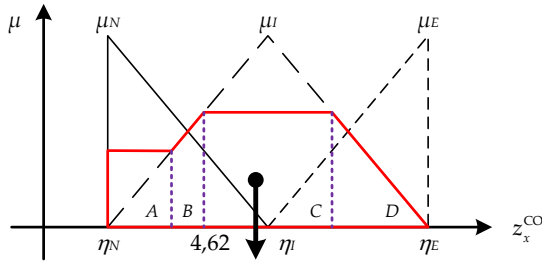
Durée de la mission i (σ_i)	Durée de la mission j (σ_j)	Cond. opérationnelle de la mission i (z_i^{CO})	Cond. opérationnelle de la mission j (z_j^{CO})
$\mu_C(\sigma_i)=0,2$	$\mu_C(\sigma_j)=0$	$\mu_N(z_i^{\text{CO}})=0$	$\mu_N(z_j^{\text{CO}})=0,4$
$\mu_M(\sigma_i)=0,8$	$\mu_M(\sigma_j)=0,4$	$\mu_I(z_i^{\text{CO}})=1$	$\mu_I(z_j^{\text{CO}})=0,6$
$\mu_L(\sigma_i)=0$	$\mu_L(\sigma_j)=0,6$	$\mu_E(z_i^{\text{CO}})=0$	$\mu_E(z_j^{\text{CO}})=0$

Tableau 4 : Détermination des facteurs d'appartenance

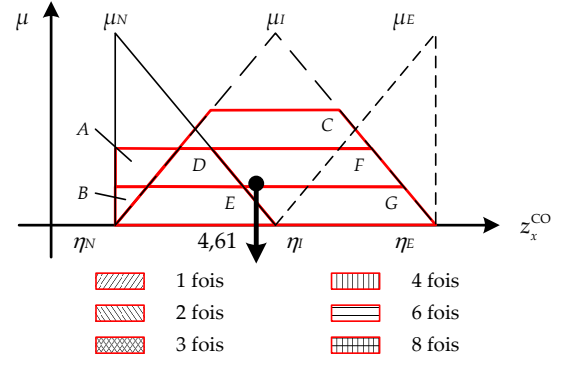
Ces niveaux d'appartenances induisent l'activation des huit règles suivantes dont le niveau d'activation est précisé entre parenthèses. Nous allons donc déterminer la condition opérationnelle pour la mission x .

1. Si $\underbrace{\sigma_i = \mu_C}_{0,2}$ et $\underbrace{z_i^{\text{CO}} = \mu_I}_1$ et $\underbrace{\sigma_j = \mu_M}_{0,4}$ et $\underbrace{z_j^{\text{CO}} = \mu_N}_{0,4}$ alors $z_x^{\text{CO}} = \mu_N$ (0,2)
2. Si $\underbrace{\sigma_i = \mu_C}_{0,2}$ et $\underbrace{z_i^{\text{CO}} = \mu_I}_1$ et $\underbrace{\sigma_j = \mu_M}_{0,4}$ et $\underbrace{z_j^{\text{CO}} = \mu_I}_{0,6}$ alors $z_x^{\text{CO}} = \mu_I$ (0,2)
3. Si $\underbrace{\sigma_i = \mu_C}_{0,2}$ et $\underbrace{z_i^{\text{CO}} = \mu_I}_1$ et $\underbrace{\sigma_j = \mu_L}_{0,6}$ et $\underbrace{z_j^{\text{CO}} = \mu_N}_{0,4}$ alors $z_x^{\text{CO}} = \mu_I$ (0,2)
4. Si $\underbrace{\sigma_i = \mu_C}_{0,2}$ et $\underbrace{z_i^{\text{CO}} = \mu_I}_1$ et $\underbrace{\sigma_j = \mu_L}_{0,6}$ et $\underbrace{z_j^{\text{CO}} = \mu_I}_{0,6}$ alors $z_x^{\text{CO}} = \mu_I$ (0,2)
5. Si $\underbrace{\sigma_i = \mu_M}_{0,8}$ et $\underbrace{z_i^{\text{CO}} = \mu_I}_1$ et $\underbrace{\sigma_j = \mu_M}_{0,4}$ et $\underbrace{z_j^{\text{CO}} = \mu_N}_{0,4}$ alors $z_x^{\text{CO}} = \mu_N$ (0,4)
6. Si $\underbrace{\sigma_i = \mu_M}_{0,8}$ et $\underbrace{z_i^{\text{CO}} = \mu_I}_1$ et $\underbrace{\sigma_j = \mu_M}_{0,4}$ et $\underbrace{z_j^{\text{CO}} = \mu_I}_{0,6}$ alors $z_x^{\text{CO}} = \mu_I$ (0,4)
7. Si $\underbrace{\sigma_i = \mu_M}_{0,8}$ et $\underbrace{z_i^{\text{CO}} = \mu_I}_1$ et $\underbrace{\sigma_j = \mu_L}_{0,6}$ et $\underbrace{z_j^{\text{CO}} = \mu_N}_{0,4}$ alors $z_x^{\text{CO}} = \mu_I$ (0,4)
8. Si $\underbrace{\sigma_i = \mu_M}_{0,8}$ et $\underbrace{z_i^{\text{CO}} = \mu_I}_1$ et $\underbrace{\sigma_j = \mu_L}_{0,6}$ et $\underbrace{z_j^{\text{CO}} = \mu_I}_{0,6}$ alors $z_x^{\text{CO}} = \mu_I$ (0,6)

La détermination de la valeur résultant z_x^{CO} est obtenue par la méthode du centre de gravité.



Centre de gravité par la formation du maximum



Centre de gravité par le modèle proposé

Figure 16 : Détermination de la condition opérationnelle associée à la mission x

Le calcul du centre de gravité par la formation du maximum est donné par :

$$z_x^{CO} = \frac{\begin{pmatrix} A & B & C & D \\ 0,8 \cdot 1 + 0,5 \cdot \mathbf{2,532} + 2,4 \cdot 5 + 0,9 \cdot 8 \end{pmatrix}}{(0,8 + 0,5 + 2,4 + 0,9)} \approx 4,62 \quad (\text{II. 9})$$

Concernant l'approche que nous proposons, le centre de gravité au niveau de l'axe des abscisses est obtenu par :

$$z_x^{CO} = \frac{\begin{pmatrix} A & B & C & D & E & F & G \\ 1 \cdot 0,3 \cdot \mathbf{0,78} + 2 \cdot 0,1 \cdot \mathbf{0,33} + 1 \cdot 1 \cdot 5 + 4 \cdot 0,4 \cdot \mathbf{2,5} + 8 \cdot 0,8 \cdot \mathbf{2,5} + 3 \cdot 1 \cdot 6 + 6 \cdot 1 \cdot 7 \end{pmatrix}}{(0,3 + 2 \cdot 0,1 + 1 + 4 \cdot 0,4 + 8 \cdot 0,8 + 3 \cdot 1 + 6 \cdot 1)} \approx 4,61 \quad (\text{II. 10})$$

Dans les équations (II. 9) et (II. 10), les moments sont représentés en caractères gras et italique, et sont associés aux surfaces qui les précèdent. Nous pouvons constater que pour chacune des deux méthodes proposées, les valeurs de sortie sont approximativement identiques. Par la méthode du centre de gravité appliquée à la surface obtenue par la formation du maximum, nous obtenons des conditions opérationnelles de 4,62 pour la mission x . En considérant l'ensemble des surfaces, la valeur de sortie est de 4,61. En réitérant cette procédure entre la nouvelle mission x et la mission k , nous allons chercher à obtenir les conditions opérationnelles pour la mission y , mission qui équivaut aux missions i, j et k .

La détermination des facteurs d'appartenance pour chaque fonction d'appartenance des variables d'entrées est présentée dans le tableau 5.

Durée de la mission x (σ_x)	Durée de la mission k (σ_k)	Cond. opérationnelle de la mission x (z_x^{CO})	Cond. opérationnelle de la mission k (z_k^{CO})
$\mu_C(\sigma_x) = 0$	$\mu_C(\sigma_k) = 0$	$\mu_N(z_x^{CO}) = 0,08$	$\mu_N(z_k^{CO}) = 0$
$\mu_M(\sigma_x) = 0$	$\mu_M(\sigma_k) = 0,6$	$\mu_I(z_x^{CO}) = 0,92$	$\mu_I(z_j^{CO}) = 0,4$
$\mu_L(\sigma_x) = 1$	$\mu_L(\sigma_k) = 0,4$	$\mu_E(z_x^{CO}) = 0$	$\mu_E(z_k^{CO}) = 0,6$

Tableau 5 : Détermination des facteurs d'appartenance

La mission x étant équivalente aux missions i et j , la durée de celle-ci correspond à la somme des durées des deux missions. Par conséquent, la mission x a une durée de 220 ut (90 + 130) et des conditions opérationnelles de 4,61.

La détermination de ces facteurs d'appartenance induit l'activation simultanée des huit règles linguistiques présentées ci-dessous :

$$1. \text{ Si } \underbrace{\sigma_x = \mu_L}_1 \text{ et } \underbrace{z_x^{\text{CO}} = \mu_N}_{0,08} \text{ et } \underbrace{\sigma_k = \mu_M}_{0,6} \text{ et } \underbrace{z_k^{\text{CO}} = \mu_I}_{0,4} \text{ alors } z_x^{\text{CO}} = \mu_N \quad (0,08)$$

$$2. \text{ Si } \underbrace{\sigma_x = \mu_L}_1 \text{ et } \underbrace{z_x^{\text{CO}} = \mu_N}_{0,08} \text{ et } \underbrace{\sigma_k = \mu_M}_{0,6} \text{ et } \underbrace{z_k^{\text{CO}} = \mu_E}_{0,6} \text{ alors } z_x^{\text{CO}} = \mu_I \quad (0,08)$$

$$3. \text{ Si } \underbrace{\sigma_x = \mu_L}_1 \text{ et } \underbrace{z_x^{\text{CO}} = \mu_N}_{0,08} \text{ et } \underbrace{\sigma_k = \mu_L}_{0,4} \text{ et } \underbrace{z_k^{\text{CO}} = \mu_I}_{0,4} \text{ alors } z_x^{\text{CO}} = \mu_N \quad (0,08)$$

$$4. \text{ Si } \underbrace{\sigma_x = \mu_L}_1 \text{ et } \underbrace{z_x^{\text{CO}} = \mu_N}_{0,08} \text{ et } \underbrace{\sigma_k = \mu_L}_{0,4} \text{ et } \underbrace{z_k^{\text{CO}} = \mu_E}_{0,6} \text{ alors } z_x^{\text{CO}} = \mu_I \quad (0,08)$$

$$5. \text{ Si } \underbrace{\sigma_x = \mu_L}_1 \text{ et } \underbrace{z_x^{\text{CO}} = \mu_I}_{0,92} \text{ et } \underbrace{\sigma_k = \mu_M}_{0,6} \text{ et } \underbrace{z_k^{\text{CO}} = \mu_I}_{0,4} \text{ alors } z_x^{\text{CO}} = \mu_I \quad (0,4)$$

$$6. \text{ Si } \underbrace{\sigma_x = \mu_L}_1 \text{ et } \underbrace{z_x^{\text{CO}} = \mu_I}_{0,92} \text{ et } \underbrace{\sigma_k = \mu_M}_{0,6} \text{ et } \underbrace{z_k^{\text{CO}} = \mu_E}_{0,6} \text{ alors } z_x^{\text{CO}} = \mu_I \quad (0,6)$$

$$7. \text{ Si } \underbrace{\sigma_x = \mu_L}_1 \text{ et } \underbrace{z_x^{\text{CO}} = \mu_I}_{0,92} \text{ et } \underbrace{\sigma_k = \mu_L}_{0,4} \text{ et } \underbrace{z_k^{\text{CO}} = \mu_I}_{0,4} \text{ alors } z_x^{\text{CO}} = \mu_I \quad (0,4)$$

$$8. \text{ Si } \underbrace{\sigma_x = \mu_L}_1 \text{ et } \underbrace{z_x^{\text{CO}} = \mu_I}_{0,92} \text{ et } \underbrace{\sigma_k = \mu_L}_{0,4} \text{ et } \underbrace{z_k^{\text{CO}} = \mu_E}_{0,6} \text{ alors } z_x^{\text{CO}} = \mu_E \quad (0,4)$$

Les conditions opérationnelles obtenues pour la mission y s'élèvent à 5,32 par la méthode des centroïdes appliquée à la surface obtenue par la formation du maximum. Quand toutes les surfaces sont dénombrées suivant les règles appliquées, la valeur de sortie obtenue est égale à 5,26.

Ainsi, l'ensemble des missions réalisées entre ces deux instants de défaillances peut être modélisé par une seule mission dont les conditions opérationnelles sont de 5,26 pour une durée de 340 ut.

Au sujet des conditions environnementales, elles sont obtenues de la même manière, c'est-à-dire en réitérant le raisonnement par logique floue entre les missions i et j , puis entre les missions x (mission intermédiaire issue des deux premières missions) et la mission k .

Dans notre étude, lorsque les différents temps de bon fonctionnement seront représentés par une seule mission, l'application de la vraisemblance partielle amènera la résolution numérique de deux équations afin d'obtenir les coefficients associés aux variables explicatives (conditions opérationnelles et environnementales). Ces deux équations sont données ci-dessous :

$$\frac{\partial L^*}{\partial \beta^{\text{CO}}} = 0 \quad (\text{II. 11})$$

$$\frac{\partial L^*}{\partial \beta^{\text{CE}}} = 0 \quad (\text{II. 12})$$

Où L^* représente la log-vraisemblance définie par :

$$L^* = \ln(V^*) = \sum_{i=1}^n \left[\left(\beta^{CO} \cdot z_i^{CO} + \beta^{CE} \cdot z_i^{CE} \right) - \ln \left(\sum_{j \in n(t_i)} e^{\left(\beta^{CO} \cdot z_j^{CO} + \beta^{CE} \cdot z_j^{CE} \right)} \right) \right] \quad (\text{II. 13})$$

Dans cette équation, V^* caractérise le produit des contributions pour chaque date de défaillance :

$$V^* = \prod_{i=1}^n V_i = \prod_{i=1}^n \frac{e^{\left(\beta^{CO} \cdot z_i^{CO} + \beta^{CE} \cdot z_i^{CE} \right)}}{\sum_{k \in n(t_i)} e^{\left(\beta^{CO} \cdot z_k^{CO} + \beta^{CE} \cdot z_k^{CE} \right)}} \quad (\text{II. 14})$$

- Avec :
- $t_i \ i = \{1, \dots, n\}$: les n défaillances observées,
- $n(t_i)$: l'ensemble de la population « à risque », c'est-à-dire l'ensemble des composants dont la durée de vie est supérieure à t_i .

VIII. Conclusion

Dans cette étude, le navire est amené à réaliser différentes missions dont les caractéristiques de fonctionnement peuvent différer suivant l'objectif (surveillance côtière, démonstration de force, etc.) et selon le lieu de réalisation. Les missions devant être choisies lorsque le navire se trouve encore à quai et ne disposant d'aucun modèle physique décrivant le fonctionnement du système, seul le pronostic basé sur l'expérience peut être utilisé. Pour ce pronostic, l'évolution des dégradations est modélisée par la loi de défaillance du système. Afin de tenir compte de l'influence des différentes missions, nous avons opté pour le modèle à risques proportionnels suivant lequel la fonction de défaillance est influencée par un ensemble de variables explicatives (conditions de fonctionnement des missions). Ce modèle suppose que les variables explicatives demeurent constantes entre deux défaillances consécutives. Cette contrainte ne pouvant être satisfaite, nous avons opté pour un raisonnement basé sur la logique floue afin de modéliser des conditions de fonctionnement équivalentes à l'ensemble des différentes missions réalisées durant ces périodes de bon fonctionnement. Un exemple, à la fin de ce chapitre, est proposé pour présenter le fonctionnement d'un raisonnement par logique floue et l'application de celui-ci à cette problématique.

Chapitre III

Études et analyses de politiques intégrées

APRES l'étude des différentes stratégies permettant l'optimisation conjointe du plan d'exploitation et de maintenance, différentes politiques de maintenances seront proposées. La première politique étudiée sera minimaliste car uniquement constituée d'actions correctives. Les deux politiques suivantes se caractériseront par la réalisation de maintenances préventives à la fin des missions. Pour la première politique, appelée systématique, des maintenances préventives seront réalisées après chaque mission, alors que pour la politique sporadique, les missions qui doivent être suivies d'actions préventives devront être déterminées. Par la suite, les politiques de types périodiques et séquentiels seront étudiées en vue de définir les intervalles de temps optimaux.

Sommaire

I. Introduction	52
II. Les différentes stratégies d'optimisation	52
III. La politique minimaliste	59
IV. La politique systématique	64
V. La politique sporadique	75
VI. La politique périodique	81
VII. La politique séquentielle	90
VIII. Conclusion	96

I. Introduction

D'après E.M. Goldratt et J. Cox [Goldratt, 2006], le principal but recherché dans le domaine industriel concerne la maximisation des profits. Pour atteindre ce but, il est nécessaire que les différentes activités, telles que la production et la maintenance, coopèrent. Néanmoins, jusqu'à la fin des années 1990, les travaux de recherches traitaient indépendamment les problématiques liées à la gestion de la production [Shapiro, 1993] et à la gestion de la maintenance [Dekker, 1996], comme si aucune relation ne coexistait entre ces deux entités. Depuis, de nombreux travaux se sont intéressés à l'étude conjointe des activités de production et de maintenance parmi lesquels nous pouvons citer les travaux de Bennour [Bennour, 2001] dont l'objectif visait à proposer une modélisation intégrée. L'optimisation de ces deux activités peut être réalisée suivant des stratégies séquentielles ou intégrées. Les stratégies séquentielles se présentent comme l'optimisation successive des deux activités en considérant la première activité comme contrainte pour l'optimisation de la seconde. Généralement, au plan de production optimal est associé à un plan de maintenance satisfaisant au mieux l'objectif recherché [Benbouzid, 2003]. La stratégie intégrée consiste à planifier conjointement les tâches de production et de maintenance [Sortraku, 2005], [Naderi, 2008].

II. Les différentes stratégies d'optimisation

L'optimisation conjointe des plans de production et des plans de maintenance peut être réalisée suivant deux stratégies, à savoir les stratégies séquentielles et intégrées. Le qualificatif de séquentiel ou d'intégré fait référence à la manière de résoudre le problème et non au résultat obtenu.

II.1. La stratégie séquentielle

La stratégie séquentielle se caractérise par une planification successive des plans. Après avoir planifié une première activité, celle-ci est considérée comme une contrainte forte pour la détermination de la seconde activité. L'application de cette stratégie à notre étude implique que la première activité à planifier correspond au plan d'exploitation. En effet, la loi de défaillance associée au système intègre une fonction de risque (modèle à risques proportionnels) qui, pour rappel, permet de modéliser l'influence des missions sur l'évolution du taux de défaillance. Par conséquent, la détermination du plan de maintenance, qui résulte de la loi de défaillance, nécessite la connaissance préalable des missions à réaliser.

L'objectif étant de maximiser les gains, l'optimisation du plan d'exploitation consiste donc à choisir les missions à réaliser de manière à ce que le profit total généré soit maximal. Ce problème d'optimisation est semblable au problème du type *Knapsack*. Ce dernier « modélise une situation analogue au remplissage d'un sac à dos, ne pouvant supporter plus d'un certain poids, avec tout ou partie d'un ensemble d'objets ayant chacun un poids et une valeur. Les objets mis dans le sac à dos doivent maximiser la valeur totale, sans dépasser le poids maximum ». Dans cette étude, le sac à dos représente le navire et les objets symbolisent les missions. Les poids et valeurs associés aux missions correspondent respectivement aux durées et profits de chaque mission.

Mathématiquement, on recherche le vecteur X qui maximise la fonction suivante :

$$Y(\mathcal{X}) = \sum_{i=1}^{\dim(\mathcal{X})} \psi_{x_i} \quad (\text{III. 1})$$

Cependant, la durée de planification étant limitée, il convient de respecter la contrainte ci-dessous :

$$W(\mathcal{X}) = \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} \leq H \quad (\text{III. 2})$$

Ce problème est bien connu de la communauté scientifique car il fait partie des 21 problèmes NP-complets de Karp [Karp, 1972]. Par conséquent, il a fait l'objet de nombreux travaux concernant les résolutions approchées ou exactes. Comme le laisse sous-entendre leur dénomination, les résolutions exactes permettent d'obtenir la solution optimale au problème. À ce titre, l'implémentation de la procédure de séparation et d'évaluation (*Branch and Bound*) proposée par Martello et Toth [Martello, 1990] est considérée comme une référence. Cependant, le nombre de possibilités étant très important, il est souvent plus judicieux d'obtenir des solutions satisfaisantes, certes non optimales, mais en un temps raisonnable. Parmi les différentes résolutions approchées, l'algorithme glouton est certainement le plus simple à implanter. Le principe consiste simplement à ajouter en priorité les missions les plus profitables jusqu'à saturation. On définit par le terme « profitabilité », le profit généré par unité de temps et non le profit engendré par la réalisation de la mission. Parmi les résolutions approchées, les méta-heuristiques permettent d'obtenir une approximation satisfaisante (généralement meilleure que la solution obtenue par l'algorithme glouton) tout en évitant de monopoliser trop de ressources.

Dès lors qu'un plan d'exploitation satisfaisant est déterminé, l'étape suivante consiste à déterminer le plan de maintenance optimal suivant la politique de maintenance choisie. Les différentes politiques de maintenances seront présentées ultérieurement. Mathématiquement, le gain total sera obtenu après la maximisation successive des équations (III. 2) (III. 3).

$$Z(\mathcal{X}) = Y(\mathcal{X}) - \Gamma(\mathcal{X}) \quad (\text{III. 3})$$

Où $\Gamma(\mathcal{X})$ représente, de manière générale, les coûts de maintenance pour le plan d'exploitation défini par le vecteur $\mathcal{X}$, abstraction faite de la politique de maintenance considérée.

Dans les hypothèses, il a été précisé que la réalisation de maintenances préventives nécessitait une durée non-négligeable. Le plan de maintenance est conséquemment contraint par un nombre maximal de maintenances préventives, qui s'exprime par :

$$N^{\max} = 1 + \left\lfloor \frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\mu_p} \right\rfloor \quad (\text{III. 4})$$

Avec :

- $N^{\max}$: Nombre maximal d'intervalles de maintenance ; le nombre de maintenances préventives étant de $N^{\max} - 1$,
- $\lfloor \cdot \rfloor$: Partie entière.

L'optimisation séquentielle présentée pour cette stratégie permet effectivement de maximiser les profits générés par le plan d'exploitation. Cependant, la minimisation des coûts de maintenance étant dépendante du plan d'exploitation et compte tenu du nombre limité d'actions de maintenances préventives, il est peu probable que le plan de maintenance s'avère efficace.

De ce fait, il est raisonnable de penser qu'un plan d'exploitation, qui génère des bénéfices moindres, mais dont le plan de maintenance minimise réellement les coûts, peut s'avérer plus satisfaisant.

II.2. La stratégie intégrée

Idéalement, la stratégie intégrée chercherait à optimiser simultanément le plan d'exploitation et le plan de maintenance. Néanmoins, dans la mesure où le plan de maintenance doit tenir compte des missions à réaliser pour être optimal, il est indispensable que le plan d'exploitation soit défini par avance. Contrairement à la stratégie séquentielle, le choix de la meilleure solution sera obtenu en se basant sur la fonction objective globale.

$$Z(\mathcal{X}) = \sum_{i=1}^{\dim(\mathcal{X})} \psi_{x_i} - \Gamma(\mathcal{X}) \quad (\text{III. 5})$$

Contrairement à la politique séquentielle, la contrainte temporelle devra intégrer la durée des actions de maintenance :

$$W(\mathcal{X}) = \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} + \Upsilon(\mathcal{X}) \leq H \quad (\text{III. 6})$$

Où $\Upsilon(\mathcal{X})$ caractérise, de manière générale, la durée liée aux actions de maintenances pour le plan d'exploitation défini par le vecteur $\mathcal{X}$, abstraction faite de la politique de maintenance considérée.

Comme notre étude ne permet pas une approche « totalement » intégrée, l'idée consiste à utiliser une méta-heuristique et à sélectionner parmi les solutions obtenues, celle qui satisfait au mieux l'équation (III. 5). Ainsi, l'outil d'aide à la décision développé permettra de déterminer des solutions efficaces avec un temps de calcul modéré.

II.2.a. Recherche de solutions efficaces

En parcourant la littérature, il apparaît que de nombreux problèmes de planification et d'ordonnancement sont résolus à l'aide de méta-heuristiques ou algorithmes évolutionnaires. Par exemple, McMullan [McMullan, 2007] se base sur l'algorithme du Grand déluge en raison de sa simplicité et de l'obtention de solutions de bonne qualité. Rajendran [Rajendran, 2004] et Benbouzid-Sitayeb [Benbouzid-Sitayeb, 2006] s'intéressent aux colonies de fourmis pour le problème du PFSP. De nombreux travaux utilisent les algorithmes génétiques pour les problèmes d'ordonnancement [Sortrakul, 2005], [Chung, 2009].

Pour la recherche de solutions efficaces, nous avons choisi d'utiliser un algorithme évolutionnaire de type algorithme génétique en raison de la bonne adaptation à la résolution de ce type de problème. L'application de cet algorithme à notre problème est illustrée par la figure 17.

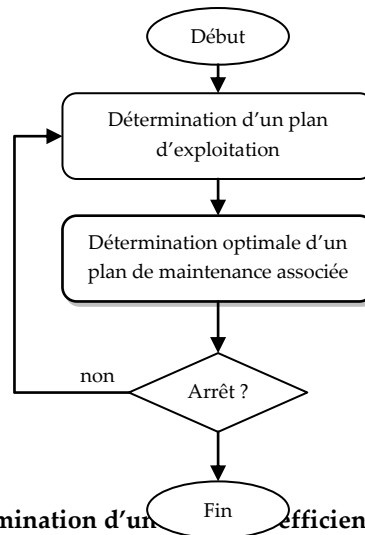


Figure 17 : Principe de détermination d'un plan efficace pour la stratégie intégrée

II.2.b. Principe de fonctionnement des algorithmes génétiques

Les algorithmes génétiques sont des algorithmes évolutionnaires inspirés par la théorie darwinienne. Ils reposent sur le principe de la sélection naturelle permettant aux individus de s'adapter davantage à leur milieu au cours du temps. Partant de cette théorie de l'évolution, Holland [Holland, 1975] fut le premier à proposer l'utilisation de ce principe pour la résolution de problèmes scientifiques. Il aura fallu attendre les travaux de Goldberg [Goldberg, 1989] afin que cette méthode de résolution soit réellement popularisée. Le principe de fonctionnement de ce type d'algorithme peut être représenté schématiquement par la figure suivante :

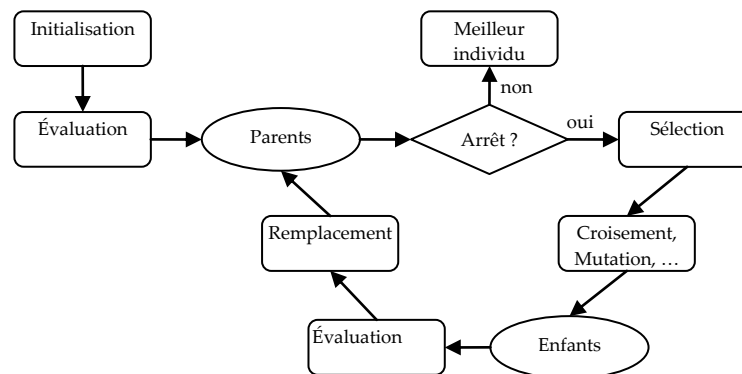


Figure 18 : Principe général de fonctionnement de l'algorithme génétique [Siarry, 2005]

La phase d'initialisation sert à fixer l'ensemble des paramètres utilisés dans l'algorithme génétique, tels que le nombre de générations, la taille des populations, les taux des opérateurs génétiques, etc. Durant cette étape, le type de sélection, les opérateurs génétiques ainsi que les critères d'arrêts sont également définis. Cette phase d'initialisation se termine par la génération de la population initiale qui constituera les premiers individus parents. Par la suite, les parents seront sélectionnés, stochastiquement et en fonction de leurs performances, afin d'engendrer la population suivante grâce aux opérateurs génétiques. Tant que le critère d'arrêt n'est pas atteint, le cycle évolutionnaire continue.

La phase de sélection a pour but de choisir les individus qui participeront à la création de la prochaine génération d'individus. Habituellement, les individus les plus adaptés ont une plus grande probabilité d'être sélectionnés. Pour ces travaux, l'objectif de chaque individu étant de maximiser les profits des plans d'exploitation, nous évaluons donc la fitness de chaque individu par le profit engendré. Dans la littérature, de nombreux modèles de sélection ont été proposés, dont les plus courants sont présentés ci-dessous :

- La sélection proportionnelle (*Roulette Wheel Selection*) est la première méthode implémentée pour sélectionner les individus. Proposée par Goldberg [Goldberg, 1989], elle fait allusion aux roulettes des casinos, dont la surface serait décomposée en autant de secteurs que d'individus contenus dans la population et dont la taille de chacun de ces secteurs serait proportionnelle au fitness des individus concernés. Avec cette méthode, la chance d'être sélectionné dépend directement du fitness de l'individu. Cependant, la variance étant élevée, des méthodes de sélections en sont dérivées, telles que la sélection avec reste stochastique (*Stochastic Universal Sampling*) [Baker, 1987].
- La sélection par tournoi (*Tournament selection*) consiste à choisir aléatoirement un groupe d'individus dans la population et de sélectionner de manière déterministe le meilleur individu de ce groupe. Parmi les sélections par tournoi, la plus couramment utilisée est le tournoi binaire, dont le groupe est constitué simplement de deux individus.
- La sélection par élitisme n'est pas à proprement parler une méthode de sélection. Elle permet néanmoins de conserver les meilleurs individus sur les différentes générations.

Les opérateurs génétiques sont employés pour élargir l'horizon de recherche. Ils se composent d'opérations de croisement et de mutation.

- L'opération de croisement est utilisée pour générer les individus enfants à partir des individus parents. Habituellement, cet opérateur génétique se décline en croisement en un point, en plusieurs points ou uniforme. Le principe consiste à échanger les séquences de gènes des deux parents entre le ou les points de croisement.
- L'opérateur génétique de mutation est utilisé pour éviter le piège des optimaux locaux. Le principe consiste à échanger un gène. Pour les codages binaires, la mutation s'effectue simplement en inversant un gène actif en un gène inactif, ou inversement. Pour les codages basés sur des entiers naturels, le processus de mutation varie.

L'intérêt principal de l'algorithme génétique consiste à obtenir rapidement une solution satisfaisante. Aussi, afin de toujours conserver les meilleures solutions, une procédure élitiste sera instaurée pour la création des nouvelles générations ; les meilleurs individus survivront au cours du temps sans être nécessairement remplacés par des individus enfants.

Pour déterminer les plans d'exploitation et de maintenance optimaux, un algorithme génétique est utilisé. Par la suite, nous allons décrire le processus de cet algorithme génétique pour la détermination du plan d'exploitation. L'optimisation du plan de maintenance sera également développée, suivant chaque politique de maintenance.

II.2.c. Application de l'algorithme génétique à la détermination des plans d'exploitation

L'utilisation de l'algorithme génétique pour l'obtention des plans optimaux nécessite la définition particulière du codage des individus ou des opérations génétiques.

II.2.c.i Le codage

Pour cet algorithme génétique, chaque individu représente un plan d'exploitation, c'est-à-dire un ensemble de missions ordonné, où chaque mission est différente. À ce titre, le chromosome est représenté par le vecteur χ dont chaque composante (c'est-à-dire le gène, au sens des algorithmes génétiques) est constituée d'un naturel qui représente ainsi la mission à réaliser. La durée de l'horizon de temps étant fixée, le nombre de missions peut s'avérer différent d'un plan d'exploitation à un autre. Par conséquent, la taille des chromosomes est variable.

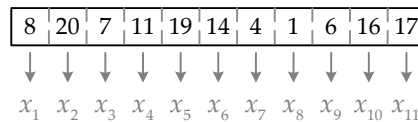


Figure 19 : Codage du chromosome pour l'A.G. associé à la détermination des plans d'exploitation

Dans l'exemple présenté ci-dessus, le plan d'exploitation est constitué de onze missions. La première mission à réaliser (x_1) correspond à la mission 8, la deuxième mission à réaliser (x_2) est la mission 20, et ainsi de suite pour l'ensemble des missions restantes. La viabilité de ce chromosome suppose que la durée cumulée de ces onze missions ($\delta_8 + \delta_{20} + \dots + \delta_6 + \delta_{17}$) est inférieure à H . Elle peut également s'exprimer sous la forme de l'équation (III. 2) où $\dim(\chi) = 11$.

II.2.c.ii La stratégie de reproduction

Pour former une nouvelle génération, certains individus doivent être sélectionnés parmi les parents en vue d'être modifiés. La sélection par tournoi déterministe est appliquée afin de faire apparaître les individus qui auront le privilège d'engendrer la génération suivante. Cette sélection se caractérise par la réalisation de tournois réalisés parmi des individus choisis aléatoirement. Pour chaque tournoi, l'individu ayant le meilleur fitness participera à la reproduction des individus enfants par l'intermédiaire des opérations génétiques de croisement et/ou mutation. Dans notre cas, le fitness correspond au profit dégagé par le plan d'exploitation minoré des coûts liés aux actions de maintenances. Par ailleurs, afin de toujours conserver les meilleurs individus, une sélection élitiste sera appliquée et pour élargir l'horizon de recherches (et ainsi éviter le piège des optimaux locaux), quelques nouveaux individus seront générés aléatoirement.

II.2.c.iii L'opérateur de croisement

Parmi les différentes méthodes de croisement présentées précédemment, le croisement en un point sera employé. Ce croisement consiste à sélectionner aléatoirement un gène et à inverser, à partir de ce point de croisement, les gènes entre les deux parents. Dans le codage employé, les individus peuvent être de tailles différentes. Il convient donc de s'assurer que le point de croisement existe dans les deux chromosomes des individus parents. Le gène, après lequel les séquences seront inversées, devra être compris dans l'intervalle $]0, \min(\dim(\chi_1), \dim(\chi_2))]$ où $i = \{1, 2\}$ représente

l'individu « Parent i ». La borne supérieure de l'intervalle correspond à la dimension du vecteur du plus petit parent. Ainsi, le point de croisement existera chez les deux parents.

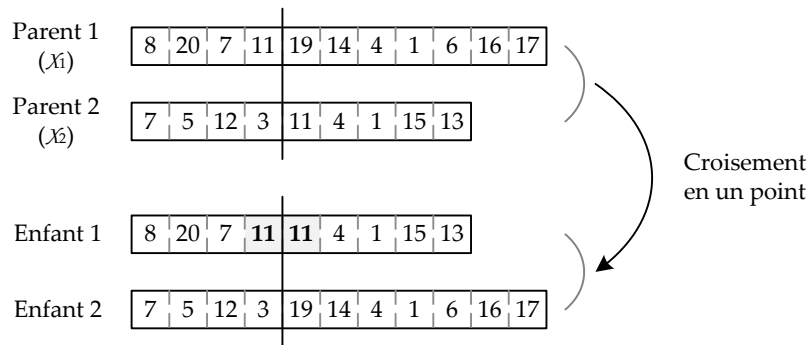


Figure 20 : Croisement en un point, après le quatrième gène

Dans l'exemple présenté sur la figure 20, l'individu « Enfant 1 » est constitué de la partie gauche du « Parent 1 » ainsi que de la partie droite du « Parent 2 ». Au contraire, le chromosome « Enfant 2 » est issu du croisement opposé : le début du « Parent 2 » et la fin du « Parent 1 ». Au niveau de l'individu « Enfant 1 », nous pouvons noter la présence d'une erreur génétique. En effet, la mission 11 est présente deux fois dans le plan d'exploitation, ce qui est contraire aux hypothèses de travail. De manière moins perceptible, une autre erreur peut survenir. Bien que les deux parents respectent la contrainte temporelle, il se peut que les durées des plans d'exploitation définies par les nouveaux plans d'exploitations soient supérieures à l'horizon de temps alloué. Ainsi, nous pouvons constater que l'opération génétique de croisement peut engendrer des individus non fiables, en ne respectant ni l'unicité des missions, ni la durée maximale d'un plan d'exploitation. Pour pallier ces erreurs, il convient d'ajouter, à ces opérations de croisement, une phase de vérification et de correction.

Lorsque l'erreur génétique se caractérise par la duplication d'une ou plusieurs missions au sein d'un chromosome, la correction s'effectue en remplaçant les missions qui apparaissent plusieurs fois par d'autres missions possibles. On entend par missions possibles, des missions appartenant au vecteur $\mathcal{M}$, mais bien évidemment absentes du chromosome. Lorsque des missions apparaissent plusieurs fois, le choix de celles à remplacer ainsi que celui des missions remplaçantes seront réalisés aléatoirement. En supposant que dans notre exemple, vingt missions soient proposées, l'individu 'Enfant 1' peut être corrigé en remplaçant le 4^{ème} ou le 5^{ème} gène par une des missions suivantes {2, 3, 5, 6, 9, 10, 12, 14, 16, 17, 18, 19}, absente de l'individu considéré. Pour la correction de cet individu, 24 possibilités sont offertes.

Dans le cas où l'individu n'est pas viable à cause de la durée du plan d'exploitation supérieure à l'horizon de temps fini, la méthode consiste à supprimer des missions aléatoirement jusqu'à satisfaire cette contrainte temporelle. Cependant et contrairement à la violation de cette contrainte temporelle, l'opérateur de croisement peut engendrer des individus viables, mais dont la différence entre la durée totale du plan d'exploitation et l'horizon de temps permettrait la réalisation de missions supplémentaires. Pour palier ce manque de missions et donc améliorer le profit, la méthode consiste à ajouter aléatoirement des missions possibles jusqu'à ce que la durée cumulée des missions se rapproche de la durée de l'horizon de temps, tout en respectant la contrainte temporelle.

II.2.c.iv L'opérateur de mutation

Comme l'opérateur de croisement, la mutation est une opération génétique. Toutefois, l'opérateur de mutation est généralement utilisé pour éviter et sortir des pièges liés aux optimaux locaux. Compte tenu que la structure des individus est constituée de naturels, l'opérateur génétique de mutation consiste à remplacer un gène (une mission) choisi aléatoirement, par une mission possible, c'est-à-dire présente dans le vecteur $\mathcal{M}$ mais absente du chromosome. Bien que les missions présentes dans l'individu soumis à cet opérateur génétique soient toutes différentes, la viabilité du plan d'exploitation n'est pas assurée. En effet, le changement d'une mission peut induire une durée supérieure du plan d'exploitation et par conséquent une violation de la contrainte temporelle. Pour corriger cet incident, la même procédure appliquée à l'opérateur de croisement peut être utilisée. À l'inverse, si la durée est inférieure au plan d'exploitation précédent, il convient de vérifier si une mission supplémentaire peut être ajoutée afin d'améliorer le profit.

Dans cette première partie, nous nous sommes intéressés aux deux stratégies fréquemment rencontrées dans la littérature. Pour la stratégie séquentielle, nous avons proposé une approche basée sur un algorithme génétique et dont l'optimisation porte uniquement sur les plans d'exploitation. Dans la suite de ce chapitre, nous allons nous intéresser à l'étude et à l'analyse de différentes politiques de maintenances constituées d'actions de maintenances correctives et préventives. Dans les hypothèses de travail, il a été précisé qu'au départ du quai, le système était considéré dans un état *As Good As New*. Cet état du système est lié à une maintenance préventive parfaite réalisée à la fin du plan d'exploitation. Cette action préventive étant obligatoire et indépendante des politiques de maintenances, nous ignorerons volontairement cette action dans la suite du mémoire.

III. La politique minimaliste

La première politique de maintenance considérée est une politique minimaliste. Son but est d'évaluer l'apport des politiques de maintenances améliorées. Cette politique de maintenance considérée est exclusivement constituée d'actions de maintenances correctives. Ces activités sont réalisées pour pallier les défaillances immobilisantes qui surgissent au cours du temps. De par leur réalisation en mer, l'effet de ces maintenances correctives est minimal ; leur rôle étant uniquement de restaurer le navire dans un état de fonctionnement. Le coût de cette politique, qui dépend du plan d'exploitation défini par le vecteur χ , s'exprime par la relation suivante :

$$\Gamma_{Min}(\chi) = M_C \cdot \phi_{Min}(\chi) \quad (\text{III. 7})$$

Avec :

- $\Gamma_{Min}(\chi)$: coût de maintenance associé à la politique minimale, pour le plan d'exploitation défini par le vecteur χ ,
- $\phi_{Min}(\chi)$: nombre moyen de pannes pouvant survenir pendant le plan d'exploitation défini par le vecteur χ et pour la politique minimaliste.

Généralement, le nombre moyen de pannes, dans le cas de maintenances minimales, s'exprime pendant une durée définie et sous des conditions de fonctionnement supposées constantes au cours

du temps. Sous ces hypothèses, le nombre moyen de pannes pendant une durée T s'exprime par la relation :

$$\phi(T) = \int_0^T \lambda(t) dt \quad (\text{III. 8})$$

La différence majeure avec notre travail repose sur l'utilisation de conditions de fonctionnement qui peuvent varier au cours du temps. En supposant les différentes missions définies par le vecteur $\mathcal{X}$, il semblerait « logique » de définir le nombre moyen de pannes, comme suit :

$$\phi_A(\mathcal{X}) = \int_0^{t_{x_1}} \lambda_{x_1}(t) dt + \int_{t_{x_1}}^{t_{x_2}} \lambda_{x_2}(t) dt + \dots + \int_{t_{x_{\dim(\mathcal{X})-1}}}^{t_{x_{\dim(\mathcal{X})}}} \lambda_{x_{\dim(\mathcal{X})}}(t) dt \quad (\text{III. 9})$$

Avec :

- $t_{x_j} = t_{x_{j-1}} + \delta_{x_j}$, avec $t_{x_0} = 0$, date caractérisant l'âge du système à la fin de la mission x_j ,
- $\lambda_{x_j}(t)$: Loi de défaillance associée à la mission x_j , c'est-à-dire aux conditions de fonctionnement Z_{x_j} , $j = \{1, 2, \dots, \dim(\mathcal{X})\}$.

Cependant, la relation définie par l'équation (III. 9) s'avère inexacte ; elle ne vérifie pas la continuité de la fonction de fiabilité. En effet, si deux missions successives, notées x_i et x_{i+1} , possèdent des conditions de fonctionnement différentes, la continuité de la fonction de fiabilité, sous le modèle à risques proportionnels, imposerait que la relation suivante soit vérifiée.

$$\left[R(t_{x_i}) \right]^{g(Z_{x_i})} = \left[R(t_{x_i}) \right]^{g(Z_{x_{i+1}})} \quad (\text{III. 10})$$

Or, pour vérifier cette équation tout en sachant que les conditions de fonctionnement sont immuables et différentes entre ces deux missions, la seule modification possible concerne la date de début de la mission x_{i+1} . En d'autres termes, si le navire a fonctionné sous des conditions de fonctionnement Z_{x_i} pendant une durée δ_{x_i} , c'est équivalent au fonctionnement du navire sous des conditions $Z_{x_{i+1}}$, mais pendant une durée différente, notée d_{x_i} . Cette nouvelle durée caractérise l'âge fonctionnel du système et l'équation (III. 10) devient :

$$\left[R(t_{x_i}) \right]^{g(Z_{x_i})} = \left[R(\tau_{x_i}) \right]^{g(Z_{x_{i+1}})} \quad (\text{III. 11})$$

Où $\tau_{x_i} = \sum_{j=1}^i d_{x_j}$, $j = \{1, 2, \dots, \dim(\mathcal{X})\}$ est une date représentative de l'âge fonctionnel du système au début de la mission x_{i+1} .

III.1. L'âge fonctionnel

Bien que la fonction du taux de défaillance varie en fonction des missions réalisées, la fiabilité du système se doit d'être continue au cours du temps. Le changement de mission impliquant un changement de la fonction de défaillance, cette variation se répercute sur la fonction de fiabilité. Ainsi, l'âge fonctionnel a pour rôle de garantir la continuité de la fiabilité en définissant une durée

représentative de l'âge du système. Par conséquent, la détermination de l'âge fonctionnel nécessite la connaissance de la fiabilité à la fin de la mission précédente, ainsi que les conditions de fonctionnement pour la mission antécédente et la mission subséquente. L'âge fonctionnel τ_{x_i} , associé à la mission x_{i+1} , doit donc vérifier la relation :

$$\left[R(t_{x_i}) \right]^{g(Z_{x_i})} = \left[R(\tau_{x_i}) \right]^{g(Z_{x_{i+1}})} \quad (\text{III. 12})$$

$$\tau_{x_i} = \begin{cases} 0 & \text{si } i = 0 \\ R^{-1} \left(\left[R(t_{x_i}) \right]^{g(Z_{x_i})} \right)^{g(Z_{x_{i+1}})} & \text{sinon} \end{cases} \quad (\text{III. 13})$$

La figure suivante illustre donc la notion d'âge fonctionnel. Deux types de missions, avec des caractéristiques de fonctionnement différentes, sont représentés par les courbes de fiabilité $R_1(t)$ et $R_2(t)$; les notations $R_i(t)$ et $[R(t)]^{g(Z_i)}$ étant équivalentes. La fiabilité du système, à la fin de la réalisation de la première mission (« Mission 1 »), s'élève à $R_1(t_1)$. Pour garantir cette fiabilité au début de la seconde mission (« Mission 2 »), nous supposons que le système a fonctionné durant une période d_1 sous les conditions de fonctionnement associées à la seconde mission.

Les différentes étapes pour la détermination de l'âge fonctionnel se décomposent comme suit :

- Détermination de la date t_i à partir de la durée de la mission δ_i .
- Détermination de la fiabilité $R_1(t_1)$.
- Détermination de la date τ_1 à partir de la fiabilité $R_1(t_1) = R_2(\tau_1)$.

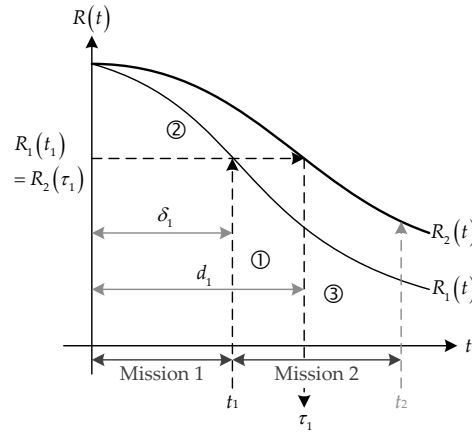


Figure 21 : Détermination de l'âge fonctionnel

À titre informatif, nous pouvons constater, à travers les deux figures suivantes, que la continuité de la fonction de fiabilité n'implique pas la continuité de la fonction du taux de défaillance.

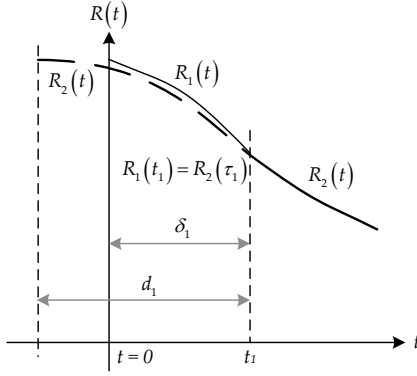


Figure 22 : Évolution de la fiabilité

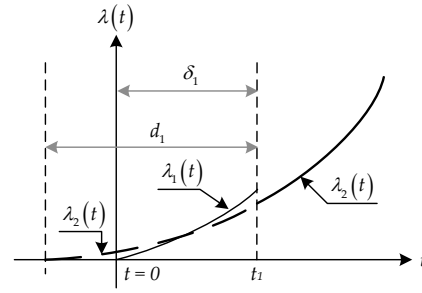


Figure 23 : Évolution du taux de défaillance

III.1.a. Exemple illustratif

Considérons un système soumis à une loi de défaillance nominale de type Weibull, de paramètre de forme (β) 2,5 et d'échelle (η) 600. Soit deux missions, notées m_1 et m_2 , à réaliser et dont les différentes caractéristiques sont données ci-dessous :

- Durée des missions : $\delta_{m_1} = 80$ ut et $\delta_{m_2} = 50$ ut ,
- Fonction de risque de défaillance : $g(Z_{m_1}) = 1,2$ et $g(Z_{m_2}) = 1,5$,

Le taux de fiabilité à la fin de la première mission s'élève à 0,992.

$$\begin{aligned} R_{m_1}(t_{m_1}) &= e^{-g(Z_{m_1}) \left(\frac{t_{m_1}}{\eta} \right)^\beta} \\ &= e^{-1,2 \left(\frac{80}{600} \right)^{2,5}} = 0,992 \end{aligned} \quad (\text{III. 14})$$

Compte tenu de la fonction de risque pour la seconde mission, l'âge fonctionnel est de 73,169 ut. La fonction de risque de défaillance associée à la seconde mission étant plus importante que celle de la première mission, il est logique que l'âge fonctionnel soit inférieur.

$$\begin{aligned} \tau_{m_1} &= R^{-1} \left(\left[R(t_{m_1}) \right]^{\frac{g(Z_{m_1})}{g(Z_{m_2})}} \right) = t_{m_1} \cdot \left(\frac{g(Z_{m_1})}{g(Z_{m_2})} \right)^{\left(\frac{1}{\beta} \right)} \\ &= 80 \cdot \left(\frac{1,2}{1,5} \right)^{1/2,5} = 73,169 \end{aligned} \quad (\text{III. 15})$$

En se basant sur les valeurs obtenues, le taux de défaillance à la fin de la première mission s'élève à $7,698 \cdot 10^{-3}$, alors qu'au début de la mission suivante, il est de $8,417 \cdot 10^{-3}$. Ainsi, il n'y a pas continuité de la fonction de défaillance.

$$\begin{aligned} \lambda_{m_1}(t_{m_1}) &= g(Z_{m_1}) \left(\frac{\beta}{\eta} \right) \cdot \left(\frac{t_{m_1}}{\eta} \right)^{\beta-1} \\ &= 1,2 \cdot \left(\frac{2,5}{600} \right) \cdot \left(\frac{80}{600} \right)^{1,5} = 7,698 \cdot 10^{-3} \end{aligned} \quad (\text{III. 16})$$

$$\begin{aligned}
\lambda_{m_2}(\tau_{m_1}) &= g(Z_{m_2}) \left(\frac{\beta}{\eta} \right) \cdot \left(\frac{\tau_{m_1}}{\eta} \right)^{\beta-1} \\
&= 1,5 \cdot \left(\frac{2,5}{600} \right) \cdot \left(\frac{73,169}{600} \right)^{1,5} = 8,417 \cdot 10^{-3}
\end{aligned}
\tag{III. 17}$$

À partir de la définition de l'âge fonctionnel, le nombre moyen de pannes pouvant survenir durant le plan d'exploitation $\mathcal{X}$ et pour la politique minimaliste se définit comme suit :

$$\phi_{Min}(\mathcal{X}) = \sum_{i=1}^{\dim(\mathcal{X})} \int_{\tau_{x_{i-1}}}^{\tau_{x_i-1} + \delta_{x_i}} \lambda_{x_i}(t) dt
\tag{III. 18}$$

Avec τ_{x_i} l'âge fonctionnel comme défini à l'équation (III. 13).

Cette politique minimaliste étant constituée uniquement d'actions correctives, aucune optimisation n'est possible. Les différentes politiques de maintenances proposées par la suite devront, pour être avantageuses, avoir un coût total inférieur à celui de cette politique de maintenance.

Comme le modèle utilisé pour représenter les conditions de fonctionnement est semi-paramétrique et ne dépend donc pas du temps, les coûts liés à cette politique de maintenance pour un plan d'exploitation donné sont valables pour les différents plans d'exploitations constitués des mêmes missions. En considérant les valeurs numériques précédentes, nous allons comparer le nombre moyen de pannes dans le cas où la mission m_2 succède à la mission m_1 , et la situation inverse.

En considérant la fonction de défaillance de type Weibull, le nombre moyen de pannes entre deux dates s'exprime par :

$$\int_X^Y \lambda_{x_i}(t) = \frac{g(Z_{x_i})}{\eta^\beta} \cdot (Y^\beta - X^\beta)
\tag{III. 19}$$

Ainsi, lorsque la mission m_1 est suivie de la mission m_2 , le nombre moyen de pannes s'élève à $2,864 \cdot 10^{-2}$. Dans le cas inverse, le nombre moyen de pannes reste identique.

$$\begin{aligned}
m_1 \rightarrow m_2 : & \frac{g(Z_{m_1})}{\eta^\beta} \cdot (t_{m_1}^\beta) + \frac{g(Z_{m_2})}{\eta^\beta} \cdot \left((\tau_{m_1} + \delta_{m_2})^\beta - \tau_{m_1}^\beta \right) \\
& \frac{1,2}{600^{2,5}} \cdot (80^{2,5}) + \frac{1,5}{600^{2,5}} \cdot \left((73,169 + 50)^{2,5} - 73,169^{2,5} \right) = 2,864 \cdot 10^{-2}
\end{aligned}
\tag{III. 20}$$

$$\begin{aligned}
m_2 \rightarrow m_1 : & \frac{g(Z_{m_2})}{\eta^\beta} \cdot (t_{m_2}^\beta) + \frac{g(Z_{m_1})}{\eta^\beta} \cdot \left((\tau_{m_2} + \delta_{m_1})^\beta - \tau_{m_2}^\beta \right) \\
& \frac{1,5}{600^{2,5}} \cdot (50^{2,5}) + \frac{1,2}{600^{2,5}} \cdot \left((54,668 + 80)^{2,5} - 54,668^{2,5} \right) = 2,864 \cdot 10^{-2}
\end{aligned}
\tag{III. 21}$$

III.2. Modélisation des maintenances préventives

Dans le but de réduire l'occurrence des défaillances et de minimiser les coûts relatifs aux actions de maintenance, des politiques de maintenances améliorées vont être proposées. Ces politiques se

caractérisent par la réalisation de maintenance préventive. Cependant, elles seront considérées comme imparfaites, car les conditions de réalisation en mer ne sont pas optimales. Par conséquent, après une maintenance préventive, l'état du système ne sera pas considéré comme étant neuf, mais se situera dans un état intermédiaire, c'est-à-dire entre l'état avant la maintenance préventive et l'état neuf du système.

Parmi les différents modèles de maintenance préventive existants, nous avons décidé d'utiliser le modèle ARA_x [Doyen, 2004B]. Pour mémoire, cette modélisation consiste à réduire l'âge du système d'une quantité proportionnelle à son âge juste avant cette action de maintenance. En supposant que l'âge du système juste avant la maintenance préventive s'élève à A_i^- , cette maintenance préventive a pour effet de ramener le système dans un état précédent, c'est-à-dire que l'âge « virtuel » du système est donné par $(1-\rho) \cdot A_i^-$ où ρ représente le facteur d'amélioration. L'intérêt de ce facteur ρ est de sécuriser la prise de décision dans la génération du plan de maintenance. C'est au décideur d'apprécier la qualité des actions de maintenances réalisées, par ses équipes, à un instant donné et d'intégrer cette appréciation dans la génération de la suite du plan de maintenance. Autrement dit, il s'agit d'un facteur de sécurité au niveau de la génération du planning des opérations de maintenance. Dans le reste du mémoire, on parlera de la notion de $\rho^{\max}$ pour représenter la valeur maximale admissible pour que la politique de maintenance générée soit économiquement intéressante et réalisable dans le temps par rapport à la politique de maintenance minimaliste. Le ρ^* optimal caractérisera, quant à lui, la valeur optimale prise par le facteur d'amélioration pour les différentes politiques de maintenances adoptées, à savoir systématique, sporadique, périodique et séquentielle.

IV. La politique systématique

Cette première politique de maintenance améliorée consiste à planifier des maintenances préventives systématiquement après la réalisation de chaque mission. Plus exactement, il s'agit de réaliser une maintenance préventive après les $(\dim(\mathcal{X}) - 1)$ premières missions. En effet, une maintenance préventive parfaite étant réalisée au retour à quai, il est inutile de prévoir une maintenance préventive après la dernière mission. Pour cette politique, les coûts totaux de maintenance sont spécifiés par la relation suivante :

$$\Gamma_{Sys}(\mathcal{X}) = M_C \cdot \phi_{Sys}(\mathcal{X}) + M_P \cdot (\dim(\mathcal{X}) - 1) \quad (\text{III. 22})$$

Avec :

- $\Gamma_{Sys}(\mathcal{X})$: coût de maintenance associé à la politique systématique, pour le plan d'exploitation défini par le vecteur $\mathcal{X}$,
- $\phi_{Sys}(\mathcal{X})$: nombre moyen de pannes pouvant survenir pendant le plan d'exploitation défini par le vecteur $\mathcal{X}$ et pour la politique systématique.

Dans cette équation (III. 22), le nombre moyen de pannes s'exprime comme suit :

$$\phi_{Sys}(\mathcal{X}) = \sum_{i=1}^{\dim(\mathcal{X})} \int_{\tau_{x_{i-1}}}^{\tau_{x_i-1} + \delta_{x_i}} \lambda_{x_i}(t) dt \quad (\text{III. 23})$$

D'après l'équation (III. 23), l'expression du nombre moyen de pannes peut sembler identique à l'expression du nombre moyen de pannes de la politique minimaliste (III. 18). Cependant, la différence se situe au niveau de l'expression de l'âge fonctionnel. Les maintenances préventives étant réalisées après les missions, l'âge fonctionnel avant chaque mission s'exprime donc par :

$$\tau_{x_i} = \begin{cases} 0 & \text{si } i = 0 \\ (1 - \rho) \cdot R^{-1} \left(\left[R(t_{x_i}) \right]^{\frac{g(Z_{x_i})}{g(Z_{x_{i+1}})}} \right) & \text{sinon} \end{cases} \quad (\text{III. 24})$$

Le terme $(1 - \rho)$ réduit l'âge fonctionnel de chaque mission suivant le facteur d'amélioration constant ρ . Par conséquent, la fiabilité du système est améliorée. Les figures suivantes permettent d'illustrer l'influence de ces maintenances préventives sur le taux de défaillance ainsi que la fiabilité du système. Pour permettre une meilleure compréhension, les illustrations de gauche représentent l'âge fonctionnel dans le cas de la politique minimaliste, alors que les figures de droite illustrent l'âge fonctionnel après la réalisation d'une maintenance préventive à la fin d'une mission.

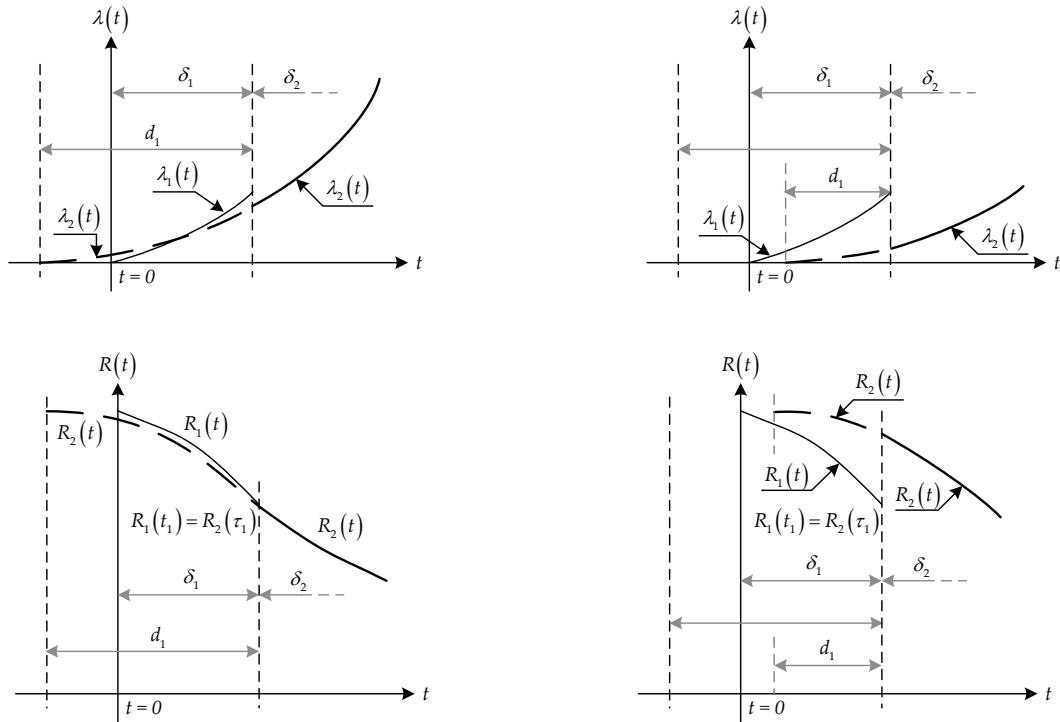


Figure 24 : Modélisation de l'âge fonctionnel pour des missions suivies de maintenances préventives

Pour l'applicabilité de cette politique de maintenance aux stratégies séquentielle et intégrée, il convient donc, lors de la phase de détermination du plan d'exploitation, de réserver une période pour la réalisation des maintenances préventives. La relation temporelle à satisfaire est donc exprimée par :

$$\sum_{i=1}^{\dim(\mathcal{K})} \delta_{k_i} + \mu_p \cdot (\dim(\mathcal{K}) - 1) \leq H \quad (\text{III. 25})$$

Contrairement à la politique minimaliste, lorsque des plans d'exploitation sont constitués des mêmes missions, le nombre moyen de pannes n'est pas constant, en raison des maintenances préventives. Pour illustrer ce propos, nous considérons les valeurs numériques présentées dans la politique minimaliste et nous considérons un facteur d'amélioration $\rho = 0,5$.

Ainsi, lorsque la mission m_1 est suivie de la mission m_2 , alors qu'une maintenance préventive a été réalisée après la mission m_1 , le nombre moyen de pannes s'élève à $1,148 \cdot 10^{-2}$. Dans le cas inverse, le nombre moyen de pannes est de $0,754 \cdot 10^{-2}$.

$$m_1 \rightarrow m_2 : \frac{g(Z_{m_1})}{\eta^\beta} \cdot (t_{m_1}^\beta) + \frac{g(Z_{m_2})}{\eta^\beta} \cdot \left(\left[(1-\rho) \cdot (\tau_{m_1} + \delta_{m_2}) \right]^\beta - \left[(1-\rho) \cdot \tau_{m_1} \right]^\beta \right) \quad (\text{III. 26})$$

$$\frac{1,2}{600^{2,5}} \cdot (80^{2,5}) + \frac{1,5}{600^{2,5}} \cdot \left(\left[0,5 \cdot (73,169 + 50) \right]^{2,5} - \left[0,5 \cdot 73,169 \right]^{2,5} \right) = 1,148 \cdot 10^{-2}$$

$$m_2 \rightarrow m_1 : \frac{g(Z_{m_2})}{\eta^\beta} \cdot (t_{m_2}^\beta) + \frac{g(Z_{m_1})}{\eta^\beta} \cdot \left(\left[(1-\rho) \cdot (\tau_{m_2} + \delta_{m_1}) \right]^\beta - \left[(1-\rho) \cdot \tau_{m_2} \right]^\beta \right) \quad (\text{III. 27})$$

$$\frac{1,5}{600^{2,5}} \cdot (50^{2,5}) + \frac{1,2}{600^{2,5}} \cdot \left(\left[0,5 \cdot (54,668 + 80) \right]^{2,5} - \left[0,5 \cdot 54,668 \right]^{2,5} \right) = 0,754 \cdot 10^{-2}$$

Pour que cette politique de maintenance s'avère plus intéressante que la politique minimaliste, il est nécessaire que les coûts de maintenances de cette politique soient inférieurs à ceux de la politique précédente. Mathématiquement, il convient de vérifier l'expression suivante :

$$M_C \cdot (\phi_{Min}(\mathcal{K}) - \phi_{Sys}(\mathcal{K})) > M_P \cdot (\dim(\mathcal{K}) - 1) \quad (\text{III. 28})$$

Pour cette politique, le facteur d'amélioration et le coût d'une action préventive sont supposés déterminés. Cependant, Arango *et al.* [Arango, 2006] ont montré que l'efficacité d'une maintenance préventive dépend d'un ensemble de facteurs. Aussi, cette politique de maintenance pourrait être améliorée en recherchant la valeur optimale prise par le facteur d'amélioration.

IV.1. Recherche du facteur d'amélioration optimal

La variation du facteur d'amélioration peut dépendre de différents paramètres. Par exemple, la compétence des opérateurs peut intervenir sur la qualité des activités de maintenances, mais le coût de l'action de maintenance varie également en fonction de ces compétences. Ainsi, la variation de la valeur du facteur d'amélioration impacte le coût des actions de maintenance préventive. Suivant ce même concept, la variation du facteur d'amélioration peut avoir une incidence sur la durée des actions de maintenances préventives.

IV.1.a. Modélisation du coût et de la durée des actions de maintenances préventives

Il est raisonnable de penser que le coût d'une action de maintenance est fonction de son efficacité. Cependant, ces coûts ne sont pas proportionnels à l'efficacité ; plus une maintenance tendra vers une action parfaite, plus la différence de coût sera élevée. Par ailleurs, le coût d'une action qui n'améliore pas le système ne peut pas être considéré comme nul. De ce fait, une action préventive est constituée

d'un coût fixe (pièces et produits utilisés) et d'un coût variable (niveau de qualification, expérience). Pour satisfaire ces contraintes, le modèle utilisé est défini par :

$$M_p(\rho) = M_p^f + (\rho \cdot M_p^v)^2 \quad (\text{III. 29})$$

Avec :

- $M_p(\rho)$: coût de l'action de maintenance préventive en fonction du facteur d'amélioration ρ ,
- ρ : facteur d'amélioration dont la valeur appartient à l'intervalle]0, 1[,
- M_p^f : coût fixe d'une maintenance préventive,
- M_p^v : coût variable d'une maintenance préventive.

Similairement à l'expression du coût des actions de maintenances préventives, la durée peut se constituer d'une durée fixe et d'une durée variable, comme l'illustre l'équation ci-dessous :

$$\mu_p(\rho) = \mu_p^f + (\rho \cdot \mu_p^v)^2 \quad (\text{III. 30})$$

Avec :

- $\mu_p(\rho)$: durée de la maintenance préventive en fonction du facteur d'amélioration ρ ,
- ρ : facteur d'amélioration dont la valeur appartient à l'intervalle]0, 1[,
- μ_p^f : durée fixe d'une maintenance préventive,
- μ_p^v : durée variable d'une maintenance préventive.

L'utilisation d'un coût de maintenance préventive dépendant du facteur d'amélioration induit quelques changements au niveau de l'expression du coût total de maintenance :

$$\Gamma_{Sys}(\mathcal{X}, \rho) = M_c \cdot \phi_{Sys}(\mathcal{X}, \rho) + M_p(\rho) \cdot (\dim(\mathcal{K}) - 1) \quad (\text{III. 31})$$

Avec :

- $\Gamma_{Sys}(\mathcal{X}, \rho)$: coût de maintenance associé à la politique systématique, pour un plan d'exploitation défini par le vecteur $\mathcal{X}$ et un facteur d'amélioration ρ ,
- $\phi_{Sys}(\mathcal{X}, \rho)$: nombre moyen de pannes lié à la politique stratégique, pour un plan d'exploitation défini par le vecteur $\mathcal{X}$ et un facteur d'amélioration ρ .

Le nombre moyen de pannes s'exprime par :

$$\phi_{Sys}(\mathcal{K}, \rho) = \sum_{i=1}^{\dim(\mathcal{X})} \int_{\tau_{x_{i-1}}(\rho)}^{\tau_{x_i}(\rho) + \delta_{x_i}} \lambda_{x_i}(t) dt \quad (\text{III. 32})$$

Où $\tau_{k_{i-1}}(\rho)$ représente l'âge fonctionnel en fonction du facteur d'amélioration ρ et dont l'expression est donnée ci-dessous :

$$\tau_{x_i}(\rho) = \begin{cases} 0 & \text{si } i = 0 \\ (1-\rho) \cdot R^{-1} \left(\left[R(t_{x_i}) \right]^{\frac{g(Z_{x_i})}{g(Z_{x_{i+1}})}} \right) & \text{sinon} \end{cases} \quad (\text{III. 33})$$

Dans cette équation (III. 33), et en raison des maintenances préventives réalisées, $t_{x_i} = \tau_{x_{i-1}} + \delta_{x_i}$.

Lorsque les coûts et/ou les durées dépendent du facteur d'amélioration, l'objectif vise à déterminer la valeur optimale de ce facteur d'amélioration, notée ρ^* . La valeur optimale est celle qui minimise les coûts de maintenance et est obtenue analytiquement par la résolution de l'équation ci-dessous :

$$\frac{\partial \Gamma_{Sys}(\mathcal{X}, \rho)}{\partial \rho} = 0 \quad (\text{III. 34})$$

Nonobstant, la complexité des bornes de l'intégrale de $\phi_{Sys}(\mathcal{X}, \rho)$ ne permet pas la résolution analytique de l'équation (III. 34). Pour pallier cette difficulté, ce facteur sera estimé à l'aide d'une procédure numérique. Le lemme suivant assure l'existence d'un minimum local :

$$\exists \rho^* \text{ si } \frac{(\phi_{Sys}(\mathcal{X}, \rho) - \phi_{Sys}(\mathcal{X}, \rho^-))}{((\rho)^2 - (\rho^-)^2)} \leq -(\dim(\mathcal{X}) - 1) \cdot \frac{(M_p^v)^2}{M_c} \leq \frac{(\phi_{Sys}(\mathcal{X}, \rho^+) - \phi_{Sys}(\mathcal{X}, \rho))}{((\rho^+)^2 - (\rho)^2)} \quad (\text{III. 35})$$

La démonstration de ce lemme est donnée par :

$$\begin{aligned} & \Gamma_{Sys}(\mathcal{X}, \rho^+) - \Gamma_{Sys}(\mathcal{X}, \rho) \geq 0 \\ & \left[M_c \cdot \phi_{Sys}(\mathcal{X}, \rho^+) + M_p(\rho^+) \cdot (\dim(\mathcal{X}) - 1) \right] - \left[M_c \cdot \phi_{Sys}(\mathcal{X}, \rho) + M_p(\rho) \cdot (\dim(\mathcal{X}) - 1) \right] \geq 0 \\ & \left[M_c \cdot (\phi_{Sys}(\mathcal{X}, \rho^+) - \phi_{Sys}(\mathcal{X}, \rho)) \right] + (\dim(\mathcal{X}) - 1) \cdot (M_p^v)^2 \cdot \left[(\rho^+)^2 - (\rho)^2 \right] \geq 0 \\ & \frac{(\phi_{Sys}(\mathcal{X}, \rho^+) - \phi_{Sys}(\mathcal{X}, \rho))}{((\rho^+)^2 - (\rho)^2)} \geq -(\dim(\mathcal{X}) - 1) \cdot \frac{(M_p^v)^2}{M_c} \end{aligned} \quad (\text{III. 36})$$

$$\begin{aligned} & \Gamma_{Sys}(\mathcal{X}, \rho) - \Gamma_{Sys}(\mathcal{X}, \rho^-) \leq 0 \\ & \left[M_c \cdot \phi_{Sys}(\mathcal{X}, \rho) + M_p(\rho) \cdot (\dim(\mathcal{X}) - 1) \right] - \left[M_c \cdot \phi_{Sys}(\mathcal{X}, \rho^-) + M_p(\rho^-) \cdot (\dim(\mathcal{X}) - 1) \right] \leq 0 \\ & \left[M_c \cdot (\phi_{Sys}(\mathcal{X}, \rho) - \phi_{Sys}(\mathcal{X}, \rho^-)) \right] + (\dim(\mathcal{X}) - 1) \cdot (M_p^v)^2 \cdot \left[(\rho)^2 - (\rho^-)^2 \right] \leq 0 \\ & \frac{(\phi_{Sys}(\mathcal{X}, \rho) - \phi_{Sys}(\mathcal{X}, \rho^-))}{((\rho)^2 - (\rho^-)^2)} \leq -(\dim(\mathcal{X}) - 1) \cdot \frac{(M_p^v)^2}{M_c} \end{aligned} \quad (\text{III. 37})$$

Par ailleurs, les limites aux bornes de $\Gamma_{Sys}(\mathcal{X}, \rho)$ donnent :

$$\begin{aligned} \lim_{\rho \rightarrow 0} \Gamma_{Sys}(\mathcal{X}, \rho) &= \lim_{\rho \rightarrow 0} \left(\underbrace{M_C \cdot \phi_{Sys}(\mathcal{X}, \rho)}_{\rightarrow \text{constante : } M_C \cdot \phi^0(\mathcal{X})} + \underbrace{\left(M_P^f + (\rho \cdot M_P^v)^2 \right) \cdot (\dim(\mathcal{X}) - 1)}_{\rightarrow \text{constante : } M_P^f \cdot (\dim(\mathcal{X}) - 1)} \right) \\ &= M_C \cdot \phi_{mIN}(\mathcal{X}) + M_P^f \cdot (\dim(\mathcal{X}) - 1) \end{aligned} \quad (III. 38)$$

$$\begin{aligned} \lim_{\rho \rightarrow 1} \Gamma_{Sys}(\mathcal{X}, \rho) &= \lim_{\rho \rightarrow 1} \left(\underbrace{M_C \cdot \phi_{Sys}(\mathcal{X})}_{\rightarrow \text{constante : } M_C \cdot \sum_{i=1}^{\dim(\mathcal{X})} \int_0^{\delta_{x_i}} \lambda_{x_i}(t) dt} + \underbrace{\left(M_P^f + (\rho \cdot M_P^v)^2 \right) \cdot (\dim(\mathcal{X}) - 1)}_{\rightarrow \text{constante : } \left(M_P^f + (M_P^v)^2 \right) \cdot (\dim(\mathcal{X}) - 1)} \right) \\ &= M_C \cdot \sum_{i=1}^{\dim(\mathcal{X})} \left(\int_0^{\delta_{x_i}} \lambda_{x_i}(t) dt \right) + \left(M_P^f + (M_P^v)^2 \right) \cdot (\dim(\mathcal{X}) - 1) \end{aligned} \quad (III. 39)$$

Il existe donc un facteur d'amélioration ρ^* qui satisfait les relations suivantes :

$$\rho^* = \begin{cases} \Gamma_{Sys}(\mathcal{X}, \rho^+) - \Gamma_{Sys}(\mathcal{X}, \rho) \geq 0 \\ \Gamma_{Sys}(\mathcal{X}, \rho) - \Gamma_{Sys}(\mathcal{X}, \rho^-) \leq 0 \\ \lim_{\rho \rightarrow 0} \Gamma_{Sys}(\mathcal{X}, \rho) = \text{constante} \\ \lim_{\rho \rightarrow 1} \Gamma_{Sys}(\mathcal{X}, \rho) = \text{constante} \end{cases} \quad (III. 40)$$

Dans les équations (III. 35) à (III. 37), les facteurs d'amélioration notés ρ^- et ρ^+ représentent respectivement les valeurs antécédente et subséquente à la valeur de ρ . La valeur du facteur d'amélioration étant un réel appartenant à l'intervalle $]0, 1[$, les facteurs ρ^- et ρ^+ sont obtenus respectivement par $(\rho - \rho^{inc})$ et $(\rho + \rho^{inc})$ où ρ^{inc} définit la précision souhaitée.

La procédure numérique utilisée pour déterminer le facteur ρ^* est présentée par la figure suivante :

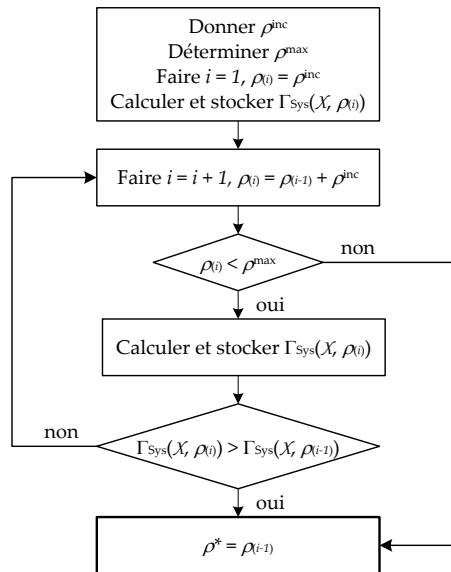


Figure 25 : Procédure numérique pour déterminer ρ^*

La première étape de cette procédure fait référence à un paramètre $\rho^{\max}$; il s'agit de la borne supérieure que peut prendre le facteur d'amélioration afin que cette politique de maintenance systématique soit plus efficace que la politique minimaliste. La valeur de ce paramètre est exprimée par :

$$\underbrace{\left(\frac{\Gamma_{Min}(\mathcal{X})}{\dim(\mathcal{X})-1} \right)}_{\text{coût maximum admissible pour une MP}} > M_p^f + (\rho^{\max} \cdot M_p^v)^2 \quad (\text{III. 41})$$

$$\rho^{\max} = \min \left(1, \frac{1}{M_p^v} \cdot \sqrt{\frac{\Gamma_{Min}(\mathcal{X})}{\dim(\mathcal{X})-1} - M_p^f} \right)$$

L'intérêt du paramètre $\rho^{\max}$ consiste simplement à minimiser l'horizon de recherche lors de la procédure numérique. Nous pouvons noter que la valeur réelle du facteur d'amélioration sera inférieure à la valeur donnée par l'équation (III. 41). En effet, les coûts liés aux actions correctives ont été négligés car ils font eux-mêmes référence à la valeur du facteur d'amélioration.

Lorsque la durée des actions préventives dépend du facteur d'amélioration, la contrainte temporelle (III. 25) se retrouve modifiée comme suit :

$$\sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} + \mu_p(\rho) \cdot (\dim(\mathcal{X}) - 1) \leq H \quad (\text{III. 42})$$

L'évolution de cette contrainte impacte directement le choix du plan d'exploitation. En effet, lorsque le facteur d'amélioration était fixé, la durée consacrée aux actions préventives était connue et était prise en compte pour le choix du plan d'exploitation. Quand la durée des actions préventives varie en fonction du facteur d'amélioration ρ , la durée totale consacrée aux actions préventives n'est plus déterministe, mais se situe dans un intervalle de temps défini comme suit :

$$(\dim(\mathcal{X}) - 1) \cdot \left[\underbrace{\mu_p^f}_{\text{MP considérées comme minimales}}, \underbrace{\mu_p^f + (\mu_p^v)^2}_{\text{MP considérées comme parfaites}} \right] \quad (\text{III. 43})$$

Suivant la valeur prise par le facteur d'amélioration, la différence maximale concernant la durée nécessaire pour l'accomplissement des maintenances préventives s'élève à $(\dim(\mathcal{X}) - 1) \cdot (\mu_p^v)^2$. Cette variabilité de la durée totale des activités de maintenances préventives engendre un dilemme quant au choix du plan d'exploitation. Deux tactiques peuvent être suivies :

- Soit le plan d'exploitation doit disposer d'une durée suffisante pour réaliser les actions de maintenances préventives considérées comme parfaites (ρ tend vers 1). Cependant, pendant la durée restante $\left(H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} - (\dim(\mathcal{X}) - 1) \cdot \left[\mu_p^f + (\rho^* \cdot \mu_p^v)^2 \right] \right)$, d'autres missions auraient pu être réalisées. La durée restante représente la différence entre la durée de l'horizon de temps H minorée de la durée des missions et des maintenances préventives réalisées avec le facteur optimal ρ^* .

- Soit le plan d'exploitation dispose d'une durée, consacrée aux actions préventives, qui appartient à l'intervalle défini par l'équation (III. 43). Dans ce cas, il existe une valeur maximale admissible pour le facteur d'amélioration.

Alors que pour la première tactique, la valeur maximale du facteur d'amélioration, notée $\rho^{\max}$, est égale à 1, il convient de déterminer, pour la seconde tactique, la valeur maximale admissible pour ce facteur d'amélioration en fonction de la durée disponible. La valeur $\rho^{\max}$ est donnée par :

$$H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} = (\dim(\mathcal{X}) - 1) \cdot \left(\mu_p^f + (\rho^{\max} \cdot \mu_p^v)^2 \right)$$

$$\rho^{\max} = \min \left(1, \frac{1}{\mu_p^v} \cdot \sqrt[2]{\frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\dim(\mathcal{X}) - 1} - \mu_p^f} \right) \quad (\text{III. 44})$$

Lorsque les durées des maintenances préventives dépendent du facteur d'amélioration, la valeur optimale ρ^* est obtenue en résolvant l'équation (III. 34). La procédure numérique définie par la figure 25 est donc utilisée pour déterminer ce facteur ; la seule différence se situant au niveau de la valeur de $\rho^{\max}$ qui est précisée par l'équation (III. 44).

Lorsque le facteur d'amélioration intervient à la fois au niveau du coût et de la durée des actions préventives, la valeur maximale admissible pour $\rho^{\max}$ s'établit à partir des équations (III. 41) et (III. 44) :

$$\rho^{\max} = \min \left(1, \frac{1}{M_p^v} \cdot \sqrt[2]{\frac{\Gamma_{Min}(\mathcal{X})}{\dim(\mathcal{X}) - 1} - M_p^f}, \frac{1}{\mu_p^v} \cdot \sqrt[2]{\frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\dim(\mathcal{X}) - 1} - \mu_p^f} \right) \quad (\text{III. 45})$$

L'application de cette politique de maintenance ne pose donc aucun problème pour la stratégie séquentielle ou intégrée. Pour chaque plan d'exploitation défini par l'algorithme génétique, il est possible de déterminer les coûts de maintenance et ce, en déterminant la valeur optimale du facteur d'amélioration.

Pour cette politique systématique, une première amélioration a pu être apportée en recherchant la valeur optimale prise par le facteur d'amélioration pour minimiser les coûts de maintenance. Néanmoins, il pourrait être intéressant de déterminer la valeur optimale prise par le facteur d'amélioration pour chaque action préventive de sorte à minimiser les coûts de maintenance.

IV.2. Recherche des facteurs d'amélioration optimaux

Pour cette nouvelle amélioration apportée à la politique systématique, la différence se situe au niveau du nombre de variables de décisions à déterminer. Ce dernier est identique au nombre de maintenances préventives, à savoir $(\dim(\mathcal{K}) - 1)$. Avec cette nouvelle amélioration, la structure des coûts et des durées reste inchangée. Par conséquent, le coût total de maintenance s'exprime sous la forme :

$$\Gamma_{Sys}(X, P) = M_C \cdot \phi_{Sys}(X, P) + \sum_{j=1}^{\dim(P)} M_P(\rho_j) \quad (\text{III. 46})$$

Avec :

- P : vecteur constitué de $(\dim(K)-1)$ composantes. La $i^{\text{ème}}$ composante représente le facteur d'amélioration associé à la $i^{\text{ème}}$ maintenance préventive,
- $\Gamma_{Sys}(X, P)$: coût de maintenance associé à la politique systématique, pour un plan d'exploitation défini par le vecteur X et le vecteur des facteurs d'amélioration P ,
- $\phi_{Sys}(X, P)$: nombre moyen de pannes pour la politique systématique, pour un plan d'exploitation défini par le vecteur X et le vecteur des facteurs d'amélioration P .

L'expression du nombre moyen de pannes reste pratiquement inchangée, la seule différence se situant au niveau des bornes de l'intégrale :

$$\phi_{Sys}(X, P) = \sum_{i=1}^{\dim(X)} \int_{\tau_{x_{i-1}}(\rho_{i-1})}^{\tau_{x_i}(\rho_i) + \delta_{x_i}} \lambda_{x_i}(t) dt \quad (\text{III. 47})$$

Où $\tau_{x_i}(\rho_i)$ représente l'âge fonctionnel en fonction du facteur d'amélioration ρ_i et dont l'expression est donnée ci-dessous :

$$\tau_{x_i}(\rho_i) = \begin{cases} 0 & \text{si } i = 0 \\ (1 - \rho_i) \cdot R^{-1} \left(\left[R(t_{x_i}) \right]^{\frac{g(Z_{x_i})}{g(Z_{x_{i+1}})}} \right) & \text{sinon} \end{cases} \quad (\text{III. 48})$$

Dans l'approche proposée précédemment, le facteur d'amélioration était commun à l'ensemble des activités de maintenances préventives. Dans ce cas, il était alors possible de déterminer une valeur maximale admissible pour ce facteur en tenant compte de la durée allouée aux actions de maintenance et du coût de la politique minimaliste. Lorsque le facteur d'amélioration peut être différent pour représenter l'effet de chaque maintenance préventive, il n'est pas possible de déterminer une valeur maximale. Cependant, pour que cette politique avec les facteurs d'amélioration optimaux soit économiquement plus intéressante que la politique minimaliste, la contrainte ci-dessous doit être vérifiée :

$$\underbrace{\dim(P) \cdot \left(\frac{\Gamma_{Min}(X)}{\dim(X) - 1} \right)}_{\text{coût maximum admissible pour une MP}} = \sum_{j=1}^{\dim(P)} \left(M_P^f + (\rho_j \cdot M_P^v)^2 \right) \quad (\text{III. 49})$$

$$\sum_{j=1}^{\dim(P)} (\rho_j^2) = \min \left(\dim(P), \frac{\dim(P)}{(M_P^v)^2} \cdot \left(\frac{\Gamma_{Min}(X)}{\dim(X) - 1} - M_P^f \right) \right)$$

Lorsque les durées des maintenances préventives dépendent également des facteurs d'amélioration, l'effet se répercute au niveau de l'expression de la contrainte temporelle, défini comme suit :

$$\sum_{i=1}^{\dim(\mathcal{X})} \delta_{k_i} + \sum_{j=1}^{\dim(\mathcal{P})} \mu_p(\rho_j) \leq H \quad (\text{III. 50})$$

De manière similaire à la recherche du facteur d'amélioration optimal commun à l'ensemble des maintenances (Paragraphe IV. 1.), la durée allouée aux maintenances préventives est variable car elle dépend des valeurs des facteurs d'amélioration. Nous sommes donc confrontés aux mêmes dilemmes, à savoir :

- Soit le plan d'exploitation conserve une durée permettant la réalisation des maintenances préventive considérées comme parfaites (ρ_j tend vers 1, avec $j = \{1, \dots, \dim(\mathcal{P})\}$). Dans ce cas, moins de missions sont réalisées, mais les facteurs d'améliorations sont optimaux.
- Soit la durée consacrée aux actions de maintenances est comprise dans l'intervalle défini par l'équation (III. 43), et dans ce cas, il convient de respecter la contrainte temporelle exprimée ci-dessous :

$$\underbrace{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{X_i}}_{\text{durée admissible pour les M.P.}} = \sum_{j=1}^{\dim(\mathcal{P})} \left(\mu_p^f + (\rho_j \cdot \mu_p^v)^2 \right) \quad (\text{III. 51})$$

$$\sum_{j=1}^{\dim(\mathcal{P})} (\rho_j^2) = \min \left(\dim(\mathcal{P}), \frac{1}{(\mu_p^v)^2} \cdot \left(H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{X_i} - \dim(\mathcal{P}) \cdot \mu_p^f \right) \right)$$

La recherche des facteurs d'amélioration optimaux (variables de décisions) nécessiterait la résolution des $\dim(\mathcal{P})$ équations ci-dessous :

$$\frac{\partial \Gamma_{sys}(\mathcal{X}, \mathcal{P})}{\partial \rho_i} = 0 \quad \forall i = \{1, 2, \dots, \dim(\mathcal{P})\} \quad (\text{III. 52})$$

Compte tenu de la complexité de l'équation du coût de maintenance (III. 46), il est très difficile d'obtenir les valeurs optimales des facteurs d'amélioration. De plus, lorsque les coûts et les durées de maintenances dépendent des facteurs d'amélioration, les contraintes exprimées par les équations (III. 49) et (III. 51) doivent être vérifiées et l'existence de solution n'est pas garantie.

Pour le cas précédent, le facteur optimal pouvait être déterminé facilement grâce à une procédure numérique, car il s'agissait de la seule variable de décision. Cependant, lorsque le nombre de variables de décision augmente, le nombre de possibilités à tester devient trop important. À titre illustratif, en considérant un plan d'exploitation constitué de 10 missions et en supposant que la précision des facteurs d'amélioration soit de 10^{-1} , il existe plus de 380 millions de cas à considérer $(10-1)^{(1/0,1)-1}$. Aussi, pour remédier à ce problème et obtenir des résultats satisfaisants en un temps de calcul raisonnable, une méta-heuristique peut être utilisée.

IV.2.a. Méta-heuristique pour la recherche des facteurs d'amélioration optimaux

Pour l'obtention de solutions satisfaisantes, nous avons à nouveau opté pour l'utilisation de l'algorithme génétique. Le processus de fonctionnement ayant déjà été expliqué lors de l'étude de la

stratégie intégrée, nous nous intéresserons directement au codage ainsi qu'aux opérateurs génétiques employés.

IV.2.a.i Le codage

Pour cet algorithme génétique, chaque individu doit représenter les facteurs d'amélioration des maintenances préventives. Par conséquent, chaque gène sera associé à une maintenance préventive et sera représenté par un entier compris dans l'intervalle $]0, 1[$. Par ailleurs, pour la valeur de chaque gène, nous nous limiterons à une précision de l'ordre de 10^{-1} permettant de réaliser une maintenance préventive suivant neuf niveaux (0,1 ; 0,2, ..., 0,9). Au-delà de cette précision, le modèle s'éloigne de la réalité car il s'avère difficile d'effectuer des maintenances conformes à la précision souhaitée.

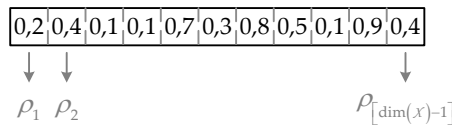


Figure 26 : Codage du chromosome pour l'A.G. associé à la détermination des facteurs optimaux

IV.2.a.ii La stratégie de reproduction

La stratégie de reproduction utilisée pour cet algorithme génétique est semblable à celle développée pour l'optimisation des plans d'exploitation. À partir des individus parents, des tournois déterministes sont réalisés afin de déterminer les individus qui, soumis aux opérateurs génétiques, engendreront les individus de la génération suivante. Une sélection élitiste sera également opérée pour garantir la sauvegarde des meilleurs individus. Enfin, quelques nouveaux individus, générés aléatoirement, seront intégrés à la nouvelle population pour maintenir une certaine diversité et éviter les optimaux locaux. Pour ces individus générés aléatoirement (ainsi que pour la population initiale), il convient de s'assurer de leur viabilité. Un individu est dit viable, si la contrainte temporelle définie par l'équation (III. 50) est vérifiée.

IV.2.a.iii L'opérateur de croisement

Pour cet algorithme génétique, le croisement en un point a été choisi. Le choix du point de croisement est réalisé aléatoirement et ne pose aucune difficulté car les chromosomes sont de tailles identiques. À partir de ce point de croisement, les gènes sont inversés entre les deux individus.

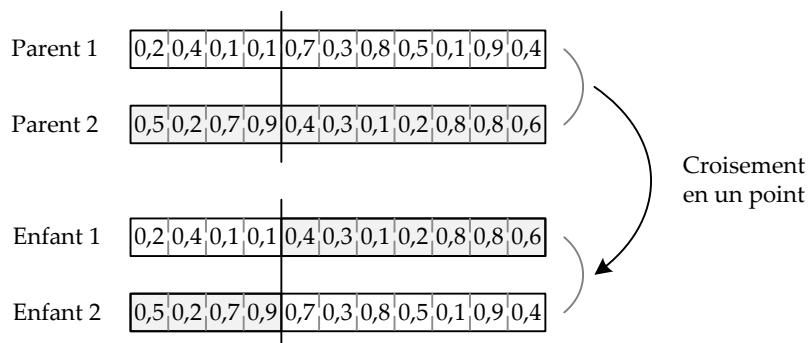


Figure 27 : Opération de croisement effectuée entre les deux individus parents

Suivant le choix effectué pour la durée réservée aux actions préventives, les individus générés peuvent ne pas être viables. En effet, si la durée laissée disponible lors de la génération du plan d'exploitation est inférieure à la durée nécessaire pour la réalisation des maintenances préventives, ces individus devront être corrigés. La correction de ces individus s'effectue en diminuant la valeur des facteurs d'amélioration, en choisissant les gènes à traiter de manière aléatoire, et ce, tant que la contrainte est violée. Néanmoins, la valeur de chaque gène devra respecter les règles imposées par le codage.

IV.2.a.iv L'opérateur de mutation

Pour éviter la non-viabilité des chromosomes, l'opérateur de mutation consiste à inverser la valeur de deux gènes choisis aléatoirement. De ce fait, l'ensemble des facteurs optimaux reste identique, la seule différence étant le placement de ces derniers. Bien évidemment, les deux facteurs d'amélioration à inverser devront être de valeurs différentes. Dans le cas contraire, l'opération génétique n'aura été d'aucun intérêt.

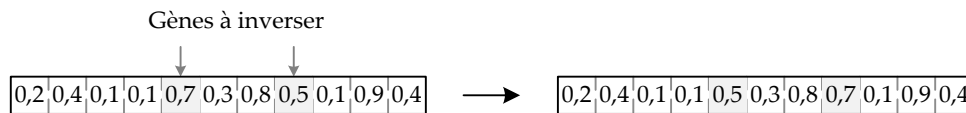


Figure 28 : Modélisation de l'opération de mutation

Cette politique de maintenance systématique est caractérisée par la réalisation de maintenances préventives après chaque mission. Afin que les coûts de maintenances soient inférieurs à la politique minimaliste, deux améliorations ont été proposées. La première consiste à déterminer le facteur d'amélioration optimal ρ^* pour l'effet des maintenances préventives, en ayant pris soin de modéliser les coûts et durées des actions préventives, en fonction de ce facteur. La deuxième amélioration est directement liée à la première : elle consiste à déterminer le facteur optimal propre à chaque action préventive. Cependant, la question qui peut être soulevée pour cette politique concerne le nombre de maintenances préventives. Est-il réellement judicieux de réaliser une maintenance préventive après chaque mission ?

V. La politique sporadique

Similairement à la politique précédente, les maintenances préventives sont toujours réalisées à la fin des missions. Cependant, cette nouvelle politique se différencie par le fait que chaque mission ne sera pas systématiquement suivie d'une action préventive. L'objectif consiste à déterminer les missions qui doivent être suivies d'actions préventives. De manière évidente, le nombre de maintenances préventives à réaliser sera compris entre 0 (Politique minimaliste) et $(\dim(\mathcal{K}) - 1)$ (Politique systématique). Le nombre d'intervalles de maintenances préventives sera noté par N , et par conséquent $(N - 1)$ représentera le nombre de maintenances préventives. Le coût total de maintenance, pour cette politique sporadique, s'établit comme suit :

$$\Gamma_{Spo}(\mathcal{X}, P) = M_C \cdot \phi_{Spo}(\mathcal{K}, P) + M_P \cdot \underbrace{\sum_{i=1}^{\dim(P)} (1 - \lfloor 1 - \rho_i \rfloor)}_{(N-1)} \quad (\text{III. 53})$$

Avec :

- $\Gamma_{Spo}(\mathcal{X}, P)$: coût de maintenance associé à la politique sporadique, pour le plan d'exploitation défini par le vecteur $\mathcal{X}$ et dont les maintenances préventives sont définies par le vecteur P ,
- $\phi_{Spo}(\mathcal{K}, P)$: nombre moyen de pannes, pour la politique sporadique, pouvant survenir pendant le plan d'exploitation défini par le vecteur $\mathcal{X}$, par rapport au vecteur P définissant les maintenances préventives,
- P : vecteur constitué de $(\dim(\mathcal{X}) - 1)$ composantes où la $i^{\text{ème}}$ composante représente l'état de la maintenance,
- ρ_i : $i^{\text{ème}}$ composante du vecteur P . Si $\rho_i = 0$, la $i^{\text{ème}}$ mission n'est pas suivie d'une maintenance préventive. Par contre, si $\rho_i = \rho$, une maintenance préventive est réalisée,
- ρ : facteur d'amélioration constant pour les maintenances préventives,
- $\lfloor x \rfloor$: partie entière de x .

Le nombre moyen de pannes, pour cette politique sporadique, est donné par :

$$\phi_{Spo}(\mathcal{X}, P) = \sum_{i=1}^{\dim(\mathcal{X})} \int_{\tau_{x_{i-1}}(\rho_{i-1})}^{\tau_{x_i}(\rho_i) + \delta_{x_i}} \lambda_{x_i}(t) dt \quad (\text{III. 54})$$

Où $\tau_{x_i}(\rho_i)$ représente l'âge fonctionnel en fonction du facteur d'amélioration ρ_i et dont l'expression est donnée ci-dessous :

$$\tau_{x_i}(\rho_i) = \begin{cases} 0 & \text{si } i = 0 \\ (1 - \rho_i) \cdot R^{-1} \left(\left[R(t_{x_i}) \right]^{\frac{g(Z_{x_i})}{g(Z_{x_{i+1}})}} \right) & \text{sinon} \end{cases} \quad (\text{III. 55})$$

En raison des maintenances préventives réalisées, $t_{x_i} = \tau_{x_{i-1}} + \delta_{x_i}$.

Lors de la détermination du plan d'exploitation, le nombre de maintenances préventives est inconnu. Par conséquent, la durée nécessaire pour ces actions préventives l'est également. Ainsi, pour un plan d'exploitation $\mathcal{X}$, la durée nécessaire pour accomplir les actions préventives se situe dans l'intervalle $[0, \mu_P \cdot (\dim(\mathcal{X}) - 1)]$. Si la durée disponible pour la réalisation des maintenances préventives est inférieure à la borne supérieure de l'intervalle, il existe un nombre maximal d'activités préventives réalisables, exprimé par la relation suivante :

$$\sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} + \mu_p \cdot \sum_{i=1}^{\dim(\mathcal{P})} (1 - \lfloor 1 - \rho_i \rfloor) \leq H$$

$$\underbrace{\left(\sum_{i=1}^{\dim(\mathcal{P})} (1 - \lfloor 1 - \rho_i \rfloor) \right)^{\max}}_{(N-1)^{\max}} = \left\lfloor \frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\mu_p} \right\rfloor \quad (\text{III. 56})$$

De manière similaire, il est possible de déterminer le nombre maximal d'actions préventives pour que cette politique sporadique soit plus avantageuse que les politiques précédentes.

Connaissant le nombre maximal de maintenances préventives réalisables, ce problème peut être résolu de façon optimale à l'aide de la procédure numérique suivante. Le nombre de possibilités reste assez faible ; il s'élève à $\sum_{i=1}^{(N-1)^{\max}} C_{[\dim(\mathcal{X})-1]}^i$.

Par analogie avec la politique systématique, une première amélioration serait de déterminer le facteur optimal de maintenance qui interagirait sur les coûts et/ou la durée.

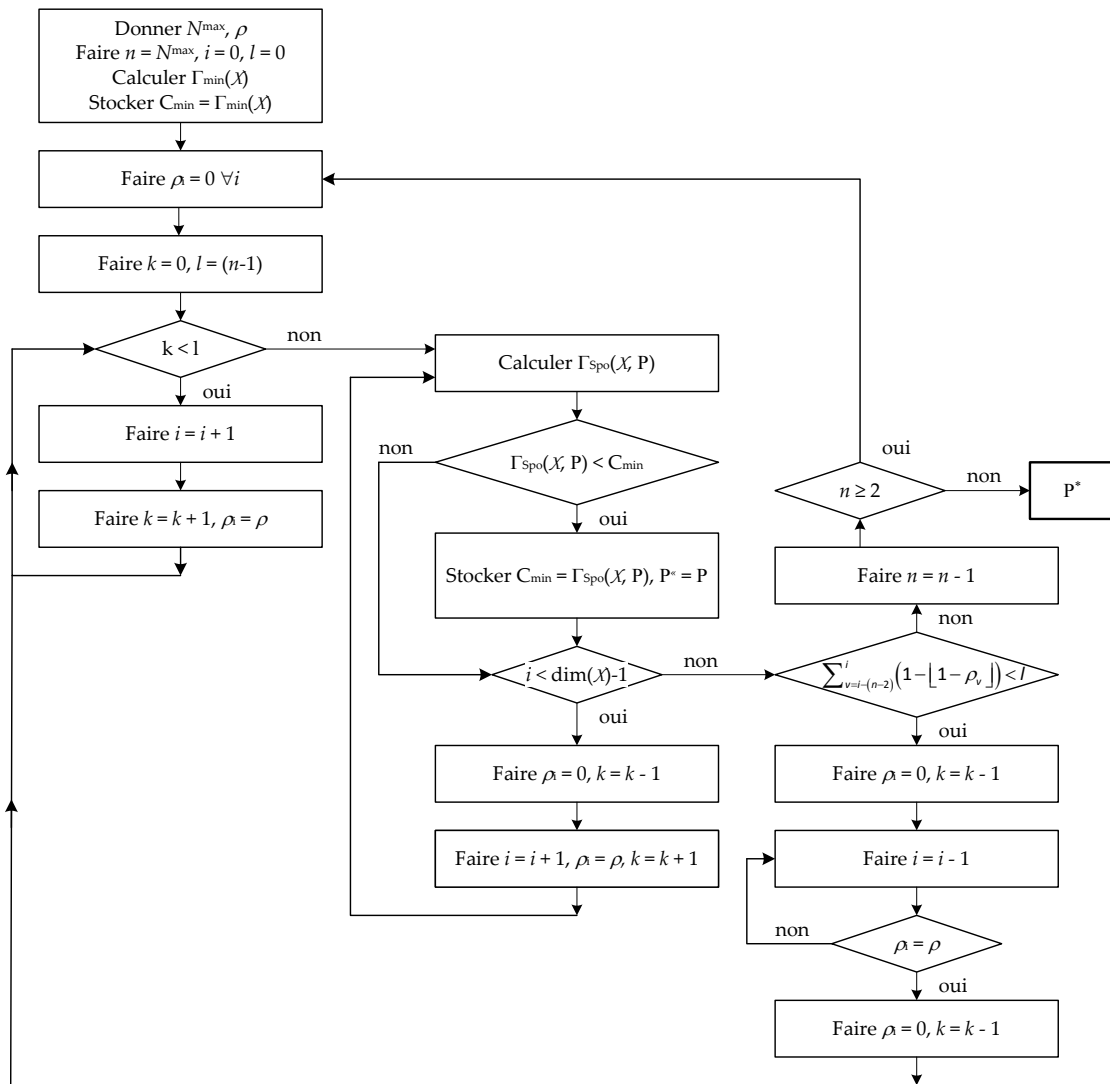


Figure 29 : Procédure numérique pour le choix des missions suivies de maintenances préventives

V.1. Recherche du facteur d'amélioration optimal

Lorsque la variation du facteur d'amélioration interfère sur le coût des actions préventives, le coût total de maintenance s'établit sous la forme suivante :

$$\Gamma_{S_{po}}(\mathcal{X}, P) = M_C \cdot \phi_{S_{po}}(\mathcal{K}, P) + \underbrace{\sum_{i=1}^{\dim(P)} M_P(\rho_i) \cdot (1 - \lfloor 1 - \rho_i \rfloor)}_{(N-1) = \sum_{i=1}^{\dim(P)} (1 - \lfloor 1 - \rho_i \rfloor)} \quad (\text{III. 57})$$

Où $\rho_i = 0$, si la $i^{\text{ème}}$ mission n'est pas suivie d'une maintenance préventive. Dans le cas contraire, la valeur de ρ_i appartient alors à l'intervalle $]0, 1[$. Les expressions du nombre moyen de pannes et de l'âge fonctionnel restent inchangées par rapport aux équations (III. 54) et (III. 55). En plus de déterminer les missions qui seront suivies d'une action préventive, il convient de trouver la valeur optimale du facteur d'amélioration. En considérant la même précision qu'utilisée précédemment pour le choix du facteur optimal, la résolution peut s'établir à partir de la procédure numérique ci-dessous :

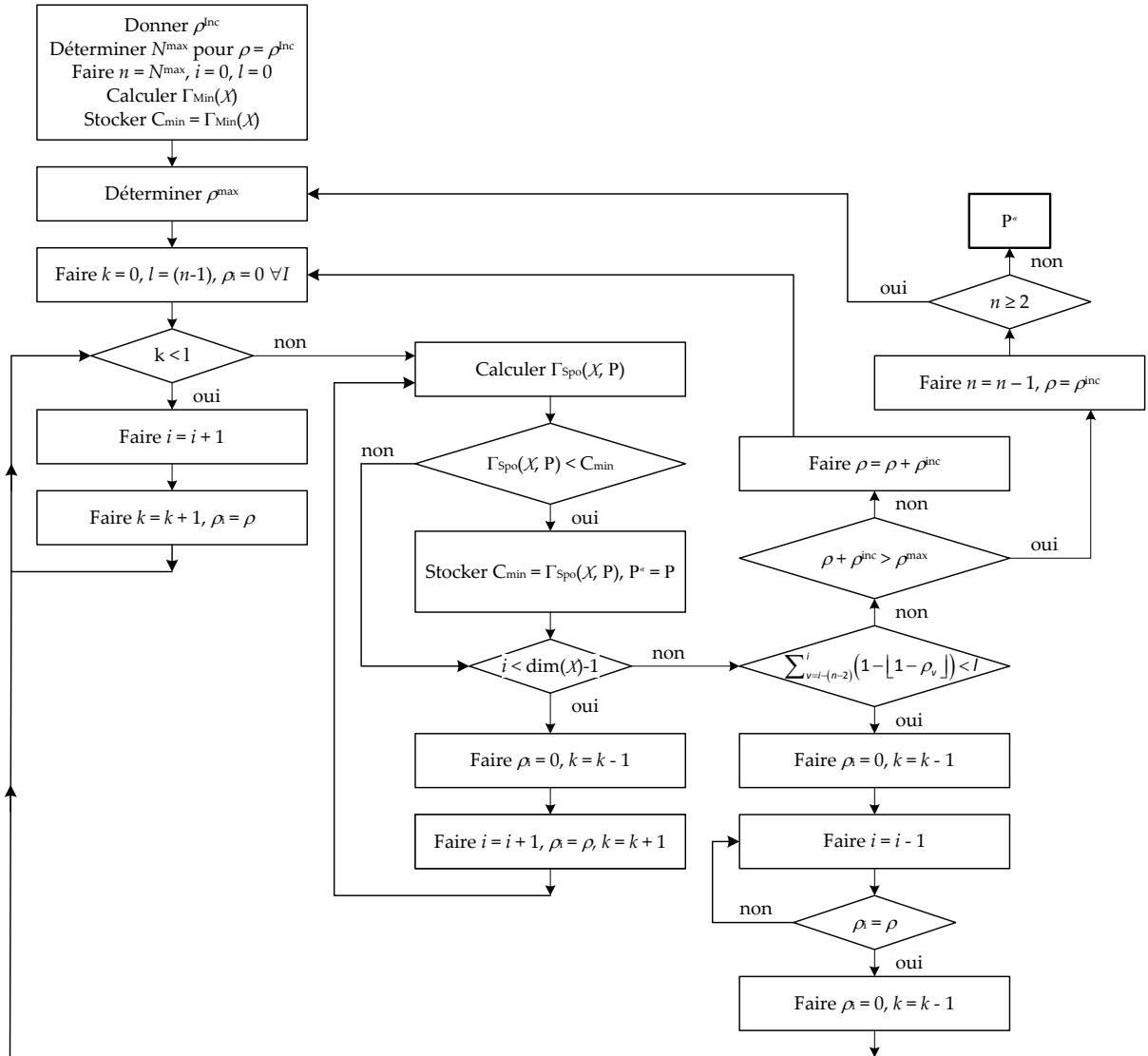


Figure 30 : Procédure numérique pour le choix du facteur optimal et des missions suivies de maintenances préventives

Dans cette procédure numérique, illustrée par la figure 30, nous notons la présence des paramètres $N^{\max}$ et $\rho^{\max}$, utilisés pour minimiser l'espace de recherche. En effet, lorsque la durée allouée pour la réalisation des maintenances préventives ne permet pas de réaliser une maintenance préventive quasi-parfaite après chaque mission (ρ tend vers 1), il est nécessaire de déterminer la valeur maximale admissible pour le facteur d'amélioration optimal.

Lorsque la durée de maintenance préventive est dépendante du facteur d'amélioration, la contrainte temporelle devient :

$$\sum_{i=1}^{\dim(X)} \delta_{x_i} + \sum_{j=1}^{\dim(P)} \left(\mu_p(\rho_j) \right) \cdot \left(1 - \lfloor 1 - \rho_j \rfloor \right) \leq H \quad (\text{III. 58})$$

Le nombre maximal de maintenances préventives est compris dans l'intervalle $]0, (\dim(X) - 1[$. Les bornes de cet intervalle sont exclues, pour éviter de retrouver les politiques minimaliste et systématique. À partir de cette équation (III. 58), la valeur maximale du facteur d'amélioration, en fonction du nombre d'actions préventives, est donnée par :

$$\sum_{i=1}^{\dim(X)} \delta_{x_i} + (N-1) \cdot \left[\mu_p^f + \left(\rho^{\max} \cdot \mu_p^v \right)^2 \right] = H$$

$$\rho^{\max} = \min \left(1, \frac{1}{\mu_p^v} \cdot \sqrt[2]{\frac{H - \sum_{i=1}^{\dim(X)} \delta_{x_i}}{(N-1)}} - \mu_p^f \right) \quad (\text{III. 59})$$

Lorsque le coût d'une action préventive est dépendant du facteur d'amélioration, il est possible de déterminer une valeur maximale, pour ce dernier, afin que cette politique soit économiquement plus intéressante que la politique minimaliste. En ne se préoccupant que des coûts liés aux actions préventives, le facteur $\rho^{\max}$ se définit par :

$$\underbrace{\left(\frac{\Gamma_{\text{Min}}(X)}{(N-1)} \right)}_{\text{coût maximum admissible pour une MP}} > M_p^f + \left(\rho^{\max} \cdot M_p^v \right)^2 \quad (\text{III. 60})$$

$$\rho^{\max} = \min \left(1, \frac{1}{M_p^v} \cdot \sqrt[2]{\frac{\Gamma_{\text{Min}}(X)}{(N-1)}} - M_p^f \right)$$

Pour cette politique sporadique, une amélioration a été apportée en déterminant le facteur d'amélioration optimal commun à toutes les maintenances préventives. Dans la partie suivante, nous chercherons à déterminer les valeurs optimales des facteurs d'amélioration pour chaque maintenance préventive.

V.2. Recherche des facteurs d'amélioration optimaux

Comme pour la politique systématique, une amélioration peut être apportée à cette politique de maintenance par la recherche des facteurs d'amélioration optimaux associés à chaque action

préventive. Avec cette modification, il est très difficile de montrer la différence entre les deux expressions mathématiques. En effet, ni l'expression du coût de maintenance pour cette politique (III. 57), ni celles du nombre moyen de pannes (III. 54) et de l'âge fonctionnel (III. 55) ne changent. La différence se situe au niveau du nombre de variables de décision et donc, des valeurs des composantes ρ . Précédemment, lorsque les maintenances préventives étaient réalisées ($\rho \neq 0$), les facteurs d'amélioration avaient tous la même valeur. Dans le cas présent, lorsque les facteurs d'amélioration sont non nuls, ils peuvent avoir des valeurs différentes. Cette politique se caractérisant par le choix de réaliser ou non une action préventive, et de choisir la valeur optimale pour facteur action préventive, la résolution de ce problème engendre une explosion combinatoire. Pour la politique systématique, lorsque 10 missions étaient réalisées, il existait plus de 380 millions de possibilités. Avec cette nouvelle politique, et en prenant compte des mêmes caractéristiques, le nombre de cas s'élève à environ 1 milliard $\left(\sum_{i=1}^{(N-1)^{\max}} C_{[\dim(\chi)-1]}^i \cdot (1/0,1-1)^i \right)$. Bien évidemment, la résolution de ce problème s'effectuera à l'aide d'une méta-heuristique.

V.2.a. Méta-heuristique pour la recherche des facteurs d'amélioration optimaux

Contrairement à la politique systématique, un algorithme génétique ne sera pas utilisé. En effet, l'opérateur génétique de croisement risquerait de générer des individus dont le nombre de maintenances préventives resterait plus ou moins constant et ne permettrait donc pas de parcourir efficacement l'horizon de recherche. Pour remédier à ce problème, l'algorithme du grand déluge a été choisi [Schutz, 2009A], et ce, notamment pour sa facilité d'implémentation.

En 2004, Burke *et al.* [Burke, 2004] ont proposé l'algorithme du grand déluge étendu. La particularité de cette procédure de recherche locale est d'accepter des solutions moins bonnes que la solution courante. À partir d'une solution initiale, qui définit la valeur de seuil L , une nouvelle solution est choisie à partir du voisinage de cette solution. Lorsqu'une nouvelle solution est acceptée, la valeur du seuil L est réduite d'une quantité ΔL , préalablement fixée. Par la suite, une solution candidate est acceptée si son fitness est meilleur que celui de la meilleure solution actuelle. Cependant, avec cette méta-heuristique, une solution peut être acceptée même si son fitness n'est pas optimal, mais à condition d'être meilleur que la valeur du seuil L . Le critère d'arrêt de cet algorithme peut être défini par un nombre total d'itérations, ou après un certain nombre d'itérations infructueuses, ou dès lors que la valeur du seuil L devient inférieure à la meilleure solution.

V.2.a.i Le codage

Pour cet algorithme, une solution représente l'ensemble des facteurs d'amélioration. Par conséquent, une solution est représentée par un segment composé de $\dim(P)$ gènes. Chaque gène caractérisera la valeur prise par le facteur d'amélioration. Notons que si la valeur du $i^{\text{ème}}$ gène est égale à 0, dans ce cas, la $i^{\text{ème}}$ mission ne sera pas suivie d'une action préventive. Ce codage est identique à celui employé au cours du paragraphe (IV.2.a.i), la seule différence étant l'utilisation de la valeur 0 pour représenter la non réalisation d'une maintenance préventive.

V.2.a.ii Définition du voisinage

Dans l'algorithme génétique, les nouveaux individus étaient obtenus grâce à des opérations génétiques. Dans l'algorithme du grand déluge étendu, des solutions potentielles (le voisinage) sont définies à partir de la solution courante. Ces solutions sont générées à l'aide de simples mutations. Pour notre problème, nous avons choisi deux types de mutations, de fréquence d'exécution équiprobable :

- Modification de la valeur d'un gène non nul,
- Modification du nombre de maintenances préventives.

La première opération ne pose aucune difficulté particulière. Il convient, cependant, de générer une valeur différente pour le facteur d'amélioration concerné et de s'assurer de sa viabilité par le respect de la contrainte temporelle. Cependant, pour la seconde opération, plusieurs cas peuvent se distinguer. Si seule, une maintenance préventive est planifiée, une autre maintenance préventive sera ajoutée en générant la valeur du facteur d'amélioration aléatoirement. Au contraire, si la durée restante (différence entre le temps alloué pour les actions préventives et la durée nécessaire pour la réalisation des maintenances définies par le segment) ne permet pas la réalisation d'une nouvelle maintenance préventive, cette opération consistera à en supprimer une. Pour les autres situations, l'ajout ou la suppression sera effectué aléatoirement.

À travers cette politique sporadique, le but recherché était dans un premier temps de définir les missions suivies de maintenances préventives, puis de définir les facteurs d'amélioration optimaux pour minimiser les coûts relatifs aux activités de maintenance. Nous pouvons noter que les politiques systématique et sporadique se caractérisaient par des maintenances préventives réalisées spécifiquement à la fin de missions. En tenant compte de cette contrainte, il est difficile d'améliorer davantage ces politiques de maintenance préventive. Aussi, afin d'élargir le champ de recherche de ces politiques, la relaxation de cette contrainte permet la réalisation d'actions préventives à tout instant, c'est-à-dire à la fin des missions mais également durant leurs exécutions. À ce titre, une des politiques les plus usitées correspond à la politique périodique.

VI. La politique périodique

Pour mémoire, la politique périodique se caractérise par la réalisation de maintenances préventives à des intervalles de temps constants. Aussi, ces actions préventives peuvent survenir au cours de l'accomplissement de missions. La durée constante de ces intervalles de maintenances préventives, notée Δ , est déterminée à partir de la durée du plan d'exploitation (durée cumulée des missions), ainsi que du nombre d'actions préventives réalisées.

$$\Delta = \frac{\sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{N} \quad (\text{III. 61})$$

Les coûts liés à cette politique de maintenance se définissent par :

$$\Gamma_{\text{Per}}(\mathcal{X}, N) = M_C \cdot \phi_{\text{Per}}(\mathcal{X}, N) + M_P \cdot (N - 1) \quad (\text{III. 62})$$

Avec :

- $\Gamma_{Per}(\mathcal{X}, N)$: coût de maintenance associé à la politique périodique, pour le plan d'exploitation défini par le vecteur $\mathcal{X}$ et suivant le nombre d'intervalles de maintenances préventives N .
- $\phi_{Per}(\mathcal{X}, N)$: Nombre moyen de pannes, fonction du plan d'exploitation $\mathcal{X}$ et du nombre d'intervalles de maintenances préventives.
- N : nombre d'intervalles de maintenances préventives.

L'expression du nombre moyen de pannes est donnée par :

$$\phi_{Per}(\mathcal{X}, N) = \sum_{n=1}^N \underbrace{\left(\sum_{m=d(n)}^{f(n)} \left(\int_{\tau_{x_{m-1}}}^{\tau_{x_m} + \sigma_{x_m}} \lambda_{x_m}(t) dt \right) \right)}_{\text{nombre moyen de pannes pour chaque intervalle de M.P.}} \quad (\text{III. 63})$$

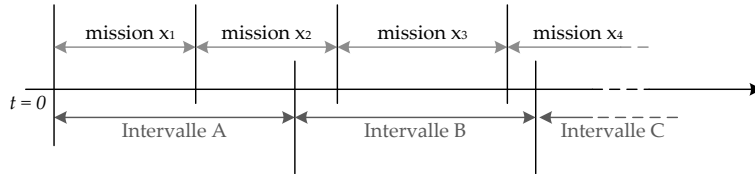


Figure 31 : Exemple de la politique périodique

Comme l'illustre la figure 31, les missions peuvent être liées à plusieurs intervalles de maintenance. À titre illustratif, la mission « x_2 » est partagée entre les intervalles de maintenance préventive « A » et « B ». Parallèlement, les intervalles de maintenance préventive peuvent être constitués de plusieurs missions. À ce titre, l'intervalle « B » est composé d'une partie de la mission « x_2 », de la mission « x_3 » et d'une partie de la mission « x_4 ».

Le calcul du nombre moyen de pannes pouvant survenir durant la réalisation du plan d'exploitation est donc réalisé pour chaque intervalle de maintenance (première sommation), qui est lui-même partagé en plusieurs missions. Par conséquent, il est nécessaire de connaître les missions réalisées durant les N intervalles de maintenances, ainsi que la durée réalisée par les missions appartenant aux intervalles. Ainsi, les termes $d(n)$ et $f(n)$ représentent respectivement, les missions qui débutent et finissent les intervalles n avec $n = \{1, 2, \dots, N\}$. Les missions pouvant intervenir sur plusieurs intervalles de maintenances, la durée de ces missions pour chaque intervalle est tronquée et est notée σ_m avec $\sigma_m \leq \delta_m$. En reprenant l'exemple illustré par la figure 31 et pour l'intervalle « B », les missions « x_2 » et « x_4 » représentent respectivement $d(B)$ et $f(B)$. La durée de la mission « x_2 » dans cet intervalle est donnée par $(\delta_{x_1} + \delta_{x_2}) - \Delta$, alors que la durée de la mission 4 résulte de $\left(2 \cdot \Delta - \left(\sum_{i=1}^3 \delta_{x_i}\right)\right)$ ou $\left(\delta_4 - \left(\sum_{i=1}^4 \delta_{x_i} - 2 \cdot \Delta\right)\right)$.

En raison des maintenances préventives qui peuvent survenir durant les missions, l'expression de l'âge fonctionnel devient :

$$\tau_{x_m} = \begin{cases} 0 & \text{si } m = 0 \\ R^{-1} \left(\left[R(t_{x_m}) \right]^{\frac{g(Z_{x_m})}{g(Z_{x_{m+1}})}} \right) & \text{si } m \neq d(n) \\ (1-\rho) \cdot R^{-1} \left(\left[R(t_{x_m}) \right]^{\frac{g(Z_{x_m})}{g(Z_{x_{m+1}})}} \right) & \text{sinon} \end{cases} \quad (\text{III. 64})$$

Où $t_{x_m} = \tau_{x_{m-1}} + \sigma_{x_m}$

La durée des intervalles de maintenances préventives étant fixe, il importe de déterminer la durée des missions réalisées pendant chaque intervalle. Quatre cas peuvent se présenter, comme précisés ci-dessous :

- Cas 1 : Si la mission x_m ne débute pas l'intervalle de maintenances préventives, ni ne le termine, la mission est donc réalisée complètement ($d(n) \neq m \neq f(n)$).

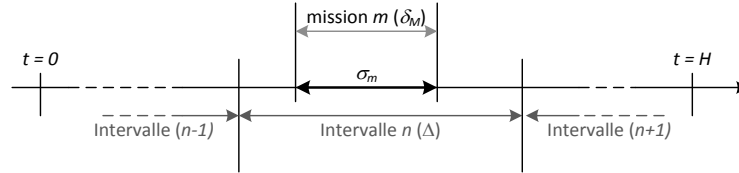


Figure 32 : Illustration du premier cas

La durée de la mission x_m dans l'intervalle n est donnée par :

$$\sigma_{x_m} = \delta_{x_m} \quad (\text{III. 65})$$

- Cas 2 : Si la mission x_m débute l'intervalle de maintenances préventives, mais ne le clôt pas ($d(n) = m \neq f(n)$)

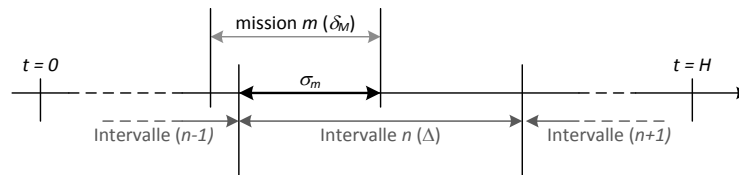


Figure 33 : Illustration du deuxième cas

La durée de la mission x_m dans l'intervalle n s'exprime par :

$$\sigma_{x_m} = \sum_{l=1}^{d(n)} \delta_{x_l} - (n-1) \cdot \Delta \quad (\text{III. 66})$$

- Cas 3 : Si la mission x_m ne débute pas l'intervalle de maintenances préventives, mais le clôt ($d(n) \neq m = f(n)$)

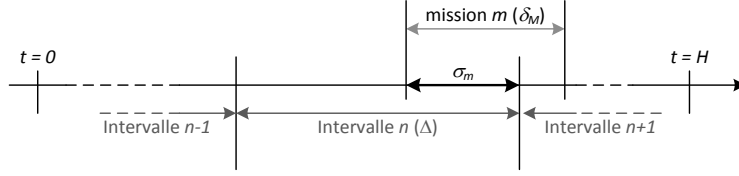


Figure 34 : Illustration du troisième cas

La durée de la mission x_m dans l'intervalle n est donnée par :

$$\sigma_{x_m} = n \cdot \Delta - \sum_{l=1}^{f(n)-1} \delta_{x_l} \text{ ou } \delta_{x_{f(n)}} - \left(\sum_{l=1}^{f(n)} \delta_{x_l} - n \cdot \Delta \right) \quad (\text{III. 67})$$

- Cas 4 : Si la mission x_m débute et clôt l'intervalle de maintenances préventives ($d(n) = m = f(n)$)

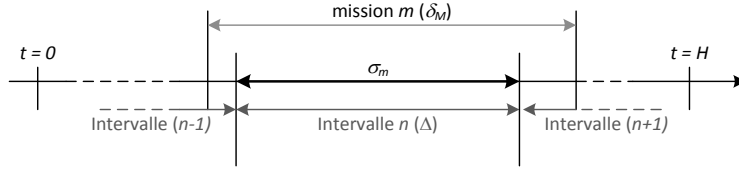


Figure 35 : Illustration du quatrième cas

La durée de la mission x_m dans l'intervalle n s'exprime par :

$$\sigma_{x_m} = \Delta \quad (\text{III. 68})$$

Au final, suivant les cas, la durée d'une mission s'exprime par :

$$\sigma_{x_m} = \begin{cases} \delta_{x_m} & \text{si } d(n) \neq m \neq f(n) \\ \sum_{l=1}^{d(n)} \delta_{x_l} - (n-1) \cdot \Delta & \text{si } d(n) = m \neq f(n) \\ n \cdot \Delta - \sum_{l=1}^{f(n)-1} \delta_{x_l} & \text{si } d(n) \neq m = f(n) \\ \Delta & \text{si } d(n) = m = f(n) \end{cases} \quad (\text{III. 69})$$

Les derniers termes restant à définir dans l'équation (III. 63) correspondent aux missions situées aux extrémités des intervalles de maintenances préventives, à savoir $x_{d(n)}$ et $x_{f(n)}$.

$$d(n) = \begin{cases} 1 & \text{si } n = 1 \\ m \text{ solution de } \begin{cases} \sum_{i=1}^m \delta_{x_i} > (n-1) \cdot \Delta \\ \sum_{i=1}^{m-1} \delta_{x_i} \leq (n-1) \cdot \Delta \end{cases} & \text{sinon} \end{cases} \quad (\text{III. 70})$$

$$f(n) = \begin{cases} \dim(\mathcal{X}) & \text{si } n = N \\ m \text{ solution de } \begin{cases} \sum_{i=1}^m \delta_{x_i} \geq n \cdot \Delta \\ \sum_{i=1}^{m-1} \delta_{x_i} < n \cdot \Delta \end{cases} & \text{sinon} \end{cases} \quad (\text{III. 71})$$

Pour cette politique périodique, la seule variable de décision recherchée correspond au nombre d'intervalles de maintenance préventive. De par la complexité de l'expression des coûts, il est impossible de déterminer analytiquement la solution optimale, qui revient à résoudre l'équation suivante :

$$\frac{\partial \Gamma_{Per}(\mathcal{X}, N)}{\partial N} = 0 \quad (\text{III. 72})$$

Le nombre optimal de maintenances préventives peut être obtenu à l'aide d'une procédure numérique. Le lemme suivant assure l'existence d'un minimum local :

$$\exists N^* \text{ si } \phi_{Per}(\mathcal{X}, N) - \phi_{Per}(\mathcal{X}, N-1) \leq -\frac{M_P}{M_C} \leq \phi_{Per}(\mathcal{X}, N+1) - \phi_{Per}(\mathcal{X}, N) \quad (\text{III. 73})$$

La démonstration de ce lemme est donnée par :

$$\begin{aligned} & \Gamma_{Per}(\mathcal{X}, N+1) - \Gamma_{Per}(\mathcal{X}, N) \geq 0 \\ & [M_C \cdot \phi_{Per}(\mathcal{X}, N+1) + M_P \cdot (N)] - [M_C \cdot \phi_{Per}(\mathcal{X}, N) + M_P \cdot (N-1)] \geq 0 \\ & \phi_{Per}(\mathcal{X}, N+1) - \phi_{Per}(\mathcal{X}, N) \geq -\frac{M_P}{M_C} \end{aligned} \quad (\text{III. 74})$$

$$\begin{aligned} & \Gamma_{Per}(\mathcal{X}, N) - \Gamma_{Per}(\mathcal{X}, N-1) \leq 0 \\ & [M_C \cdot \phi_{Per}(\mathcal{X}, N) + M_P \cdot (N-1)] - [M_C \cdot \phi_{Per}(\mathcal{X}, N-1) + M_P \cdot (N-2)] \leq 0 \\ & \phi_{Per}(\mathcal{X}, N) - \phi_{Per}(\mathcal{X}, N-1) \leq -\frac{M_P}{M_C} \end{aligned} \quad (\text{III. 75})$$

Par ailleurs, les limites aux bornes de $\Gamma_{Per}(\mathcal{X}, N)$ donnent :

$$\begin{aligned} \lim_{N \rightarrow 1} \Gamma_{Per}(\mathcal{X}, N) &= \lim_{\rho \rightarrow 0} \left(\underbrace{M_C \cdot \phi_{Per}(\mathcal{X}, N)}_{\rightarrow \text{constante : } M_C \cdot \phi_{Min}(\mathcal{X})} + \underbrace{M_P \cdot (N-1)}_{\rightarrow 0} \right) \\ &= M_C \cdot \phi_{Min}(\mathcal{X}) \end{aligned} \quad (\text{III. 76})$$

$$\begin{aligned} \lim_{N \rightarrow +\infty} \Gamma_{Per}(\mathcal{X}, N) &= \lim_{\rho \rightarrow 0} \left(\underbrace{M_C \cdot \phi_{Per}(\mathcal{X}, N)}_0 + \underbrace{M_P \cdot (N-1)}_{\rightarrow +\infty} \right) \\ &= +\infty \end{aligned} \quad (\text{III. 77})$$

Il existe donc un nombre d'intervalles de maintenance N^* qui satisfait les relations suivantes :

$$N^* = \begin{cases} \Gamma_{Per}(X, N+1) - \Gamma_{Per}(X, N) \geq 0 \\ \Gamma_{Per}(X, N) - \Gamma_{Per}(X, N-1) \leq 0 \\ \lim_{N \rightarrow 1} \Gamma_{Per}(X, N) = \text{constante} \\ \lim_{N \rightarrow +\infty} \Gamma_{Per}(X, N) = +\infty \end{cases} \quad (\text{III. 78})$$

La résolution de cette politique de maintenance à l'aide d'une procédure numérique s'effectue en incrémentant le nombre d'intervalles de maintenance jusqu'à trouver un N^* satisfaisant les deux première relations de l'équation (III. 78) [Schutz, 2009B]. Cependant, les durées des actions préventives n'étant pas négligeables, la durée allouée pour ces maintenances n'est pas assurément suffisante pour atteindre ce nombre optimal. Il existe donc un nombre maximal de maintenances préventives donné par la relation suivante, issue de la contrainte temporelle :

$$\sum_{i=1}^{\dim(X)} \delta_{x_i} + \mu_p \cdot (N-1) \leq H$$

$$(N-1)^{\max} = \left\lfloor \frac{H - \sum_{i=1}^{\dim(X)} \delta_{x_i}}{\mu_p} \right\rfloor \quad (\text{III. 79})$$

La procédure numérique est présentée par la figure 36. Ayant prouvé l'existence d'un minimum local, cette procédure numérique incrémente le nombre d'intervalles de maintenances préventives jusqu'à satisfaire la relation (III. 80) ou arriver jusqu'à la valeur maximale admissible, $N^{\max}$.

$$\begin{cases} \Gamma_{Per}(X, N^* - 1) > \Gamma_{Per}(X, N^*) \\ \Gamma_{Per}(X, N^* + 1) > \Gamma_{Per}(X, N^*) \end{cases} \quad (\text{III. 80})$$

VI.1. Recherche du facteur d'amélioration optimal

Comme pour les politiques précédentes, l'effet des maintenances préventives dépend du facteur d'amélioration ρ . Bien évidemment, la modification de ce facteur impacte le coût et la durée des actions préventives.

Lorsque le coût d'une action préventive dépend du facteur d'amélioration, le coût total de maintenance est donné par :

$$\Gamma_{Per}(X, N, \rho) = M_c \cdot \phi_{Per}(X, N, \rho) + M_p(\rho) \cdot (N-1) \quad (\text{III. 81})$$

Avec :

- $\Gamma_{Per}(X, N, \rho)$: coût de maintenance associé à la politique périodique, constitué de N intervalles de maintenances préventives, pour le plan d'exploitation défini par le vecteur X et dont le facteur d'amélioration est défini par ρ ,
- $\phi_{Per}(X, N, \rho)$: nombre moyen de pannes, pour la politique périodique, pouvant s'associer au plan d'exploitation défini par le vecteur X , avec un facteur d'amélioration ρ ,
- ρ : facteur d'amélioration utilisé pour chaque action préventive.

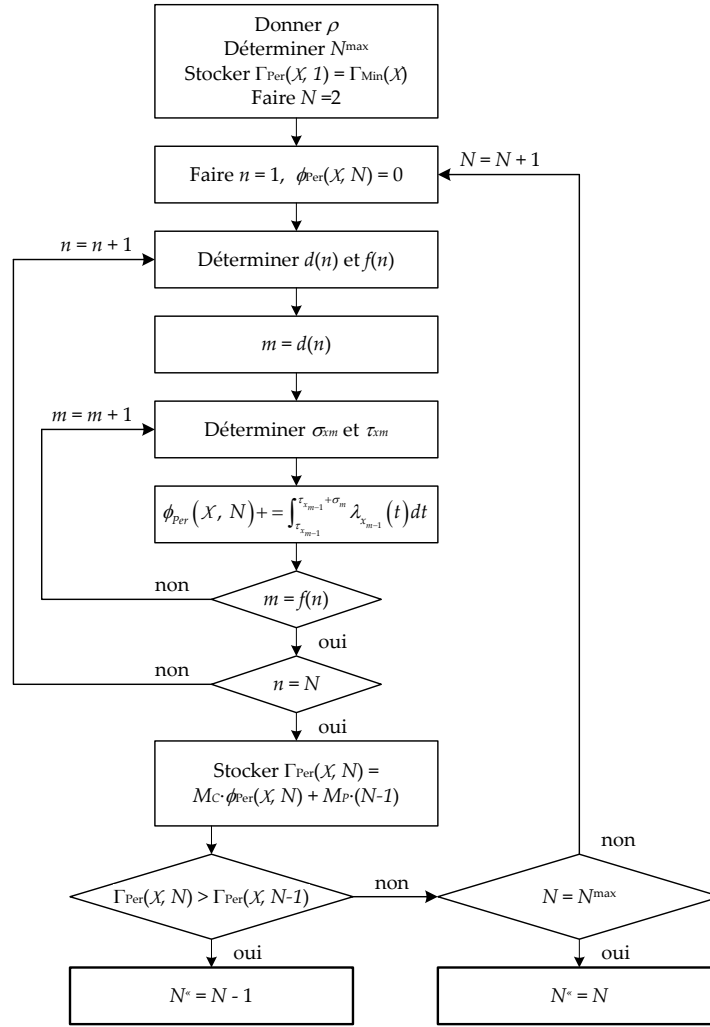


Figure 36 : Procédure numérique pour l'obtention du nombre optimal d'intervalles de maintenances préventives

Afin que cette politique soit plus rentable que la politique minimaliste, il existe une valeur maximale admissible pour le facteur d'amélioration, notée $\rho^{\max}$, qui doit satisfaire la relation (III. 82). Le nombre moyen de pannes faisant également référence au facteur d'amélioration, le facteur d'amélioration optimal sera évidemment inférieur à :

$$\Gamma_{Min}(X) > (N-1) \cdot M_p(\rho^{\max})$$

$$\rho^{\max} = \min \left(1, \frac{1}{M_p^v} \cdot \sqrt[2]{\frac{\Gamma_{Min}(X)}{(N-1)} - M_p^f} \right) \quad (\text{III. 82})$$

Dès lors que les durées de maintenances préventives dépendent du facteur d'amélioration, la contrainte temporelle se trouve modifiée comme suit :

$$\sum_{i=1}^{\dim(X)} \delta_{x_i} + \left[\mu_p^f + (\rho \cdot \mu_p^v)^2 \right] \cdot (N-1) \leq H \quad (\text{III. 83})$$

Par rapport à cette contrainte temporelle et donc à la durée de temps allouée par le plan d'exploitation pour la réalisation des actions préventives, le nombre maximal de maintenances

préventives est obtenu lorsque ρ tend vers 0. Dans ce cas, le nombre maximal d'intervalles de maintenances préventives est donné par :

$$(N-1)^{\max} = \begin{cases} \left\lfloor \frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\mu_p^f} \right\rfloor - 1 & \text{si } \frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\mu_p^f} = \left\lfloor \frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\mu_p^f} \right\rfloor > 0 \\ \left\lfloor \frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\mu_p^f} \right\rfloor & \text{sinon} \end{cases} \quad (\text{III. 84})$$

Par corollaire, pour chaque nombre d'intervalles de maintenances préventives, il existe une valeur maximale admissible pour le facteur d'amélioration, définie par :

$$\rho^{\max} = \min \left(1, \frac{1}{\mu_p^v} \cdot \sqrt[2]{\frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{(N-1)}} - \mu_p^f \right) \quad (\text{III. 85})$$

La solution, pour laquelle le nombre de possibilités est assez faible, peut être obtenue grâce à une procédure numérique dérivée de la procédure présentée par la figure 36.

Suivant cette politique de maintenance, plusieurs actions préventives seront réalisées durant le plan d'exploitation. L'étape suivante s'intéresse aux valeurs optimales prises par les facteurs d'amélioration associés à chaque action préventive en vue de minimiser les coûts de maintenance.

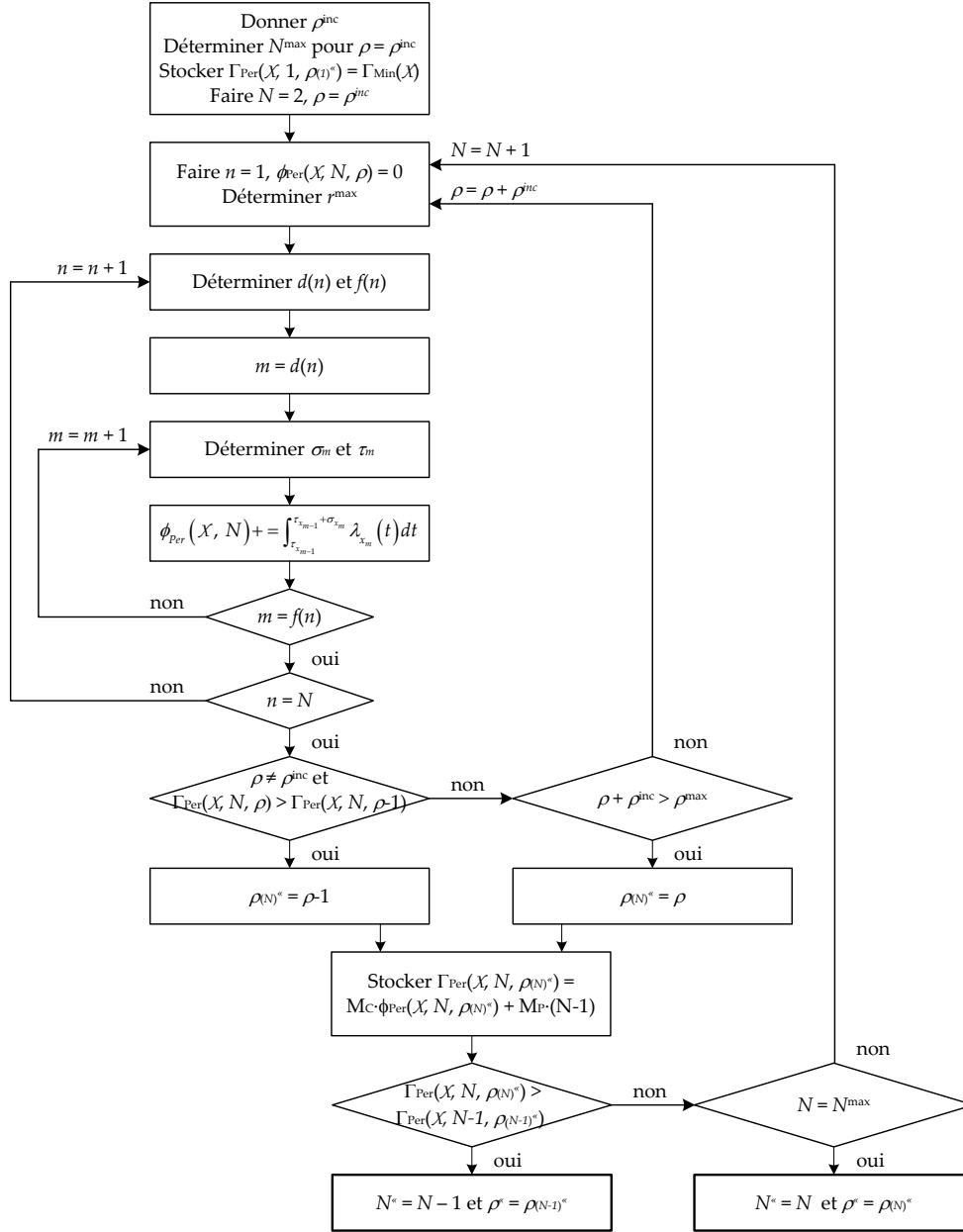
VI.2. Recherche des facteurs d'amélioration optimaux

Lorsque le facteur d'amélioration diffère suivant chaque action préventive, le coût total de maintenance s'exprime comme suit :

$$\Gamma_{\text{Per}}(\mathcal{X}, N, P) = M_C \cdot \phi_{\text{Per}}(\mathcal{X}, N, P) + \sum_{i=1}^{\dim(P)} M_p(\rho_i) \quad (\text{III. 86})$$

Avec :

- $\Gamma_{\text{Per}}(\mathcal{X}, N, P)$: coût de maintenance pour la politique périodique, constitué de N intervalles de maintenances préventives, pour le plan d'exploitation défini par le vecteur $\mathcal{X}$ et dont les facteurs d'amélioration sont donnés par le vecteur P ,
- $\phi_{\text{Per}}(\mathcal{X}, N, P)$: nombre moyen de pannes, pour la politique périodique, pouvant s'associer au plan d'exploitation défini par le vecteur $\mathcal{X}$, avec le vecteur des facteurs d'amélioration P ,
- P : vecteur précisant le facteur d'amélioration utilisé pour chaque action préventive.

Figure 37 : Procédure numérique pour l'obtention du N^* et ρ^*

La contrainte temporelle à respecter est donnée par :

$$\sum_{i=1}^{\dim(X)} \delta_{x_i} + \sum_{i=1}^{\dim(P)} \mu_p^f + (\rho_i \cdot \mu_p^v)^2 \leq H \quad (\text{III. 87})$$

La détermination du nombre optimal de maintenances préventives et des facteurs d'amélioration associés sera réalisée à l'aide d'une méta-heuristique du type algorithme du grand déluge, du fait du codage employé.

VI.2.a. Méta-heuristique pour la recherche des facteurs d'amélioration optimaux

Pour cette politique de maintenance, le nombre de variable de décision n'est pas fixe. En effet, la valeur de la variables de décision précisant le nombre de maintenances préventives à réaliser

correspond au nombre de variables de décision relatives aux facteurs d'amélioration. En d'autres termes, si X maintenances préventives doivent être réalisées, ils convient de déterminer les X facteurs d'amélioration $\rho_1, \rho_2, \dots, \rho_X$. Par conséquent, le nombre de facteurs d'amélioration à déterminer est variable.

Le codage employé consiste à utiliser un segment où chaque gène correspond à un facteur d'amélioration. Ainsi, la dimension du segment définit le nombre de maintenances préventives à réaliser.

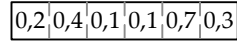


Figure 38 : Codage employé pour représenter les variables de décisions

Dans l'exemple présenté par la figure 38, le segment est constitué de six gènes. Par conséquent, le plan de maintenance sera composé de six maintenances préventives.

À partir de ce codage, le voisinage d'une solution, c'est-à-dire les solutions potentielles, peut être obtenu suivant deux types d'opérations. La première opération consiste à changer la valeur d'un gène. Avec cette modification, il convient de s'assurer du respect de la contrainte temporelle (III. 87). La seconde opération a pour but de modifier le nombre de maintenances préventives en en ajoutant ou, au contraire, en en supprimant. La suppression de maintenances préventives ne sera possible que si la dimension du vecteur est strictement supérieure à 2, auquel cas, nous retomberions sur la politique minimaliste. Bien évidemment, l'ajout de maintenances préventives devra toujours satisfaire la contrainte temporelle.

VII. La politique séquentielle

La politique séquentielle se rapproche de la politique périodique dans la mesure où les maintenances préventives peuvent être réalisées durant l'accomplissement des missions. La différence avec la politique périodique se situe au niveau des intervalles de temps qui ne sont plus inéluctablement constants pour la politique séquentielle. Au contraire, il est logique de penser que plus le système vieillit, plus les actions de maintenances seront rapprochées. Cette affirmation est vraie lorsque la loi de défaillance est constante au cours du temps [Nakawaga, 2009]. Par contre, lorsque les conditions de fonctionnement interviennent dans l'expression de cette loi de défaillance, cette expression ne peut plus être vérifiée [Schutz, 2009B]. L'objectif consiste donc à déterminer les tailles optimales des différents intervalles de maintenances préventives. Sous cette politique, le coût total de maintenance est donné par :

$$\Gamma_{Seq}(\mathcal{X}, \mathcal{N}) = M_C \cdot \phi_{Seq}(\mathcal{X}, \mathcal{N}) + M_P \cdot (\dim(\mathcal{N}) - 1) \quad (\text{III. 88})$$

Avec :

- $\Gamma_{Seq}(\mathcal{X}, \mathcal{N})$: coût total de maintenance pour la politique séquentielle et pour un plan d'exploitation défini par le vecteur $\mathcal{X}$, et dont les intervalles de maintenances préventives sont donnés par le vecteur $\mathcal{N}$,

- $\phi_{Seq}(\mathcal{X}, \mathcal{N})$: nombre moyen de pannes pour la politique séquentielle, associé au plan d'exploitation défini par le vecteur $\mathcal{X}$ et dont les intervalles de maintenances préventives sont donnés par le vecteur $\mathcal{N}$,
- $\mathcal{N}$: vecteur précisant les durées des intervalles de maintenances,
- $\dim(\mathcal{N})$: dimension du vecteur $\mathcal{N}$.

Le nombre moyen de pannes pour cette politique est présenté ci-dessous :

$$\phi_{Seq}(\mathcal{X}, \mathcal{N}) = \sum_{n=1}^{\dim(\mathcal{N})} \left(\sum_{m=d(n)}^{f(n)} \left(\int_{\tau_{x_{m-1}}}^{\tau_{x_{m-1}} + \sigma_{x_m}} \lambda_{x_m}(t) dt \right) \right) \quad (\text{III. 89})$$

Compte tenu de la politique séquentielle, les durées des intervalles de maintenances préventives ne sont pas constantes. Il en résulte une modification de l'expression des durées de chaque mission, comme exprimée par les relations suivantes :

$$\sigma_{x_m} = \begin{cases} \delta_{x_m} & \text{si } d(n) \neq m \neq f(n) \\ \sum_{l=1}^{d(n)} \delta_{x_l} - \sum_{i=1}^{n-1} \Delta_i & \text{si } d(n) = m \neq f(n) \\ \sum_{i=1}^n \Delta_i - \sum_{l=1}^{f(n)-1} \delta_{x_l} & \text{si } d(n) \neq m = f(n) \\ \Delta_n & \text{si } d(n) = m = f(n) \end{cases} \quad (\text{III. 90})$$

L'équation (III. 90) entraîne une modification de l'expression des missions $d(n)$ et $f(n)$:

$$d(n) = \begin{cases} 1 & \text{si } n = 1 \\ m \text{ solution de } \begin{cases} \sum_{i=1}^m \delta_{x_i} > \sum_{i=1}^{(n-1)} \Delta_i \\ \sum_{i=1}^{m-1} \delta_{x_i} \leq \sum_{i=1}^{(n-1)} \Delta_i \end{cases} & \text{sinon} \end{cases} \quad (\text{III. 91})$$

$$f(n) = \begin{cases} \dim(\mathcal{X}) & \text{si } n = N \\ m \text{ solution de } \begin{cases} \sum_{i=1}^m \delta_{x_i} \geq \sum_{i=1}^n \Delta_i \\ \sum_{i=1}^{m-1} \delta_{x_i} < \sum_{i=1}^n \Delta_i \end{cases} & \text{sinon} \end{cases} \quad (\text{III. 92})$$

Dans l'équation (III. 89), la valeur de $\tau_{x_{m-1}}$ reste inchangée par rapport à l'expression de cette durée fonctionnelle exposée pour la politique périodique.

Lorsque le facteur d'amélioration, les coûts et les durées sont fixés pour l'ensemble des maintenances et qu'un plan d'exploitation est défini, il existe un nombre maximal de maintenances préventives réalisables, issu de la contrainte temporelle :

$$\sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} + \mu_p \cdot \dim(\mathcal{N}) \leq H$$

$$\dim(\mathcal{N})^{\max} = 1 + \left\lfloor \frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i}}{\mu_p} \right\rfloor \quad (\text{III. 93})$$

Pour la détermination des variables de décisions, un principe similaire à la politique périodique est employé. Connaissant le nombre maximal de maintenances préventives envisageables, le principe consiste à déterminer pour chaque valeur de $\dim(\mathcal{N})$, compris dans l'intervalle $]1, \dim(\mathcal{N})^{\max}]$ les intervalles de temps optimaux.

Malheureusement, le nombre de cas différents ne permet pas une résolution à l'aide de procédures numériques. En supposant que les durées des intervalles de maintenances préventives soient multiples de l'unité de temps, la planification de trois actions de maintenances pour un horizon de temps, constitué de 1000 unités de temps, offre plus de 165 millions de possibilités. Sachant que la détermination du plan de maintenance est une étape de l'algorithme génétique du choix du plan d'exploitation (stratégie intégrée), il convient d'obtenir une solution satisfaisante en un temps raisonnable, c'est-à-dire en utilisant une méta-heuristique.

VII.1. Méta-heuristique pour l'obtention de solutions efficaces

Comme ce fut le cas pour la procédure numérique développée pour la politique périodique, le principe utilisé consiste à déterminer, pour chaque valeur de $\dim(\mathcal{N})$, les solutions efficaces jusqu'à satisfaire la relation (III. 94) ou atteindre la valeur de $\dim(\mathcal{N})^{\max}$.

$$\begin{cases} \Gamma_{Seq} \left(X, \mathcal{N}^* \right)_{\dim(\mathcal{N}^*)=K-1} > \Gamma_{Seq} \left(X, \mathcal{N}^* \right)_{\dim(\mathcal{N}^*)=K} \\ \Gamma_{Seq} \left(X, \mathcal{N}^* \right)_{\dim(\mathcal{N}^*)=K+1} > \Gamma_{Seq} \left(X, \mathcal{N}^* \right)_{\dim(\mathcal{N}^*)=K} \end{cases} \quad (\text{III. 94})$$

Où K est une valeur comprise entre $]2, \dim(\mathcal{N})^{\max} - 1]$.

VII.1.a. Le codage

Pour cet algorithme génétique, l'objectif étant de déterminer les intervalles de maintenances optimaux, le codage employé doit être représentatif de ces durées. Cependant, si les gènes représentent effectivement les durées des intervalles de maintenances, la majorité des individus issus des opérations génétiques ne sera pas viable. En effet, l'inversion de gènes entre deux parents ne garantit pas que la durée cumulée des intervalles de maintenances préventives reste égale à la durée de l'horizon de temps. Pour pallier ces problèmes, le codage des individus portera sur la date de réalisation de l'action de maintenance. Par conséquent, chaque gène correspond à un entier compris dans l'intervalle $[1, H-1]$ correspondant à la date à laquelle une action préventive devra être réalisée. Cependant, il conviendra de s'assurer de l'unicité des valeurs des gènes. Si deux gènes ont la même valeur, cela signifierait que deux actions de maintenances préventives ont lieu simultanément, ce qui s'avère être une situation improbable.

121	228	347	500	609	723	845	914
-----	-----	-----	-----	-----	-----	-----	-----

Figure 39 : Codage utilisé pour la politique séquentielle

Dans le codage, les gènes doivent être ordonnés suivant la date de réalisation de la maintenance préventive. De ce fait, la durée du $i^{\text{ème}}$ intervalle de maintenance préventive ($i = \{2, \dots, \dim(\mathcal{N})-1\}$) est déterminée par la différence entre les $i^{\text{ème}}$ et $(i-1)^{\text{ème}}$ gènes. Concernant le premier intervalle, sa durée correspond à la valeur du premier gène et la durée du dernier intervalle est donnée par $H - (\dim(\mathcal{N}) - 1)^{\text{ème}}$ gène.

En prenant l'exemple du codage proposé par la figure 39 et en supposant que la durée de l'horizon de temps est de 1000 ut, les durées des neuf intervalles de maintenances préventives sont {121, 107 (228 - 121), 119 (347 - 228), 153, 109, 114, 122, 69, 86 (1000-914)}.

VII.1.b. L'opérateur de croisement

L'opération de croisement utilisée pour cet algorithme génétique reste inchangée par rapport aux précédents algorithmes génétiques utilisés. À partir du point de croisement sélectionné aléatoirement, les gènes des deux individus parents sont échangés afin de générer deux nouveaux individus. Parmi les deux individus générés, il peut arriver que l'un des deux présente des erreurs génétiques. Les deux erreurs possibles sont des dates de maintenances préventives en double ou un mauvais ordonnancement de ces dates.

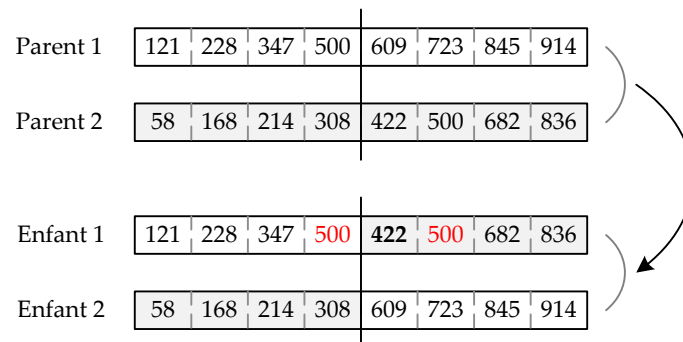


Figure 40 : Erreurs génétiques

L'exemple présenté par la figure 40 illustre les deux erreurs génétiques qui peuvent intervenir. Dans le chromosome « Enfant 1 », le gène portant la valeur « 500 » est présent deux fois. Par ailleurs, nous pouvons remarquer qu'après la première valeur « 500 », le gène suivant indique une valeur inférieure. Après les opérations génétiques de croisement, il convient de vérifier la viabilité des nouveaux individus et de les corriger, si nécessaire.

Lorsque des valeurs de gènes apparaissent en double, la correction s'effectue simplement en régénérant une des deux valeurs. Lorsque l'unicité de chaque gène est vérifiée, la solution résultante peut, néanmoins, présenter la seconde erreur génétique, à savoir un mauvais ordonnancement des dates de maintenances. Cette seconde erreur est rectifiée en réorganisant les gènes par valeur croissante.

VII.1.c. L'opérateur de mutation

Pour l'opérateur génétique de mutation, deux types d'opérations sont utilisés. La première opération consiste à supprimer un gène choisi aléatoirement et à régénérer une nouvelle valeur différente. À la suite de cette opération, il convient de réorganiser l'ensemble des dates de maintenances préventives. La seconde opération de mutation utilisée consiste à modifier la valeur d'un gène choisi aléatoirement. La nouvelle date générée doit appartenir à l'intervalle bornant la valeur initiale. Si le premier gène est choisi, la nouvelle date devra appartenir à l'intervalle entre 0 et la valeur du deuxième gène (exclu). Par contre, si la sélection porte sur le dernier gène, la nouvelle date sera comprise entre la valeur du pénultième gène et la valeur de H .

VII.2. Recherche du facteur d'amélioration optimal

Comme pour les politiques de maintenances précédentes, une amélioration concerne la recherche du facteur d'amélioration optimal ρ^* , suivant lequel les coûts et/ou les durées des maintenances peuvent varier. De ce fait, l'expression du coût total de maintenance s'exprime par :

$$\Gamma_{Seq}(\mathcal{X}, \mathcal{N}, \rho) = M_C \cdot \phi_{Seq}(\mathcal{X}, \mathcal{N}, \rho) + M_P(\rho) \cdot (\dim(\mathcal{N}) - 1) \quad (\text{III. 95})$$

Avec :

- $\Gamma_{Seq}(\mathcal{X}, \mathcal{N}, \rho)$: coût total de maintenance pour la politique séquentielle suivant plan d'exploitation défini par le vecteur $\mathcal{X}$, avec les intervalles de maintenances préventives donnés par le vecteur $\mathcal{N}$, et suivant le facteur d'amélioration ρ ,
- $\phi_{Seq}(\mathcal{X}, \mathcal{N}, \rho)$: nombre moyen de pannes pour la politique séquentielle, associé au plan d'exploitation défini par le vecteur $\mathcal{X}$, et dont les intervalles de maintenances préventives sont donnés par le vecteur $\mathcal{N}$ et le facteur d'amélioration par ρ ,
- ρ : facteur d'amélioration associé aux maintenances préventives.

Lorsque les durées dépendent du facteur ρ , la contrainte temporelle devient :

$$\sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} + \mu_P(\rho) \cdot \dim(\mathcal{N}) \leq H \quad (\text{III. 96})$$

De manière identique à la politique périodique, il existe un nombre maximum de maintenances préventives réalisable pour satisfaire la durée allouée par le plan d'exploitation. La relation donnée par l'équation (III. 84) reste valable. De la même manière, la valeur maximale admissible pour le facteur d'amélioration suivant chaque valeur de $\dim(\mathcal{N})$ résulte de l'équation (III. 85).

Cependant, le facteur d'amélioration ρ^* et les différents intervalles de maintenances préventives ne peuvent pas être résolus à partir d'une procédure numérique à cause de l'explosion combinatoire précisée auparavant. Cependant, l'algorithme présenté pour la détermination des intervalles optimaux peut être utilisé en l'exécutant séparément pour chaque valeur du facteur d'amélioration. Compte tenu de la précision souhaitée (10^{-1}), neuf itérations de l'algorithme génétique seront réalisées, c'est-à-dire une pour chaque valeur du facteur d'amélioration. La solution optimale sera obtenue parmi les meilleures solutions déterminées pour chaque valeur du facteur d'amélioration.

Pour optimiser davantage cette politique de maintenance, il s'avèrerait intéressant de déterminer, en plus du nombre et des durées des intervalles de maintenances préventives, les facteurs d'amélioration optimaux associés à chaque action préventive.

VII.3. Recherche des facteurs d'amélioration optimaux

Lorsque les facteurs d'amélioration varient d'une maintenance à une autre, l'expression du coût est donnée par :

$$\Gamma_{Seq}(\mathcal{X}, \mathcal{N}, \mathbf{P}) = M_C \cdot \phi_{Seq}(\mathcal{X}, \mathcal{N}, \mathbf{P}) + \sum_{i=1}^{\dim(\mathbf{P})} M_p(\rho_i) \quad (\text{III. 97})$$

Avec :

- $\mathbf{P}$: vecteur définissant les facteurs d'amélioration de chaque maintenance préventive,
- ρ_i : composante du vecteur $\mathbf{P}$; le facteur ρ_i est associé à la $i^{\text{ème}}$ maintenance préventive.

Pour tenir compte des différentes valeurs des facteurs d'amélioration, la contrainte temporelle s'écrit comme :

$$\sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} + \sum_{i=1}^{\dim(\mathbf{P})} \mu_p(\rho_i) \leq H \quad (\text{III. 98})$$

Quand le facteur d'amélioration est constant pour toutes les maintenances, il est possible de déterminer la valeur maximale prise par ce facteur. Par contre, quand la valeur peut varier suivant chaque action préventive, seule la relation suivante peut être établie :

$$\sum_{i=1}^{\dim(\mathbf{P})} ((\rho_i)^2) \leq \frac{H - \sum_{i=1}^{\dim(\mathcal{X})} \delta_{x_i} - \dim(\mathbf{P}) \cdot \mu_p^f}{(\mu_p^v)^2} \quad (\text{III. 99})$$

Sans grande surprise, les différentes variables de décisions (nombre d'intervalles, durées des intervalles, valeurs des facteurs d'amélioration) seront estimées à l'aide d'un algorithme génétique.

VII.3.a. Méta-heuristique pour la recherche des variables de décision

Dans l'ensemble des autres cas développés, les variables de décisions portaient sur les valeurs d'un vecteur, voire sur sa dimension. Pour la situation présente, les variables de décisions portent sur deux vecteurs. Par conséquent, pour le codage des individus, nous aurons recours à l'utilisation de deux chromosomes. Le premier segment sera associé aux dates de maintenances alors que le second fera référence aux facteurs d'amélioration.

Compte tenu du codage adopté, les opérations génétiques porteront simultanément sur les deux chromosomes. En considérant l'opérateur de croisement, le point de croisement sera commun aux deux segments. Cette possibilité est offerte car ces deux segments sont de même taille. Pareillement, le gène choisi pour l'opération de mutation aura la même place dans les deux chromosomes.

Lors de ces opérations génétiques, il conviendra de s'assurer de la viabilité des individus. Pour le segment associé aux facteurs d'amélioration, l'individu sera considéré viable, si la relation (III. 99) est respectée. Concernant le chromosome représentant les dates de maintenances préventives, la viabilité de ce dernier est assurée si toutes les dates spécifiées sont uniques et l'ordre à l'intérieur du chromosome croissant.

Dans le cas où des erreurs génétiques sont présentes, un processus de réparation doit être appliqué. Si le problème concerne les facteurs d'amélioration, la correction s'effectuera en diminuant la valeur d'un gène. Le niveau de réduction du facteur d'amélioration ainsi que le choix du gène sont aléatoires. Néanmoins, la valeur du gène ne pourra être inférieure à 0,1 (précision souhaitée). Ce processus est répété tant que la relation (III. 99) est violée. Lorsque la non-viabilité est liée au segment associé au vecteur $\mathcal{N}$, la correction du chromosome s'effectuera suivant la manière présentée dans la partie « VII.1 Méta-heuristique pour l'obtention de solutions efficaces ».

VIII. Conclusion

À travers ce chapitre, nous avons étudié plusieurs types de politiques de maintenances. Les premières politiques se caractérisent par des maintenances préventives dont les dates de réalisation coïncident avec des dates de fin de mission. Ainsi, la politique systématique se caractérise par l'accomplissement de maintenances préventives à la fin de chaque mission, de manière similaire aux maintenances réalisées sur les avions après chaque vol. Cependant, pour éviter des maintenances préventives trop fréquentes et, de ce fait, pas nécessairement utiles, la politique sporadique se caractérise par la réalisation de maintenances préventives à la fin de missions, mais plus de façon automatique. Le choix des missions suivies de maintenances préventives est déterminée à l'aide d'une méta-heuristique. Les deux politiques de maintenances suivantes se différencient par des dates de réalisation distinctes des dates de fin de mission. Ainsi, la première politique, dite périodique, se définit par des actions préventives réalisées à des intervalles de temps constants. L'objectif consiste alors à déterminer le nombre optimal de maintenances préventives. Finalement, la dernière politique séquentielle est une extension de la politique périodique dans la mesure où l'on cherche, en plus du nombre de maintenances préventives, les durées optimales de ces intervalles. Pour chacune de ces politiques de maintenances, nous avons cherché à déterminer les facteurs d'améliorations optimaux qui permettent de minimiser les coûts liés aux actions de maintenances. Sachant que la loi de défaillance est dépendante du plan d'exploitation, le chapitre suivant considèrera un plan d'exploitation donné, pour évaluer les différentes politiques de maintenances proposées. Par ailleurs, nous pourrions également visualiser l'impact de nouvelles contraintes, telles qu'un seuil de fiabilité minimal à ne pas franchir, une disponibilité minimale du système à respecter, etc.

Chapitre IV

Exemple illustratif

CE dernier chapitre est consacré à l'étude numérique des différentes stratégies d'optimisation couplées aux différentes politiques de maintenances présentées dans le chapitre précédent. La première partie de ce chapitre sera consacrée à l'étude de la stratégie d'optimisation séquentielle en tentant d'y appliquer chacune des politiques de maintenances abordées. La seconde partie sera dédiée à l'étude de la stratégie séquentielle, permettant ainsi de comparer ces deux stratégies d'optimisation ainsi que la variation des coûts liée aux politiques de maintenances.

Sommaire

I. Introduction	98
II. Données numériques	98
III. Première étude : la stratégie séquentielle	100
IV. Seconde étude : la stratégie intégrée	102
V. Synthèse au sujet des stratégies séquentielle et intégrée	109

I. Introduction

À travers les chapitres précédents, ont été exposées différentes stratégies d'optimisation et politiques de maintenance. Dans le présent chapitre et au travers d'un cas illustratif, nous allons tenter d'appliquer chaque politique de maintenance aux stratégies étudiées. À partir d'un ensemble de missions, deux études seront menées. La première étude se caractérisera par une stratégie séquentielle pour laquelle un plan d'exploitation optimal sera déterminé. À ce plan d'exploitation, la politique de maintenance optimale sera déterminée. Pour la seconde étude, la stratégie intégrée sera appliquée. Il s'agira de déterminer conjointement le plan d'exploitation optimal, ainsi que la politique de maintenance adaptée. Bien évidemment, la détermination des politiques de maintenances optimales nécessite la recherche des différentes variables de décisions (facteur d'amélioration, durées des intervalles de maintenances préventives, etc.)

II. Données numériques

L'étude des deux stratégies se basera sur les données numériques présentées ci-dessous :

II.1. Les missions et l'horizon de temps

Mission (m)	1	2	3	4	5	6
Durée (δ_m)	106	107	82	104	108	121
Profit (ψ_m)	4654	6031	4694	4484	4719	4413
Les conditions de fonct. $Z_m = \{z_m^{\text{CO}}, z_m^{\text{CO}}\}$	{2, 4}	{4, 7}	{7, 4}	{9, 3}	{3, 4}	{1, 2}

Mission (m)	7	8	9	10	11	12
Durée (δ_m)	74	64	91	100	114	112
Profit (ψ_m)	5670	4045	4746	6031	6223	4918
Les conditions de fonct. $Z_m = \{z_m^{\text{CO}}, z_m^{\text{CO}}\}$	{2, 5}	{6, 6}	{6, 9}	{4, 5}	{5, 4}	{9, 2}

Mission (m)	13	14	15	16	17	18
Durée (δ_m)	60	110	101	138	128	108
Profit (ψ_m)	4904	6085	5567	5196	5716	4429
Les conditions de fonct. $Z_m = \{z_m^{\text{CO}}, z_m^{\text{CO}}\}$	{6, 2}	{4, 3}	{7, 5}	{5, 2}	{4, 7}	{3, 1}

Tableau 6 : Liste des missions envisageables

L'horizon de temps considéré (H) est de 1000 unités de temps alors que la durée cumulée des missions s'élève à 1826 ut.

II.2. Les données relatives à la fiabilité du système

Les données dévoilées ci-dessous concernent la fiabilité du système ainsi que les coûts et durées associés aux actions de maintenance.

La loi de défaillance caractérisant les conditions nominales est du type Weibull, définie par :

- Paramètre de forme (β) : 2,5
- Paramètre d'échelle (η) : 600
- Paramètre de position (γ) : 0

Ainsi, les fonctions du taux nominal de défaillance et de fiabilité s'expriment par :

- $\lambda_0(t) = \left(\frac{2,5}{600}\right) \cdot \left(\frac{t}{600}\right)^{(2,5-1)}$
- $R_0(t) = e^{-\left(\frac{t}{600}\right)^{2,5}}$

Concernant le modèle à risque proportionnel, l'évaluation des coefficients des conditions de fonctionnement renvoient les valeurs :

- Coefficient associé aux conditions opérationnelles (β^{CO}) : 0,025
- Coefficient associé aux conditions opérationnelles (β^{CE}) : 0,042

Les coûts et durées des actions de maintenances correctives et préventives, sont exprimés ci-dessous :

- $M_C = 3000$
- $M_P(\rho) = 500 + (\rho \cdot 40)^2$
- $\mu_P(\rho) = 1 + (\rho \cdot 2)^2$

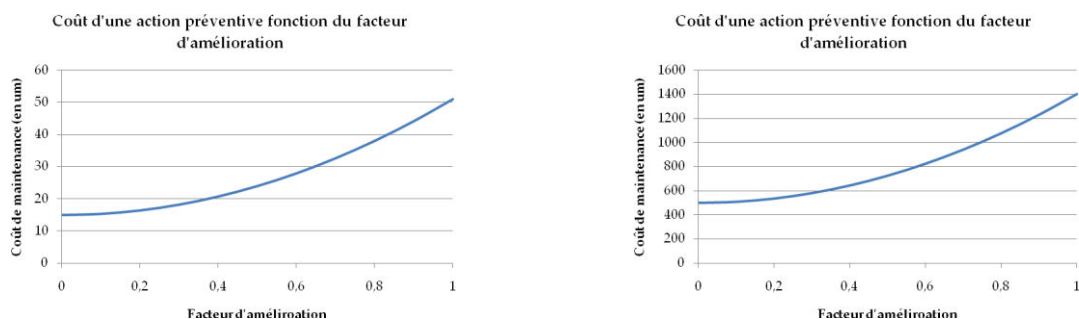


Figure 41 : Représentation du coût et de la durée d'une action préventive en fonction du facteur d'amélioration

À présent que les différentes données sont exprimées, l'étude numérique peut porter sur la stratégie séquentielle.

II.3. Les données relatives aux méta-heuristiques

Les données relatives à l'utilisation des différentes méta-heuristiques sont indiquées ci-dessous :

Pour le choix des plans d'exploitation et l'optimisation des politiques de maintenances, les caractéristiques des algorithmes génétiques sont :

- Nombre de générations : 100
- Nombre d'individus par génération : 100
- Taux de croisement : 80%
- Taux de mutation (appliqué aux individus issus du croisement) : 5%
- Taux d'élitisme : 10%
- Taux de renouvellement : 10%

Au sujet de l'algorithme du grand déluge utilisé pour l'optimisation de la politique sporadique, très peu de paramètres sont nécessaires.

- Nombre total d'itérations : 10000
- Nombre d'itérations infructueuses (généralisant l'arrêt de l'algorithme) : 500
- Quantité ΔL : 100 um

III. Première étude : la stratégie séquentielle

Comme explicitée au cours du chapitre III, cette stratégie séquentielle débute par la détermination du plan d'exploitation optimal.

III.1. Détermination du plan d'exploitation optimal

La détermination du plan d'exploitation optimal consiste dans un premier temps à déterminer l'ensemble des missions à réaliser. En effet, à cette étape, l'ordonnancement des missions n'a aucune influence puisque seule la maximisation des profits est recherchée. En conséquence, ce problème d'optimisation peut se rapprocher d'un problème du type *Knapsack*. Dans notre situation, le nombre de plans d'exploitation différents reste somme toute acceptable. En effet, en se basant sur les 18 missions, le nombre de possibilités s'élève à $2^{18} = 262\,144$. Néanmoins, la durée cumulée de toutes ces missions est très nettement supérieure à l'horizon de temps alloué. En considérant les missions ordonnées par durées croissantes, le nombre maximal de missions pour satisfaire l'horizon de temps s'élève à 11. Par conséquent, le nombre de plans d'exploitations à tester peut être exprimé par :

$$\sum_{i=1}^{11} \frac{\dim(\mathcal{M})!}{i!(\dim(\mathcal{M})-i)!} = \quad (\text{IV. 1})$$

À titre informatif, la recherche du plan d'exploitation par une résolution approchée du type algorithme glouton génère un bénéfice de 53 996 um (unités monétaires) pour une durée totale de 903 ut (unités de temps). Ce plan d'exploitation est constitué des missions suivantes : 13, 7, 8, 10, 3, 2,

14, 15, 11, 9. Pour cette approche gloutonne, les missions sont ajoutées suivant le critère de profit par unités de temps. Bien que le but recherché soit la maximisation du profit, ce critère permet de sélectionner les missions les plus intéressantes tout en considérant la contrainte temporelle. Par exemple, si deux missions A et B (de durées $\delta_A > \delta_B$) génèrent un même profit ($\psi_A = \psi_B$), la première mission à ajouter sera la mission A , suivie de la mission B car $\psi_A/\delta_A > \psi_B/\delta_B$.

Grâce à l'utilisation d'une méta-heuristique, le profit total généré s'élève à 56 911 um avec le plan d'exploitation constitué des missions 1, 2, 3, 4, 7, 8, 9, 10, 13, 14 et 15.

III.2. Les politiques de maintenances

La politique de maintenance minimaliste est exclusivement composée de maintenances correctives, dont les durées de réparations sont considérées comme négligeables. Compte tenu du plan d'exploitation optimal, le nombre moyen de défaillances immobilisantes s'élève à 4,97, engendrant de ce fait un coût total de maintenance de 14 910 um.

Mission	Age fonctionnel	Durée	Fiabilité	Nb moyen de pannes durant la mission
1	0	106	0,984	0,016
2	98,79	107	0,903	0,086
3	210,04	82	0,792	0,131
4	291,10	104	0,607	0,267
7	409,75	74	0,469	0,257
8	457,04	64	0,350	0,294
9	495,43	91	0,202	0,551
10	639,86	100	0,100	0,701
13	762,70	60	0,062	0,480
14	825,33	110	0,023	1,012
15	877,76	101	0,007	1,176

Tableau 7 : Détails de l'évolution de la fiabilité et du nombre moyen de pannes pour le plan d'exploitation $X = \{1, 2, 3, 4, 7, 8, 9, 10, 13, 14, 15\}$ soumis à la politique minimaliste

Sous cette politique de maintenance, le gain généré pour ce plan d'exploitation s'élève donc à 42 001 um. (56 911 - 14910). Afin de maximiser ce gain, l'application des maintenances préventives a pour but de minimiser les coûts liés aux actions de maintenances correctives.

Cependant, la durée nécessaire à la réalisation de ce plan d'exploitation s'élève à 999 ut. La différence entre la durée de l'horizon de temps et la durée cumulée des missions du plan d'exploitation représente le temps alloué aux actions de maintenances préventives. Malheureusement, compte tenu du plan d'exploitation optimal déterminé, la durée envisageable est de 1 ut. Cette durée étant inférieure à la durée minimale d'une action de maintenance préventive ($\mu_P(0,1) = 1,04$ ut), aucune politique préventive ne peut être appliquée. À travers cet exemple, nous observons les limites de la stratégie séquentielle dans la mesure où il est très difficile d'optimiser le plan de maintenance. Par conséquent, le bénéfice total obtenu dans le cas de la stratégie séquentielle est de 42 001 um.

La seconde étude abordée dans ce chapitre, à savoir l'étude de la stratégie intégrée, permet de résoudre cette problématique dans la mesure où différents plans d'exploitations seront déterminés

et suivant la politique de maintenance adoptée, les plans d'exploitations pourront tenir compte d'une durée minimale de maintenances préventives.

IV. Seconde étude : la stratégie intégrée

Pour la stratégie séquentielle, la maximisation des gains s'effectuait suivant deux étapes. Dans un premier temps, le plan d'exploitation générant des profits maximum était recherché, puis différentes politiques de maintenances étaient étudiées afin de minimiser les coûts de maintenances associés à ce plan d'exploitation donné. Pour la stratégie intégrée, il ne reste qu'un seul objectif à optimiser, à savoir la maximisation des gains. L'étude de cette stratégie se décomposera suivant les différentes politiques de maintenances proposées durant le chapitre précédent.

IV.1. La politique systématique

Pour cette stratégie intégrée, l'étude de la politique systématique avec un facteur d'amélioration fixé ($\rho = 0,5$) ne pose aucune difficulté. Afin d'être applicable à l'ensemble des plans d'exploitation générés au cours de l'algorithme génétique, il conviendra de s'assurer de la viabilité des plans d'exploitations par le respect de la contrainte (III. 42). La figure suivante montre l'évolution des gains générés en fonction du nombre de générations. Trois répétitions de l'algorithme génétique ont été réalisées, et nous pouvons constater que les solutions obtenues sont relativement proches.

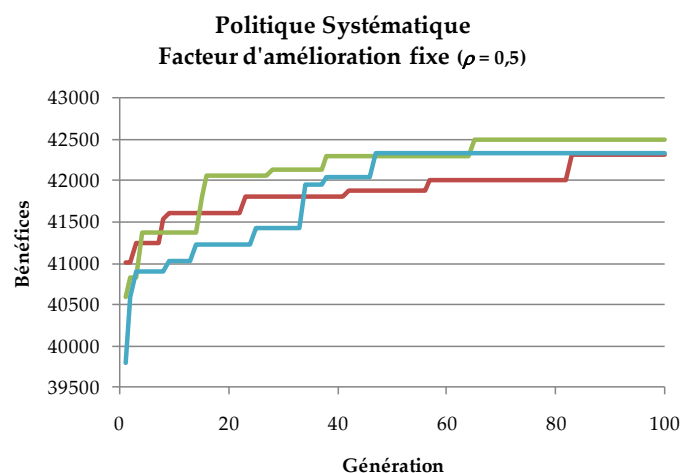


Figure 42 : Évolution des solutions pour la politique systématique avec un facteur d'amélioration fixé

La meilleure solution obtenue pour cette politique de maintenance engendre un bénéfice de 42 492 um pour une durée totale de 998 ut. Cette solution correspond au plan d'exploitation défini par le vecteur $X = \{14, 11, 4, 7, 2, 15, 3, 10, 13, 17\}$ engendrant un profit net de 55 405 um et par conséquent des coûts de maintenances de 12 913 um.

Lorsque le facteur d'amélioration est constant pour l'ensemble des maintenances préventives, mais que ce dernier peut varier entre 0,1 et 0,9 avec un incrément de 0,1, la meilleure solution est obtenue avec un facteur d'amélioration de 0,5. La solution optimale obtenue, après trois répétitions, est proche de celle déterminée juste auparavant. La figure 43 représente les meilleures solutions

obtenues suivant les différentes valeurs prises par le facteur d'amélioration. Pour chaque facteur, trois répétitions ont été réalisées.

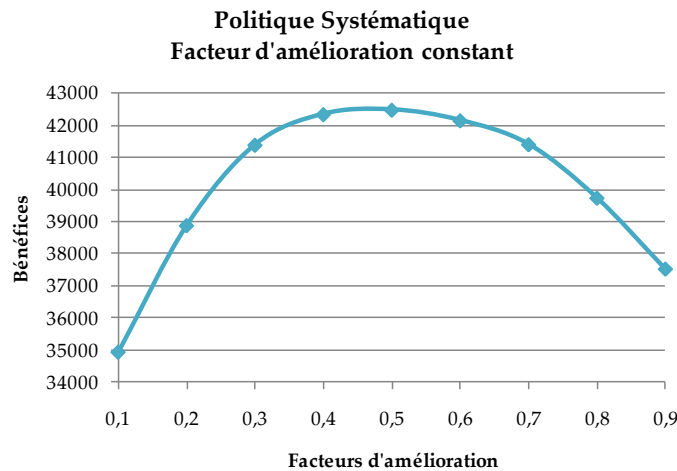


Figure 43 : Évolution des solutions pour la politique systématique suivant les valeurs des facteurs d'amélioration constants

Concernant l'évolution de ces bénéfices, il n'y a pas de grandes différences au niveau des profits générés par les divers plans d'exploitation. Cependant, la différence provient des coûts de maintenances qui varient entre 19 051 um pour le facteur d'amélioration 0,1 et 12 913 um pour un facteur d'amélioration de 0,4. Lorsque les facteurs d'améliorations sont assez faibles, les coûts de maintenances résultent principalement des actions correctives, alors que dans la situation opposée, les coûts sont liés à des maintenances préventives qui ne sont effectivement pas toutes nécessaires.

L'utilisation de facteurs d'amélioration variables suivant les actions préventives devrait permettre de remédier à ce problème dans la mesure où, au départ et à la fin du plan d'exploitation, il n'est pas nécessairement utile de réaliser des maintenances avec un facteur d'amélioration élevé.

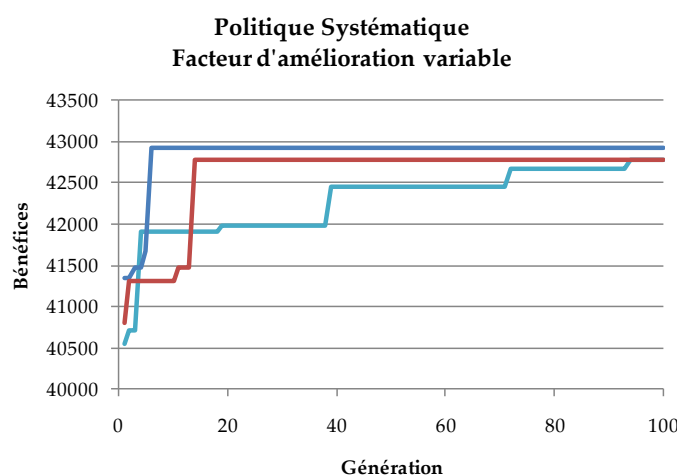


Figure 44 : Évolution des solutions pour la politique systématique avec des facteurs d'amélioration variables

Sur la figure 44, nous pouvons constater que les solutions obtenues génèrent un bénéfice supérieur à l'utilisation d'un facteur constant au cours du temps. La meilleure solution est liée au plan

d'exploitation défini par $\mathcal{X} = \{11, 17, 15, 14, 2, 13, 3, 7, 10, 1\}$ dont le profit et la durée de réalisation s'élèvent respectivement à 55 575 um et 982 ut. Les actions préventives, dont les facteurs d'améliorations après chaque mission sont donnés par $P = \{0,3, 0,6, 0,3, 0,9, 0,3, 0,2, 0,7, 0,4, 0,2\}$ induisent un coût de maintenance de 12647 um. La politique sporadique devrait permettre de minimiser les coûts de maintenances, dans la mesure où l'on recherche les missions devant être suivies d'actions préventives.

IV.2. La politique sporadique

Lorsque le facteur d'amélioration est fixé ($\rho = 0,5$), le bénéfice engendré par cette politique est en moyenne supérieur de 2 500 um par rapport à la politique systématique suivant les mêmes critères. Après les trois répétitions, le bénéfice maximal (cf. figure 45) est issu de la réalisation du plan d'exploitation défini par $\mathcal{X} = \{5, 11, 7, 3, 13, 10, 17, 15, 14, 2\}$ (Profit : 55 640 um) pour lequel les coûts de maintenance s'élèvent à 10 395 um. L'obtention de ces coûts relève de l'exécution de maintenances préventives après les missions 11, 3, 13 et 10.

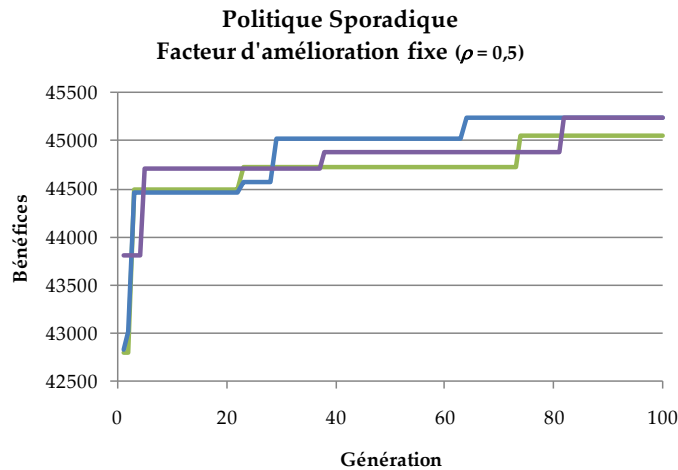


Figure 45 : Évolution des solutions pour la politique sporadique avec un facteur d'amélioration fixé

Lorsque nous observons les bénéfices maximum obtenus suivant les valeurs prises par le facteur d'amélioration, la courbe représentative n'est pas convexe, car le nombre de maintenances préventives est variable. Cependant, nous observons que les meilleurs bénéfices sont obtenus lorsque le facteur d'amélioration prend les valeurs de 0,8 ou 0,9.

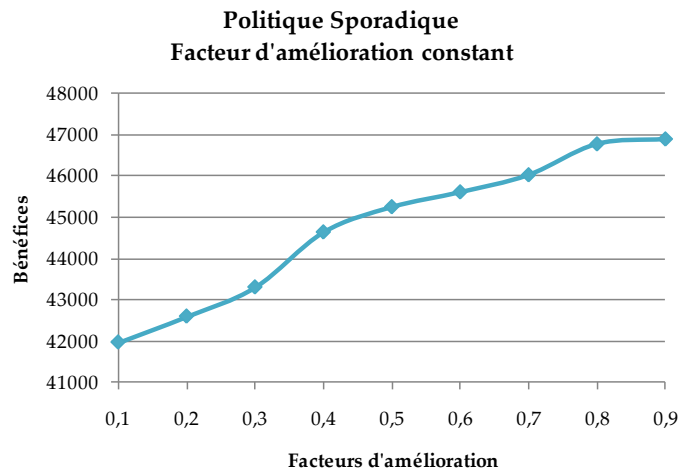


Figure 46 : Évolution des solutions pour la politique sporadique suivant les valeurs des facteurs d'amélioration constants

En recherchant les solutions optimales lorsque le facteur d'amélioration est constant, deux solutions, parmi les trois répétitions, avaient un facteur d'amélioration égal à 0,9 et une dont le facteur d'amélioration s'élevait à 0,8. La meilleure solution fut obtenue avec un facteur d'amélioration de 0,9 et pour le plan d'exploitation défini par $\mathcal{X} = \{13, 14, 15, 8, 10, 7, 2, 3, 11, 17\}$ (Profit : 54 966 um) pour lequel deux actions préventives ont été réalisées après les missions 8 et après la mission 7. Les coûts de maintenances engendrés sont de 8 078 um.

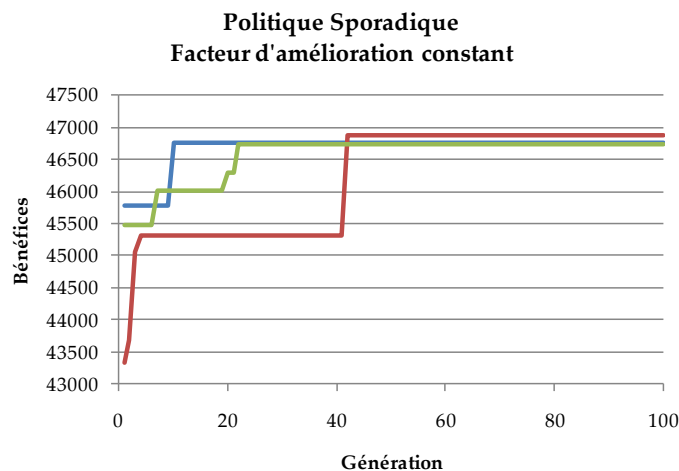


Figure 47 : Évolution des solutions pour la politique sporadique avec un facteur d'amélioration constant

Lorsque les facteurs d'améliorations peuvent être variables suivant chaque facteur d'amélioration, nous ne pouvons pas noter une grande différence au niveau des bénéfices générés. En effet, parmi les trois répétitions réalisées, la meilleure solution fut obtenue avec la réalisation de deux maintenances préventives, dont les facteurs d'amélioration étaient de 0,9, après les missions 7 et 3 du plan d'exploitation $\mathcal{X} = \{2, 13, 10, 7, 14, 3, 11, 5, 15, 17\}$. Les bénéfices s'élèvent à 47 073 um, issus des profits (55 640 um) minorés des coûts actions préventives (8567 um). Pour les deux autres répétitions, deux maintenances préventives étaient réalisées, avec des facteurs d'amélioration de 0,8 et 0,9.

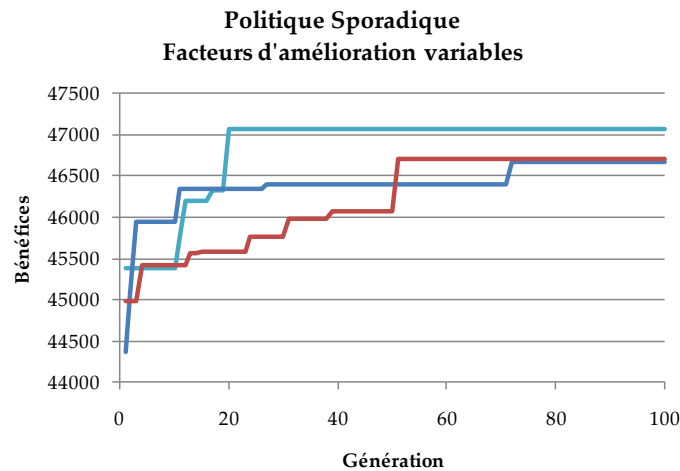


Figure 48 : Évolution des solutions pour la politique sporadique avec des facteurs d'amélioration variables

IV.3. La politique périodique

Concernant la politique périodique, lorsque le facteur d'amélioration est fixé ($\rho = 0,5$), une solution satisfaisante est obtenue avec le plan d'exploitation défini par $X = \{15, 9, 11, 14, 17, 13, 7, 2, 10, 3\}$. Ce plan d'exploitation génère un profit de 55 667 um. Les bénéfices générés s'élèvent à 46 729 um, soit un coût de maintenance de 8 938 um. Ces coûts de maintenance sont issus de la réalisation de 4 maintenances préventives. Pour les trois répétitions, les solutions finales étaient toutes constituées de 4 maintenances préventives.

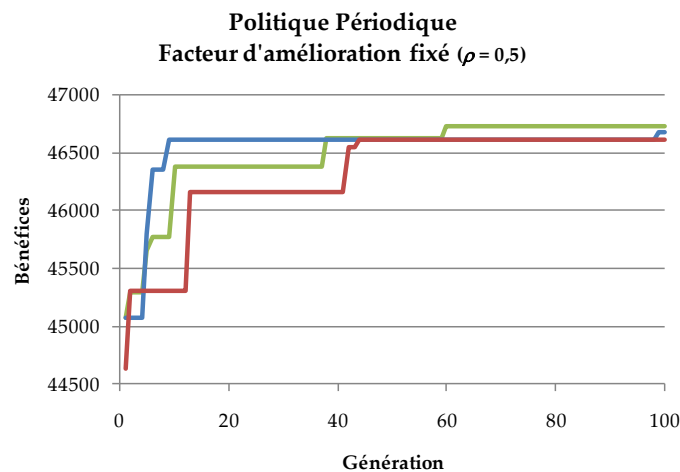


Figure 49 : Évolution des solutions pour la politique périodique avec un facteur d'amélioration fixé

Lorsque le facteur d'amélioration doit être identique pour chaque action préventive, mais que la valeur de ce dernier peut varier entre 0,1 et 0,9, il apparaît que les bénéfices augmentent avec le facteur d'amélioration. Cependant, l'augmentation des facteurs n'implique pas un nombre constant de maintenances préventives, comme l'illustre la figure suivante.

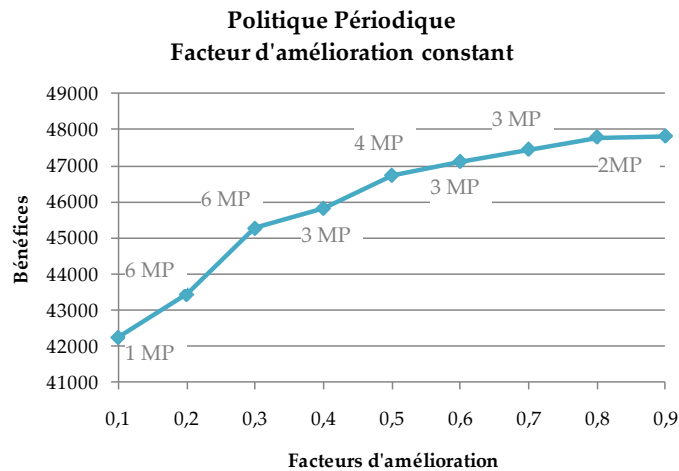


Figure 50 : Évolution des solutions pour la politique périodique suivant les valeurs des facteurs d'amélioration constants

De manière évidente, lorsque le facteur d'amélioration optimal est recherché, les solutions satisfaisantes sont obtenues lorsque le facteur d'amélioration est égal à 0,9. Les trois répétitions de l'algorithme génétique ont renvoyé des résultats proches, pour lesquels les facteurs d'amélioration étaient toujours de 0,9. Pour la meilleure solution, le bénéfice est de 47 816 um avec la réalisation de deux maintenances préventives durant le plan d'exploitation $X = \{14, 17, 8, 13, 15, 11, 7, 2, 3, 10\}$.

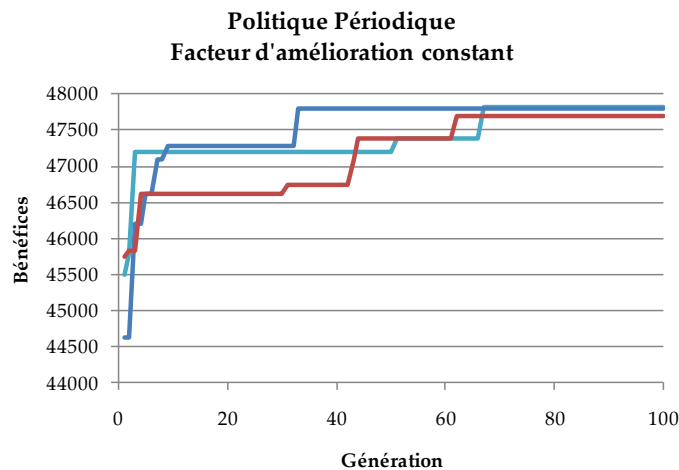


Figure 51 : Évolution des solutions pour la politique périodique avec un facteur d'amélioration constant

D'après les différents essais réalisés avec la politique périodique et les facteurs d'améliorations variables, les meilleures solutions renvoyées possèdent les mêmes facteurs d'améliorations, à savoir la valeur de 0,9. Les résultats obtenus sont donc similaires à ceux présentés sur la figure 51. Comparés à la politique sporadique, les résultats obtenus avec le facteur d'amélioration fixé ($\rho = 0,5$) sont très proches. Cependant, lorsque la valeur optimale du facteur d'amélioration est déterminée, les solutions issues de la politique périodique sont légèrement meilleures.

Comparée à la politique périodique, la politique séquentielle consiste à déterminer, en plus du nombre de maintenances préventives, les intervalles de temps optimaux. Par conséquent, il semble logique que les résultats issus de cette politique séquentielle soient encore supérieurs.

IV.4. La politique séquentielle

La figure 52 illustre les bénéfices générés suivant la politique séquentielle pour le facteur d'amélioration fixé. Les résultats obtenus sont très proches et avoisinent des bénéfices de l'ordre de 47 000 um. La meilleure solution est dérivée du plan d'exploitation $X = \{11, 15, 2, 14, 10, 13, 7, 17, 3, 9\}$ dont les profits s'élèvent à 55 667 um. Les coûts de maintenances de 8 593 um découlent de la réalisation de quatre actions préventives respectivement réalisées au bout de 363 ut, puis 80 ut, 129 ut et 122 ut.

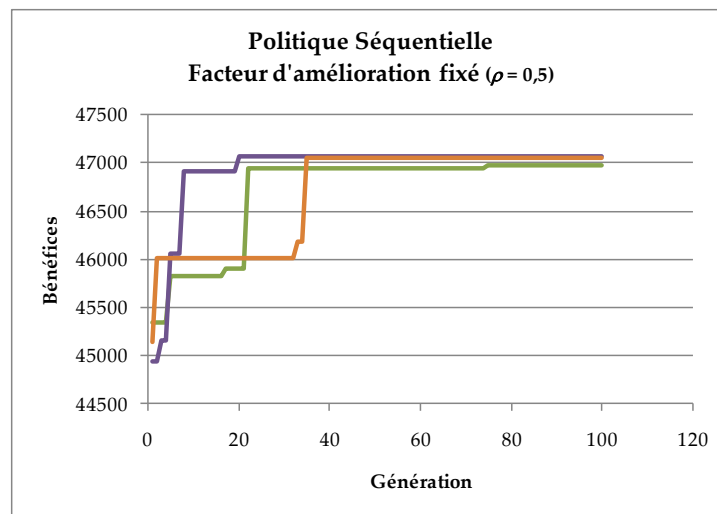


Figure 52 : Évolution des solutions pour la politique séquentielle avec un facteur d'amélioration fixé

À l'instar des deux précédentes politiques, il apparaît clairement, sur le graphique suivant, que les bénéfices générés sont croissants avec l'augmentation du facteur d'amélioration.

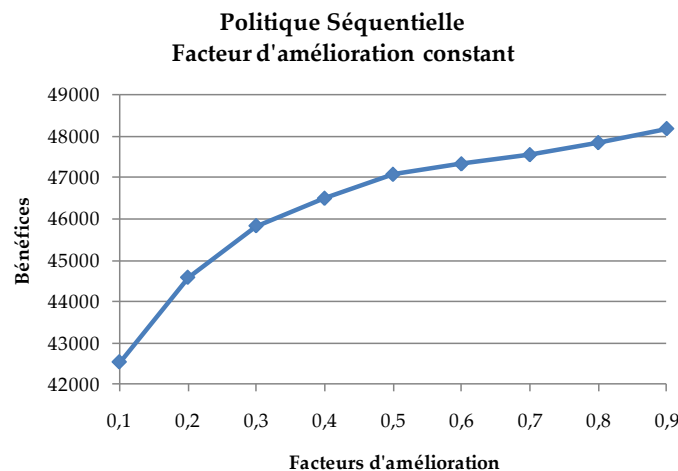


Figure 53 : Évolution des solutions pour la politique séquentielle suivant les valeurs des facteurs d'amélioration constants

Sans grande surprise, lorsque la valeur du facteur d'amélioration optimale doit être déterminée (le facteur étant commun à l'ensemble de toutes les actions préventives), les bénéfices maximums sont obtenus pour le facteur d'amélioration égal à 0,9. La meilleure solution obtenue génère un bénéfice de 48 182 um avec l'accomplissement du plan d'exploitation $\chi = \{15, 2, 7, 11, 17, 3, 9, 13, 10, 14\}$ et d'actions de maintenances préventives après 356 ut puis 313 ut. Les deux autres solutions satisfaisantes sont obtenues avec la réalisation de deux maintenances préventives dont les valeurs des facteurs d'améliorations étaient de 0,9 et 0,8 pour chacun des deux cas.

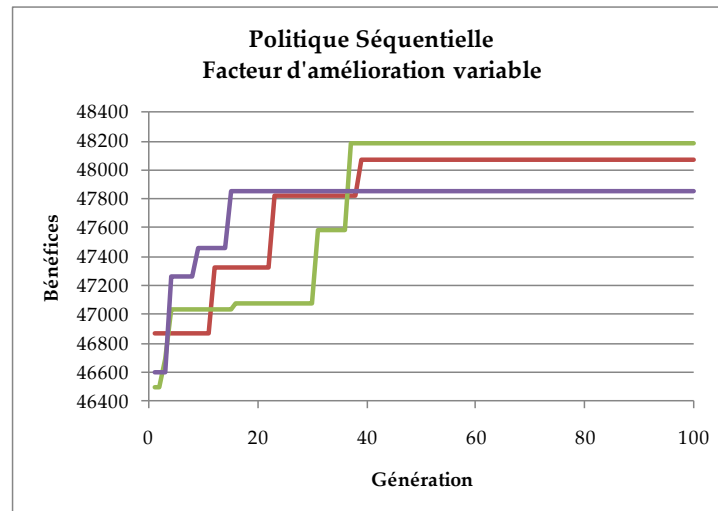


Figure 54 : Évolution des solutions pour la politique séquentielle avec des facteurs d'amélioration variables

V. Synthèse au sujet des stratégies séquentielle et intégrée

La stratégie séquentielle a certes l'avantage de renvoyer le meilleur plan d'exploitation pour la maximisation des profits générés, mais l'application de politiques de maintenances préventives peut s'avérer très difficile dans la mesure où la durée n'est pas suffisante. Par ailleurs, cette stratégie néglige totalement les conditions opérationnelles et environnementales des missions, selon lesquelles les coûts de maintenances peuvent varier, ainsi que l'ordonnancement des missions. Ces différents problèmes sont résolus en adoptant la stratégie intégrée dans la mesure où ces différents paramètres interviennent au niveau de l'optimisation globale.

Pour la stratégie intégrée, nous pouvons constater, que pour chaque politique de maintenance étudiée, les bénéfices étaient maximums dans le cas où la valeur des facteurs d'améliorations n'était pas fixée. Mise à part la stratégie systématique, pour laquelle le facteur d'amélioration optimal (constant pour chaque action préventive) était égal à 0,5, pour les autres politiques, les bénéfices étaient obtenus dans le cas où le facteur d'amélioration était maximum. Par conséquent, les différences concernant les bénéfices obtenus lorsque les facteurs d'amélioration étaient constants ou variables sont négligeables.

		Facteurs d'amélioration		
		Fixés	Constants	Variables
Politiques	Systématique	42 492 um	42 492 um	42 928 um
	Sporadique	45 245 um	46 888 um	47 073 um
	Périodique	46 729 um	47 816 um	47 816 um
	Séquentielle	47 074 um	48 182 um	48 182 um

Tableau 8 : Bénéfices obtenus pour la stratégie intégrée

Par ailleurs, nous pouvons remarquer que la réalisation de maintenances préventives systématiques génère des bénéfices assez faibles dans la mesure où les actions préventives (pas forcément utiles) monopolisent du temps qui aurait pu être accordé à la réalisation d'autres missions. D'autre part, nous constatons que la politique séquentielle génère des profits maximums dans la mesure où les durées des intervalles de maintenances préventives sont optimales.

Conclusion

DANS ce mémoire, nous nous sommes intéressés au problème de l'intégration entre des activités de production et des activités de maintenance. Concernant notre système d'étude, à savoir un navire, il est délicat de parler d'activités de production. Plus précisément, ces activités consistaient, dans un premier temps, à déterminer les missions qui devaient être réalisées, puis à trouver l'ordonnancement optimal des dites missions sélectionnées. L'étude d'un tel système diverge de l'étude des systèmes de production classiques dans la mesure où, pendant que le navire est en mer, aucun retour à quai n'est possible. De ce fait, il convient de déterminer, avant le départ, le nombre de pièces nécessaires aux différentes activités de maintenance. Par ailleurs, pendant la durée en mer, le navire est amené à accomplir plusieurs missions dont les lieux de réalisation varient, tout comme les conditions environnementales associées (température, salinité, etc.). En plus de ces différents lieux, les missions réalisées ne sont pas toutes identiques. Par exemple, les conditions opérationnelles associées à des missions de démonstration de force impactent plus rapidement la dégradation du système que de simples missions de surveillances côtières.

Nous avons donc commencé cette thèse par une étude de la littérature du point de vue des activités de maintenances, ainsi que des différentes politiques de maintenances existantes. Nous avons pu montrer l'existence de lacunes par rapport aux différentes spécificités de notre étude. En effet, l'étude des horizons de temps fini est récente et, par conséquent, peu de travaux existent. De manière similaire, très peu de travaux abordent la problématique des conditions de fonctionnement variables au cours du temps. Par la suite, nous avons étudié les différentes stratégies (séquentielle et intégrée) permettant le couplage des activités de production et de maintenance.

Dans le deuxième chapitre, nous nous sommes concentrés sur la modélisation de la loi de défaillance afin de la rendre adaptable aux différentes missions, c'est-à-dire aux différentes caractéristiques de fonctionnement. Après l'étude des deux principes de modélisation existants, notre choix s'est porté sur le modèle à risques proportionnels pour sa bonne adéquation à notre problématique. Néanmoins, ce modèle suppose que les conditions de fonctionnement sont constantes entre deux défaillances immobilisantes. Cette hypothèse de travail ne peut, dans notre cas, pas être vérifiée. En effet, différents types de mission peuvent être réalisés dans des lieux distincts entre deux défaillances consécutives. À l'aide d'un raisonnement basé sur la logique floue, nous avons proposé une méthode permettant de représenter l'ensemble des missions réalisées par une seule mission, constituée de conditions de fonctionnement équivalentes, représentative.

Le chapitre suivant concerne l'étude des deux stratégies permettant le couplage entre les plans d'exploitation et de maintenance. La première stratégie, à savoir séquentielle, se caractérise par l'optimisation du plan d'exploitation, en vue de maximiser les profits générés par l'accomplissement des missions, puis de considérer cet ordonnancement de missions comme contrainte forte pour la détermination du plan de maintenance, qui cherche à minimiser les coûts de maintenance. À l'inverse, la stratégie intégrée vise à optimiser simultanément les plans d'exploitation et de maintenance en maximisant la fonction objective : la maximisation des bénéfices engendrés. L'optimisation des bénéfices générés pour la stratégie intégrée est réalisée à l'aide d'une méta-heuristique pour les choix optimaux des plans d'exploitation et de maintenance. Concernant les plans de maintenances, différentes politiques ont été proposées au cours de ce chapitre. La première

politique de maintenance est considérée comme minimaliste dans la mesure où elle se caractérise par la réalisation exclusive d'actions correctives. À travers cette politique, la notion d'âge fonctionnel a été introduite afin de garantir la continuité de la fiabilité du système entre les différentes missions. En effet, il est illogique de penser que, suite à des modifications des conditions de fonctionnement, la fiabilité du système peut être améliorée ou au contraire dégradée. Afin de minimiser les coûts de maintenance liés à la politique minimaliste, différentes politiques de maintenances préventives ont été proposées. Les deux premières politiques présentées, systématique et sporadique, se caractérisent par l'accomplissement des actions préventives à la fin des missions. Comme le laisse sous-entendre son nom, la politique systématique se définit par des actions préventives exécutées à la fin de chaque mission. La politique sporadique se base sur la même idée, à savoir des maintenances préventives à la fin des missions, mis à part que le choix des missions suivies des actions préventives sera déterminé en vue de minimiser les coûts de maintenances en ne réalisant que les actions réellement utiles. Les deux autres politiques abordées au cours de ce chapitre se différencient par les instants des actions préventives, qui peuvent intervenir indifféremment à la fin des missions ou pendant leurs exécutions. Ainsi, la politique périodique se caractérise par la réalisation d'actions préventives à intervalles de temps constants alors que la politique séquentielle vise à déterminer, en plus du nombre optimal de maintenances préventives, les durées optimales des intervalles de maintenances préventives. Lorsque le facteur d'amélioration, qui sert à modéliser l'efficacité des maintenances préventives, est fixé, la minimisation des coûts de maintenances pour les politiques sporadique et séquentielle est réalisée à l'aide de méta-heuristiques. La détermination des facteurs optimaux pour chaque action préventive et ce, pour les différentes politiques de maintenances préventives, entraîne l'utilisation de méta-heuristiques du type algorithme génétique ou algorithme du grand déluge. Par conséquent, la maximisation des bénéfices, dans le cas de la stratégie intégrée amène à l'utilisation imbriquée de deux méta-heuristiques.

Le dernier chapitre de cette thèse est consacré à l'illustration des différentes stratégies et politiques de maintenances préventives à travers un exemple numérique. L'exemple présenté démontre les inconvénients liés à la stratégie séquentielle, et ce notamment pour le choix des politiques de maintenances préventives. Au contraire, la stratégie intégrée consent à l'utilisation de toutes les politiques de maintenances et permet donc de démontrer l'efficacité de chacune d'entre elles. Il apparaît, de façon évidente, que la politique la plus adaptée à cette problématique s'avère être la politique séquentielle avec les facteurs d'améliorations variables, dans la mesure où ces facteurs d'amélioration, le nombre de maintenances préventives ainsi que les durées des intervalles de maintenances préventives sont autant de variables de décisions qui permettent d'obtenir une solution satisfaisante, voire optimale.

Au-delà de l'application au domaine naval, ces contributions peuvent être appliquées à d'autres domaines. Par exemple, dans le domaine aéronautique, il serait possible de définir des politiques d'inspection, utiles pour vérifier l'état des sondes Pitot, qui pour mémoire, sont à l'origine du crash du vol AF 447 (RIO-PARIS) ayant provoqué la mort de 228 personnes le 1^{er} juin dernier. Par rapport à la problématique étudiée dans ce mémoire, l'horizon de temps est également considéré comme fini (durée du vol) et les conditions de fonctionnement varient suivant les conditions météorologiques (température, pluie, etc.) ainsi que les conditions de navigation (altitude, etc.). Alors que pour l'étude présentée dans ce mémoire, le plan de maintenance était proposé avant le départ du navire, dans le cas des sondes Pitot, l'outil d'aide à la décision devrait servir en temps réel et ce, en fonction des changements climatiques qui pourraient intervenir.

Bibliographie

- [AFNOR, 2001] AFNOR (2001). Norme européenne. *Terminologie de la maintenance* (NF EN 13306).
- [Aghezzaf, 2007] Aghezza, E.H., M.A. Jamali and D. Ait-Kadi (2007). An integrated production and preventive maintenance planning model. *European Journal of Operational Research*, 181(2), pp. 679–685.
- [Arango, 2006] Arango, G., S. Hennequin and N. Rezg (2006). Fuzzy modeling and optimization of imperfect maintenance. *Proceedings of International Conference on Industrial Engineering and Systems Management (IESM'07)*, Beijing, China.
- [Bagdonavičius, 1994] Bagdonavičius V. and M. Nikulin (1994). Accelerated Life Models: Modeling and Statistical Analysis. CRC Press Inc.
- [Baker, 1987] Baker, J.E. (1987). Reducing bias and inefficiency in the selection algorithm. *Proceedings of the Second International Conference on Genetic Algorithms*, pp. 14–21.
- [Balas, 2008] Balas, E., N Simonetti and A Vazacopoulos (2008). Job shop scheduling with setup times, deadlines and precedence constraints. *Journal of Scheduling*, 11(4), pp. 253–262.
- [Barlow, 1960] Barlow and Hunter, 1960 R. Barlow and L.C. Hunter, Optimum preventive maintenance policies, *Operations Research* 8 (1960) (1), pp. 90–100.
- [Barlow, 1965] Barlow, R.E. and F. Proschan (1965). *Mathematical theory of reliability*. John Wiley & Sons.
- [Benbouzid, 2003] Benbouzid, F., Y. Bessadi, S.A. Guebli, C. Varnier and N. Zerhouni (2003). Résolution du problème de l'ordonnancement conjoint maintenance/production par la stratégie séquentielle. *Proceedings of the 4^{ème} Conférence Francophone de MOdélisation et de SIMulation (MOSIM'03)*, Toulouse, France.
- [Benbouzid-Sitayeb, 2006] Benbouzid-Sitayeb, F., C. Varnier and N. Zerhouni (2006). Résolution du problème de l'ordonnancement conjoint production/maintenance par colonies de fourmis. *Proceedings of the 6^{ème} Conférence Francophone de MOdélisation et SIMulation (MOSIM'06)*, Rabat, Maroc.
- [Bennour, 2001] Bennour, M., C.Bloch and N. Zerhouni (2001). Modélisation intégrée des activités de maintenance et de production. *Proceedings of the 3^{ème} Conférence Francophone de MOdélisation et de SIMulation (MOSIM'01)*, Troyes, France.

- [Brown, 1983] Brown, M. and F. Proschan (1983). Imperfect repair. *Journal of Applied Probability*, 20, pp. 851–859.
- [Brucker, 1996] Brucker P., O. Thiele (1996). A branch-and-bound method for general shop problem with sequence-dependent setup times. *OR Spectrum*, 18(3), pp. 145-161.
- [Bühler, 1994] Bühler, H. (1994). Réglage par logique floue. *Presses polytechniques et universitaires romandes*.
- [Burke, 2004] Burke, E., Y. Bykov, J. Newall and S. Petrovic (2004). A time-predefined local search approach to exam timetabling problems, *IIE Transactions*, 36(6), pp. 509–528.
- [Byington, 2002] Byington, C., M. Roemer, G. Kacprzyński and T. Galie (2002). Prognostic Enhancements to Diagnostic Systems for Improved Condition-Based Maintenance. *Proceedings of the IEEE Aerospace Conference*, Big Sky, USA.
- [Byington, 2003] Byington, C., M. Watson, M. Roemer and T. Galie (2003). Prognostic Enhancements to Gas turbine Diagnostic Systems. *Proceedings of the IEEE Aerospace Conference*, Big Sky, USA.
- [Cassady, 2003] Cassady, C.R. and Kutanoğlu, E. (2003). Minimizing job tardiness using integrated preventive maintenance planning and production scheduling. *IIE Transactions*, 35(6), pp. 503–513.
- [Chan, 1993] Chan, J.K. and L. Shaw (1993). Modeling repairable systems with failure rates that depend on age and maintenance. *IEEE Transactions on Reliability*, 42(4), pp. 566–571.
- [Cheng, 2007] Cheng, C.Y., M. Chen and R. Guo (2007). The optimal periodic preventive maintenance policy with degradation rate reduction under reliability limit. *IEEE International Conference on Industrial Engineering and Engineering Management (IEEE'07)*, pp. 649–653, Singapore.
- [Cheung, 1997] Cheung, K.L. and W.H. Hausmann (1997). Joint optimization of preventive maintenance and safety stock in an unreliable production environment. *Naval Research Logistics*, 44, pp. 257–272.
- [Chung, 2009] Chung, S.H., F.T.S. Chan and H.K. Chan (2009). A modified genetic algorithm approach for scheduling of perfect maintenance in distributed production scheduling. *Engineering Applications of Artificial Intelligence*, 22(7), pp. 1005–1014.
- [Cocheteux, 2009] Cocheteux P., A. Voisin, E. Levrat and B. Iung (2009). Méthodologie pour l'évaluation de perte de performance système dans le cadre d'une maintenance prévisionnelle. *3^{ème} Journées Doctorales / Journées Nationales MACS (JD-JN-MACS 2009)*, Angers, France

- [Coleman, 1992] Coleman, B.J. (1992). Technical note: a simple model for optimizing the single machine early/tardy problem with sequence-dependent setups. *Production and Operations Management*, 1(2), pp. 225–228.
- [Cox, 1972] Cox, D. (1972). Regression Models and Life-Tables. *Journal of the Royal Statistical Society*, 34(2), pp. 187–220.
- [Dahane, 2007] Dahane, M. (2007). Couplage de la gestion de la maintenance et la gestion de la production sous contrainte de sous-traitance. *Thèse de doctorat*. Université Paul Verlaine – Metz, France.
- [Dedopoulos, 1998] Dedopoulos, I.T. and Y. Smeers (1998). An age reduction approach for finite horizon optimization of preventive maintenance for single units subject to random failures. *Computers and Industrial Engineering*, 34(3), pp. 643–654.
- [Dekker, 1996] Dekker R. (1996). Application of maintenance optimization models: a review and analysis, *Reliability Engineering and System Safety*, 51(3), pp. 229–240.
- [Dellagi, 2006] Dellagi, S. (2006). Développement de stratégies de maintenance dans un contexte de sous-traitance partielle de production. *Thèse de doctorat*. Université Paul Verlaine – Metz, France.
- [Doyen, 2004A] Doyen, L. and O. Gaudoin (2004). Classes of imperfect repair models based on reduction of failure intensity or virtual age. *Reliability Engineering and System Safety*, 84(1), pp. 45–56.
- [Doyen, 2004B] Doyen, L. and O. Gaudoin (2004). Modélisation de l'efficacité de la maintenance des systèmes réparables – Synthèse bibliographique. *Rapport du contrat T50L47/F00555/0 entre EDF et le LMC*.
- [Godjevac, 1999] Godjevac, J. (1999). Idées nettes sur la logique floue. Presse polytechniques et universitaires romandes.
- [Goldberg, 1989] Goldberg, D.E. (1989). Genetic algorithms in search, optimization and machine learning. *Addison Wesley*.
- [Goldratt, 2006] Goldratt, E.M. and J. Cox (2006). Le but : un processus de progrès permanent. *AFNOR* (3^{ème} édition).
- [Hajjej, 2009] Hajjej, Z., S. Dellagi and N. Rezg. (2009). An optimal production/maintenance planning under stochastic random demand, service level and failure rate. *IEEE International Conference on Automation Science and Engineering (CASE 2009)*, pp. 292 –297, Bangalore.
- [Holland, 1975] Holland, J.H. (1975). Adaptation in Natural and Artificial Systems. *University of Michigan Press, Ann Arbor*.

- [Iung, 2005] Iung, B., M. Vérona, M.C. Suhnera and A. Muller (2005). Integration of Maintenance Strategies into Prognosis Process to Decision-Making Aid on System Operation. *CIRP Annals - Manufacturing Technology*, 54(1), pp. 5–8.
- [Jayabalan, 1992A] Jayabalan, V. and D. Chaudhuri (1992). Optimal maintenance—replacement policy under imperfect maintenance. *Reliability Engineering & System Safety*, 36(2), pp. 165–169.
- [Jayabalan, 1992B] Jayabalan, V. and D. Chaudhuri (1992). An heuristic approach for finite time maintenance policy. *International Journal of Production Economics*, 27(3), pp. 251–256.
- [Ji, 2007] Ji, M., Y. He, T.C.E. Cheng (2007). Single-machine scheduling with periodic maintenance to minimize makespan. *Computers and Operations Research*, 34(6), pp. 1764–1770.
- [Jiang, 2009] Jiang, R. (2009). An accurate approximate solution of optimal sequential age replacement policy for a finite-time horizon. *Reliability Engineering and System Safety*, 94, pp. 1245–1250.
- [Johnson, 1954] Johnson, S.M. (1954). Optimal two- and three-stage production schedules with setup times included. *Naval Research Logistics Quarterly*, 1, pp. 61–68.
- [Karp, 1972] R. M. Karp (1972). Reducibility Among Combinatorial Problems. *Complexity of Computer Computations*, Plenum Press, New York, pp. 85–103.
- [Kenné, 2004] Kenné, J.P. and A. Gharbi (2004). Stochastic optimal production control problem with corrective maintenance. *Computers and Industrial Engineering*, 46, pp. 865–875.
- [Kijima, 1988] Kijima, M., H. Morimura and Y. Suzuki (1988). Periodical replacement problem without assuming minimal repair. *European Journal of Operational Research*, 37, pp. 194–203.
- [Kijima, 1989] Kijima, M. (1989). Some results for repairable systems with general repair. *Journal of Applied Probability*, 26, pp. 89–102.
- [Lebold, 2001] Lebold, M. and M. Thurston (2001). Open Standards for Condition-Based Maintenance and Prognostic Systems. *Proceedings of the 5th Annual Maintenance and Reliability Conference (MARCON 2001)*, Gatlinburg, USA.
- [Legát, 1996] Legát, V., A. H. Žaludová, V. Červenka and V. Jurča (1996). Contribution to optimization of preventive replacement. *Reliability Engineering and System Safety*, 51(3), pp. 259–266.

- [Léger, 2003] Léger, J.-B. (2003). La Sûreté Industrielle, la Maintenance Prévisionnelle et les NTIC : innovation du progiciel CASIP et perspectives de développement. 1ère édition du colloque international francophone « PERformances et Nouvelles TechnolOgies en Maintenance » (PENTOM'03), Valenciennes, France.
- [Lin, 2000] Lin, D., M. Zuo and R.C.M. Yam (2000). General sequential imperfect preventive maintenance models. *International Journal of Reliability, Quality and Safety Engineering*, 7(3), pp. 253–266.
- [Liu, 2008] Liu, H.H. (2008). Optimal Periodic Preventive Maintenance Policy(with Warranty for Finite Time Span. *Mémoire de Master*. Department and Graduate Institute of Industrial Engineering and Management, Thailand
- [Lyonnet, 2000] Lyonnet, P. (2000). La maintenance : mathématiques et méthodes. Edition Tec et Doc, Paris.
- [Lyonnet, 2006] Lyonnet, P. (2006). Ingénierie de la fiabilité. Edition Tec et Doc, Paris.
- [Malik, 1979] Malik, M.A.K. (1979). Reliable preventive maintenance scheduling. *AIIE Transactions*, 11(3), pp. 221–228.
- [Martello, 1990] Martello, S. et P. Toth (1990). Knapsack Problems: Algorithms and Computer Implementations, *John Wiley & Sons*.
- [Martorell, 1999] Martorell, S., A. Sanchez and S. Vicente, (1999). Age-dependent reliability model considering effects of maintenance and working conditions. *Reliability Engineering and System Safety*, 64(1), pp. 19–31.
- [McMullan, 2007] McMullan, P. and B. McCollum (2007). Dynamic Job Scheduling on the Grid Environment Using the Great Deluge Algorithm. *Parallel Computing Technologies*, pp. 283–292.
- [Muller, 2005] Muller, A. (2005). Contribution à la maintenance prévisionnelle des systems de production par la formalization d'un processus de prognostic. *Thèse de doctorat*. Université Henri Pointcaré, Nancy.
- [Naderi, 2008] Naderi, B., M. Zandieh and S.M.T. Fatemi Ghomi (2008). A study on integrating sequence dependent setup time flexible flow lines and preventive maintenance scheduling, *Journal of Intelligent Manufacturing*. (doi: 10.1007/s10845-008-0157-6).
- [Nakagawa, 1979] Nakagawa, T. (1979). Optimum policies when preventive maintenance is imperfect. *IEEE Transactions on Reliability*, 28(4), pp. 331–332.
- [Nakagawa, 1986] Nakagawa, T. (1986). Periodic and Sequential Preventive Maintenance Policies. *Journal of Applied Probability*, 23(2), pp. 536–542.
- [Nakagawa, 1988] Nakagawa, T. (1988). Sequential Imperfect Preventive Maintenance Policies. *IEEE Transactions on Reliability*, 37(3), pp. 295–298.

- [Nakagawa, 2005] Nakagawa, T. (2005). Maintenance Theory of Reliability. Springer, London.
- [Nakagawa, 2009] Nakagawa, T. and S. Mizutani (2009). A summary of maintenance policies for a finite interval. *Reliability Engineering and System Safety*, 94(1), pp. 89–96.
- [Nawaz, 1983] Nawaz, M., E.E. Enscore Jr. and I. Ham (1983). A heuristic algorithm for the m-machine, n-job flow-shop sequencing problem. *International Journal of Management Science*, 11 (1), pp. 91–95.
- [Osman, 1989] Osman, I.H. and C.N. Potts (1989). Simulated annealing for permutation flow-shop scheduling, OMEGA. *The International Journal of Management Science* 17, 6, pp. 551–557.
- [Özekici, 1995] Özekici, S. (1995). Optimal maintenance policies in random environments. *European Journal of Operational Research*, 82(2), pp. 283–294.
- [Pham, 1996] Pham, H. and H. Wang (1996). Imperfect maintenance. *European Journal of Operational Research*, 94(3), pp. 425–438.
- [Pinedo, 2005] Pinedo, M.L. (2005). Planning and Scheduling in Manufacturing and Services. Springer Series in Operations Research and Financial Engineering. Springer, Berlin Heidelberg New York.
- [Radhoui, 2008] Radhoui, M. (2008). Analyse des performances de systems de production sujets à des défaillances aléatoires et pouvant engendrer des produits non conformes dans un environnement incertain. *Thèse de doctorat*. niversité Paul Verlaine – Metz, France.
- [Rajendran, 2004] Rajendran, C. and H. Ziegler (2004). Ant-colony algorithms for permutation flow-shop scheduling to minimize makespan/total flowtime of jobs. *European Journal of Operational Research*, 155, pp. 426–438.
- [Rezg, 2004] Rezg, N., X. XIE, Y. Mati (2004). Joint optimization of preventive maintenance and inventory control in a production line using simulation. *International Journal of Production Research*, 42(10), pp. 2029–2046.
- [Ross, 2004] Ross, T.J. (2004). Fuzzy Logic with Engineering Applications. John Wiley & Sons, Ltd.
- [Ruiz, 2005] Ruiz, R. and C. Maroto (2005). A comprehensive review and evaluation of permutation flowshop heuristics. *European Journal of Operational Research*, 165, pp. 474–494.
- [Ruiz, 2006] Ruiz, R., C. Maroto and J. Alcaraz (2006). Two new robust genetic algorithms for the flow-shop scheduling problem. *International Journal of Management Science*, 3, pp. 461–476.

- [Ruiz, 2007] Ruiz, R., J.C. García-Díaz and C. Maroto (2007). Considering scheduling and preventive maintenance in the flowshop sequencing problem. *Computers & Operations Research*, 34(11), pp. 3314–3330.
- [Schutz, 2008A] Schutz, J., N Rezg and J.B. Léger (2008). Contribution au développement d'un plan d'exploitation et de maintenance basé sur une loi de dégradation évolutive. *Proceedings of the 7^{ème} Conférence Internationale de Modélisation et Simulation (MOSIM'08)*, pp. 1841–1850, Paris, France.
- [Schutz, 2008B] Schutz, J., N Rezg and J.B. Léger (2008). Contribution for an Optimal Choice of Business and Maintenance Plans, Based on an Adaptive Failure Law. *Proceedings of the 9th IFAC Workshop on Intelligent Manufacturing Systems (IFAC'08)*, pp. 313–318, Szczecin, Poland.
- [Schutz, 2008C] Schutz, J., N Rezg and J.B. Léger (2008). Periodic preventive maintenance policy in finite horizon with an adaptive failure law. *IEEE International Conference on Industrial Engineering and Engineering Management (IEEM 2008)*, pp. 2117–2121, Singapore.
- [Schutz, 2009A] Schutz, J., N Rezg and J.B. Léger (2009). Efficient Solution of Business and Maintenance Plans under Adaptive Failure Law. *Proceedings of the 20th International Conference on Systems Engineering (ICSE'09)*, Coventry, United Kingdom.
- [Schutz, 2009B] Schutz, J., N. Rezg and J.B. Léger. (2009). Periodic and sequential preventive maintenance policies over a finite planning horizon with a dynamic failure law. *Journal of Intelligent Manufacturing*. (doi: 10.1007/s10845-009-0313-7).
- [Schutz, 2009C] Schutz, J., N. Rezg and J.B. Léger. (2009). Détermination conjointe de plans efficaces d'exploitation et de maintenance soumis à une loi de défaillance dynamique. *4^{ème} édition du colloque international francophone « PErformances et Nouvelles TechnolOgies en Maintenance » (PENTOM'09)*, Autrans, France.
- [Shapiro, 1993] Shapiro J.F. (1993). Mathematical programming models and methods for production planning and scheduling, *Handbooks in Operations Research and Management Science, Logistics of Production and Inventory*, 4, pp. 371–439.
- [Siarry, 2005] Siarry, P., J. Dréo, A. Pétrowski and É. Taillard (2005). Métaheuristiques pour l'optimisation difficile. Edition Eyrolles.
- [Sortrakul, 2005] Sortrakul, N., H.L. Nachtmann and C.R. Cassady (2005). Genetic algorithms for integrated preventive maintenance planning and production scheduling for a single machine. *Computers in Industry*, 56(2), pp. 161–168.
- [Srinivasan, 1996] Srinivasan, M.M., H.-S. Lee (1996). Production-inventory systems with preventive maintenance. *IIE Transactions*, 28(11), pp. 879–890.

- [Suliman, 2000] Suliman, S.M.A. (2000). A two-phase heuristic approach to the permutation flow-shop scheduling problem. *International Journal of Production Economics*, 64, pp. 143–152.
- [Wang, 2002] Wang, H. (2002). A survey of maintenance policies of deteriorating systems. *European Journal of Operational Research*, 139(3), pp. 469–489.
- [Widmer, 1989] Widmer, M. and A. Hertz (1989). A new heuristic method for the flow shop sequencing problem, *European Journal of Operational Research*, 41, pp. 186–193.
- [Yeh, 2009] Yeh, R.H., K.C. Kao and WL Chang (2009). Optimal preventive maintenance policy for leased equipment using failure rate reduction. *Computers and Industrial Engineering*, 57(1), pp. 304–309.
- [Zadeh, 1965] Zadeh, L.A. (1965). Similarity relations and fuzzy ordering. *Information Sciences*, 3, pp. 177–200.
- [Zadeh, 1973] Zadeh, L.A. (1973). Outline of a new approach to the analysis of complex systems and decision processes. *IEEE Transactions on systems, Man and Cybernetic*, 3(1), pp. 28–44.
- [Zhou, 2007] Zhou, X., L. Xi and J. Lee (2007). Reliability-centered predictive maintenance scheduling for a continuously monitored system subject to degradation. *Reliability Engineering & System Safety*, 92(4), pp. 530–534.

Résumé

Dans le domaine naval, lorsqu'un navire quitte le quai, il s'engage à réaliser un ensemble de missions durant un horizon de temps déterminé. L'ensemble des missions à accomplir ainsi que l'ordonnancement de celles-ci constituent le plan d'exploitation. Un navire pouvant réaliser différents types de missions (diplomatique, océanographique, etc.), ces dernières ont des conditions de fonctionnement variables tant au niveau opérationnel qu'au niveau environnemental. L'objectif consiste donc à choisir parmi toutes les missions proposées, dont la durée cumulée dépasse l'horizon de temps alloué, celles à réaliser ainsi que l'ordonnancement. Par ailleurs, les conditions de fonctionnement étant différentes suivant les missions, l'impact sur la fiabilité du système varie. Aussi, afin de minimiser les coûts de maintenance, différentes stratégies de maintenance seront étudiées telles que les politiques périodiques, séquentielles, basées sur un seuil de fiabilité. De manière générale, les politiques de maintenance sont constituées d'actions correctives minimales et préventives imparfaites. Compte tenu des hypothèses de travail (horizon de temps fini, loi de défaillance dynamique), l'optimisation de ces politiques de maintenance sera réalisée à l'aide de procédures numériques ou de méta-heuristiques. L'optimisation des plans d'exploitation étant résolue à l'aide de méta-heuristiques, l'optimisation conjointe des deux plans peut résulter de l'intégration de deux méta-heuristiques. Le travail de recherche a donc pour objectif le développement d'un outil d'aide à la décision basé sur le pronostic. Il doit permettre de générer les plans d'exploitation et de maintenance optimaux, compte tenu de l'historique des missions réalisées, du retour d'expérience d'autres navires, mais aussi des missions génériques possibles.

Abstract

The present research work aims to develop a decision-making tool based on prognosis. Indeed, in the naval domain, when a navy ship leaves the dock, it should experience a set of missions within a finite time horizon. The business plan consist to perform as well as to schedule the overall missions. As a navy ship can achieve different type of missions (diplomatic, oceanographic, etc.), the system degradations varie according to both operational and environmental conditions. Therefore, the business plan consist on choosing among all possible missions those whose the cumulative duration does not exceed the finite time horizon. Such missions should be then scheduled and performed. On the other hand, our work deals with different maintenance policies such as systematic, sporadic, periodic and sequential. Generally, the maintenance actions are of two types, namely minimal corrective and imperfect preventive. Based on some assumptions (finite time horizon, dynamic failure law), the optimization of maintenance policy will be carried out using exact numerical procedures or meta-heuristics. To solve the joint optimization of business and maintenance plans, we use the integration of two types of meta-heuristics. The obtained solution should define the optimal business and maintenance plans given the history of missions performed, the experiences of other navy ships but also of all possible generic missions.