



Evaluation de paramètres de sûreté de fonctionnement en présence d'incertitudes et aide à la conception : application aux Systèmes Instrumentés de Sécurité

Mohamed Sallak

► To cite this version:

Mohamed Sallak. Evaluation de paramètres de sûreté de fonctionnement en présence d'incertitudes et aide à la conception : application aux Systèmes Instrumentés de Sécurité. Automatique / Robotique. Institut National Polytechnique de Lorraine - INPL, 2007. Français. NNT: . tel-00202324

HAL Id: tel-00202324

<https://theses.hal.science/tel-00202324>

Submitted on 5 Jan 2008

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

**Evaluation de paramètres de sûreté de fonctionnement
en présence d'incertitudes et aide à la conception :
Application aux Systèmes Instrumentés de Sécurité**

THÈSE

présentée et soutenue publiquement le 19 Octobre 2007

pour l'obtention du

Doctorat de l'Institut National Polytechnique de Lorraine

Spécialité Automatique, Traitement du Signal et Génie Informatique

par

MOHAMED SALLAK

Composition du jury

<i>Président :</i>	D. MAQUIN	Professeur à l'INPL
<i>Rapporteurs :</i>	L. FOULLOY	Professeur à l'Université de Savoie
	E. CHATELET	Professeur à l'Université Technologique de Troyes
<i>Examineurs :</i>	Y. DUTUIT	Professeur à l'Université de Bordeaux I
	J-F. AUBRY	Professeur à l'INPL (Directeur de Thèse)
	C. SIMON	Maître de Conférences à l'Université de Nancy II (Co-directeur de Thèse)



Centre de Recherche en Automatique de Nancy
UMR 7039 - Nancy Université - CNRS
2, Avenue de la Forêt de Haye 54516 Vandœuvre-Lès-Nancy
Tél. +33 (0)3 83 59 59 59 Fax +33 (0)3 83 59 56 44

Mis en page avec la classe thloria.

Remerciements

Le travail présenté dans ce mémoire a été effectué au sein de l'équipe SACSS (Systèmes Automatisés Contraints par la Sûreté de fonctionnement et la Sécurité) du groupe thématique SURFDIAG (Sûreté de Fonctionnement et Diagnostic des systèmes) du Centre de Recherche en Automatique de Nancy (CRAN).

J'exprime mes profonds remerciements à mes deux directeurs de thèse, Christophe Simon et Jean-Francois Aubry pour leur aide inestimable, leur patience et leurs encouragements tout au long de ce travail. Leurs compétences ont été un atout indéniable à la réussite de ces travaux et m'ont permis d'apprendre énormément durant ces trois années de collaboration.

Je remercie Monsieur Alain Richard, directeur du CRAN pour m'avoir permis de réaliser cette thèse. Je remercie aussi le responsable l'école doctorale IAEM Lorraine, Monsieur Francis Lepage, pour son engagement dans notre formation doctorale.

Je souhaiterais aussi remercier les rapporteurs Monsieur Eric Chatelet et Monsieur Laurent Foulloy qui ont accepté de se plonger dans ce sujet, ainsi que les autres membres qui ont accepté de participer au jury.

Il me reste à remercier toutes les personnes avec qui j'ai travaillé durant ces trois années.

Je dédie cette thèse à mes parents...

Table des matières

Table des figures	vii
Liste des tableaux	ix
Production personnelle	1
Acronymes	3
Introduction	5
1 Problématique	5
2 Objectifs	7
3 Organisation du rapport de thèse	8
Chapitre 1 Evaluation de la sûreté de fonctionnement des systèmes dédiés aux applications de sécurité	11
1.1 Problématique de la sécurité	12
1.1.1 Introduction	12
1.1.2 Notion de danger	13
1.1.2.1 Définitions	13
1.1.2.2 Phénomènes dangereux et situations dangereuses	13
1.1.3 Notion de risque	14
1.1.3.1 Mesure de risque	15
1.1.3.2 Classification des risques	17
1.1.3.3 Risque acceptable	18
1.1.3.4 Risque majeur	18
1.1.3.5 Risque industriel	18
1.1.4 Notion de sécurité	19
1.1.5 Principes généraux de protection	19
1.1.5.1 Sécurités passives	19
1.1.5.2 Sécurités actives	19

1.2	Sécurité fonctionnelle	20
1.2.1	Définitions	20
1.2.2	Systèmes relatifs aux applications de sécurité	20
1.3	Référentiel normatif	21
1.3.1	Introduction	21
1.3.2	Norme IEC 61508	21
1.3.2.1	Domaine d'application	21
1.3.2.2	Structure générale de la norme	21
1.3.2.3	Objectifs de la norme	22
1.3.2.4	Evaluation du niveau d'intégrité de sécurité (SIL)	22
1.3.2.5	Quelques commentaires	27
1.3.3	Norme IEC 61511	28
1.3.3.1	Domaine d'application	29
1.3.3.2	Structure générale de la norme	30
1.3.3.3	Objectifs de la norme	30
1.3.3.4	Quelques commentaires	32
1.4	Problématique de l'évaluation de la sûreté de fonctionnement des SIS	32
1.5	Conclusion	33

Chapitre 2 Evaluation de la sûreté de fonctionnement des systèmes en présence d'informations imparfaites 35

2.1	Introduction	36
2.2	Représentation des connaissances imparfaites	37
2.2.1	Théorie des probabilités	37
2.2.1.1	Définitions et propriétés	37
2.2.1.2	Cadre bayésien	38
2.2.1.3	Limitations de la théorie des probabilités	39
2.2.1.4	Conclusion	39
2.2.2	Théorie des sous-ensembles flous	39
2.2.2.1	Définitions et concepts de base	39
2.2.2.2	Opérations simples sur les sous-ensembles flous	40
2.2.2.3	Nombres flous du type $L - R$	42
2.2.2.4	Notion de α -coupes	44
2.2.3	Théorie des possibilités	45
2.2.3.1	Mesure et distribution de possibilités	45
2.2.3.2	Mesure de nécessité	46
2.2.4	Théorie des fonctions de croyance	47

2.2.4.1	Fonction de masse (ou de croyance)	47
2.2.4.2	Fonctions de crédibilité et de plausibilité	47
2.2.5	Conclusion	48
2.3	Méthodes de sûreté de fonctionnement en présence d'informations imparfaites	49
2.3.1	Arbres de défaillance	49
2.3.1.1	Principes, construction et analyse	49
2.3.1.2	Arbres de défaillances flous	53
2.3.1.3	Méthode de Singer	53
2.3.1.4	Exemple	55
2.3.1.5	Méthode des α -coupes	56
2.3.1.6	Cas des événements répétés	57
2.3.1.7	Dépendance des événements	57
2.3.1.8	Remarque	58
2.3.2	Approches markoviennes	58
2.3.2.1	Processus stochastiques	58
2.3.2.2	Chaînes de Markov	59
2.3.2.3	Processus de Markov	61
2.3.2.4	Observations et limitations	63
2.3.2.5	Chaînes de Markov à états non flous et probabilités de transition floues	63
2.3.2.6	Chaînes de Markov à états flous et non flous et probabilités de transition singulières	66
2.3.2.7	Modèle de Markov avec un système d'inférence flou (MMF)	70
2.3.2.8	Conclusion et perspectives	72
2.3.3	Réseaux de Petri	73
2.3.3.1	Introduction	73
2.3.3.2	Définitions et concepts de base	73
2.3.3.3	Réseaux de Petri temporels	74
2.3.3.4	Réseaux de Petri p-temporels	75
2.3.3.5	Réseaux de Petri stochastiques	76
2.3.3.6	Réseaux de Petri flous (RdPF)	77
2.3.3.7	Observations et limitations	81
2.4	Conclusion	81

Chapitre 3 Evaluation des SIL par une approche probabiliste floue et études de sensibilité 83

3.1	Introduction	84
-----	--------------	----

3.2	Evaluation des SIL par une approche probabiliste floue	85
3.2.1	Modélisation des taux de défaillance des composants des SIS	85
3.2.2	Evaluation du PFD_{avg} des SIS	86
3.3	Caractérisation de l'incertitude des résultats obtenus	89
3.4	Etudes de sensibilité	90
3.4.1	Facteurs d'importance probabilistes	90
3.4.1.1	Hypothèses et notations	91
3.4.1.2	Facteur d'importance marginal	92
3.4.1.3	Facteur d'importance critique	93
3.4.1.4	Facteur d'importance critique des coupes minimales	94
3.4.1.5	Facteur d'importance de diagnostic	94
3.4.1.6	Observations et limitations	95
3.4.2	Facteurs d'importance intégrant les connaissances imparfaites	95
3.4.2.1	Facteurs d'importance sur les incertitudes	95
3.4.2.2	Facteurs d'importance flous	96
3.4.3	Facteurs d'importance proposés	98
3.4.3.1	Facteur d'importance flou (FPIM)	98
3.4.3.2	Facteur d'incertitude flou (FPUM)	98
3.4.4	Conclusion partielle	98
3.5	Cas d'application	99
3.5.1	Description du système	99
3.5.2	Hypothèses	99
3.5.3	Approche probabiliste floue	101
3.5.4	Caractérisation de l'incertitude	104
3.5.5	Etudes de sensibilité et réduction de l'incertitude du SIL	104
3.5.5.1	Facteur d'importance flou (FPIM)	104
3.5.5.2	Facteur d'incertitude flou (FPUM)	105
3.6	Conclusions et perspectives	108

Chapitre 4 Conception optimale des SIS 109

4.1	Introduction	110
4.2	Etat de l'art de l'allocation de fiabilité et de redondance	111
4.2.1	Méthodes pondérées	111
4.2.1.1	Méthode AGREE	112
4.2.1.2	Méthode Karmiol	112
4.2.1.3	Méthode ARINC	113
4.2.1.4	Méthode d'égale allocation	114

4.2.2	Méthodes d'optimisation	114
4.3	Algorithmes génétiques (AG)	116
4.3.1	Introduction	116
4.3.2	Nomenclature des AG	117
4.3.3	Avantages des AG	117
4.3.4	Inconvénients des AG	118
4.3.5	Concepts de base	118
4.3.6	Opérateurs des AG	118
4.3.6.1	Codage des solutions	118
4.3.6.2	Population initiale	120
4.3.6.3	Taille des populations	120
4.3.6.4	Sélection	120
4.3.6.5	Croisement	120
4.3.6.6	Mutation	121
4.3.6.7	Remplacement	121
4.3.7	Choix des paramètres des AG	121
4.4	Modélisation des SIS	122
4.4.1	Blocs diagrammes de Fiabilité (BdF)	122
4.4.1.1	Diagramme série	122
4.4.1.2	Diagramme parallèle	122
4.4.1.3	Diagramme en redondance k/n	123
4.4.1.4	Diagramme complexe	124
4.4.1.5	Liens minimaux et coupes minimales	124
4.4.2	Réseaux de fiabilité	125
4.4.2.1	Rappel sur la théorie des graphes	125
4.4.2.2	Concepts de base des réseaux de fiabilité	127
4.4.2.3	Lien d'un réseau de fiabilité	129
4.4.2.4	Coupe d'un réseau de fiabilité	130
4.4.2.5	Lien minimal et coupe minimale d'un réseau de fiabilité	130
4.4.2.6	Remarque	130
4.5	Conception optimale	130
4.5.1	Introduction	130
4.5.2	Pourquoi les AG ?	130
4.5.3	Pourquoi les réseaux de fiabilité ?	131
4.5.4	Méthodologie générale	131
4.5.5	Cas d'application	137
4.5.5.1	Formulation mathématique du problème	139

4.5.5.2	Optimisation par l'AG	140
4.5.6	Résultats et analyse	140
4.5.6.1	Exigence SIL1	140
4.5.6.2	Exigence SIL2	140
4.5.6.3	Exigences SIL3 et SIL4	142
4.6	Conception optimale à paramètres flous	146
4.6.1	Introduction	146
4.6.2	Méthodologie générale	146
4.6.3	Cas d'application	147
4.6.4	Résultats et analyse	149
4.6.4.1	Exigence SIL1	149
4.6.4.2	Exigence SIL2	149
4.6.4.3	Exigence SIL3	151
4.6.4.4	Exigence SIL4	153
4.7	Conclusion et perspectives	154
Conclusion générale		155
Bibliographie		159

Table des figures

1.1	Scénario d'accident [Desroches, 2005]	14
1.2	Diagramme état/transition du processus accidentel [Mazouni <i>et al.</i> , 2007]	15
1.3	Courbe de Farmer [Farmer, 1967]	16
1.4	Caractérisation du risque [Gouriveau, 2003]	17
1.5	Structure générale de la norme IEC 61508 [IEC61508, 1998]	23
1.6	Graphe de risque [IEC61508, 1998]	24
1.7	Matrice de risque [IEC61508, 1998]	26
1.8	L'IEC 61508 et ses normes sectorielles	29
1.9	Utilisateurs de l'IEC 61511 et l'IEC 61508	29
1.10	Relations entre l'IEC 61511 et l'IEC 61508	30
1.11	Structure générale de la norme IEC 61511 [IEC61511, 2000]	31
2.1	Nombre flou du type L-R	43
2.2	α -coupes d'un nombre flou	44
2.3	Arbre de défaillance	49
2.4	Opérateur de condition	50
2.5	Opérateur " m/n "	50
2.6	Processus d'analyse des événements intermédiaires [Villemeur, 1998]	51
2.7	Processus d'élaboration d'un arbre de défaillance [Villemeur, 1998]	52
2.8	Arbre de défaillance avec un opérateur ET	53
2.9	Arbre de défaillance avec un opérateur OU	54
2.10	Arbre de défaillance simplifié	55
2.11	Probabilité d'occurrence de l'événement "surchauffe du fil AB"	56
2.12	Fonctions d'appartenance des valeurs Télévée et Lmoyenne	71
2.13	Modèle flou de Markov	72
2.14	Représentation d'un chien de garde	75
2.15	Modélisation de l'implication logique par une transition temporelle	78
2.16	Distribution de possibilité sur les places p_i en fonction du temps : (P, p_0, T)	79
2.17	Distribution de possibilité sur les places p_i (π, P) aux instants $t = t_1$ et $t = t_2$	80
3.1	Modélisation du taux de défaillance imprécis λ_i par un nombre flou triangulaire	86
3.2	Différentes formes possibles des fonctions d'appartenance	87
3.3	Probabilité floue de défaillance d'un composant i	88
3.4	Arbre de défaillance	88
3.5	α -coupes d'une probabilité floue	89
3.6	Traduction de l'incertitude par les mesures de possibilité et de nécessité	90
3.7	Réservoir sous pression	100

3.8	Configuration du SIS	101
3.9	Arbre de défaillance du système complet	102
3.10	PFD_{avg} du SIS	103
3.11	Probabilités de défaillance précise et imprécise de l'unité de traitement	105
3.12	PFD_{avg} du SIS selon les modifications du SIS après le calcul du FPIM	105
3.13	Probabilités de défaillance précise et imprécise du capteur de température	107
3.14	PFD_{avg} du SIS selon les modifications du SIS après le calcul du FPUM	107
4.1	Organigramme simplifié de conception des SIS	111
4.2	Algorithme génétique de base	119
4.3	BdF d'un système série	122
4.4	BdF d'un système parallèle	123
4.5	BdF d'un système k/n	123
4.6	Schéma de connexion d'un système complexe	124
4.7	Dessin d'un graphe	126
4.8	Matrice booléenne	126
4.9	Dictionnaire des suivants	126
4.10	Graphe r -appliqué	127
4.11	Graphe partiel	128
4.12	Réseau de fiabilité	128
4.13	Réseau de fiabilité partiel	129
4.14	Réseau de fiabilité général d'un SIS	132
4.15	Codage du réseau de fiabilité	133
4.16	Exemple d'un réseau de fiabilité	133
4.17	Correspondance entre réseaux de fiabilité et schémas de connexion	134
4.18	Schéma de connexion général d'un SIS	137
4.19	Organigramme de conception optimale d'un SIS	138
4.20	Schéma de connexion du SIS à concevoir	139
4.21	Architecture obtenue (SIL1 : $A_{avg} = 0.906056, C = 8100$) : Réseau de fiabilité . .	141
4.22	Architecture obtenue (SIL1 : $A_{avg} = 0.906056, C = 8100$) : schéma de connexion .	141
4.23	Architecture obtenue (SIL1 : $A_{avg} = 0.916748, C = 8800$) : Réseau de fiabilité . .	141
4.24	Architecture obtenue (SIL1 : $A_{avg} = 0.916748, C = 8800$) : schéma de connexion .	141
4.25	Architecture obtenue (SIL1 : $A_{avg} = 0.925333, C = 8700$) : Réseau de fiabilité . .	142
4.26	Architecture obtenue (SIL1 : $A_{avg} = 0.925333, C = 8700$) : schéma de connexion .	142
4.27	Architecture obtenue (SIL2 : $A_{avg} = 0.990326, C = 14200$) : Réseau de fiabilité .	143
4.28	Architecture obtenue (SIL2 : $A_{avg} = 0.990326, C = 14200$) : schéma de connexion	143
4.29	Architecture obtenue (SIL2 : $A_{avg} = 0.992653, C = 14800$) : Réseau de fiabilité .	143
4.30	Architecture obtenue (SIL2 : $A_{avg} = 0.992653, C = 14800$) : schéma de connexion	143
4.31	Architecture obtenue (SIL2 : $A_{avg} = 0.992805, C = 14800$) : Réseau de fiabilité .	144
4.32	Architecture obtenue (SIL2 : $A_{avg} = 0.992805, C = 14800$) : schéma de connexion	144
4.33	Architecture obtenue (SIL3 : $A_{avg} = 0.999003, C = 16900$) : Réseau de fiabilité .	145
4.34	Architecture obtenue (SIL3 : $A_{avg} = 0.999003, C = 16900$) : schéma de connexion	145
4.35	Probabilités de défaillance floues des capteurs	147
4.36	Probabilités de défaillance floues des unités de traitement	148
4.37	Probabilités de défaillance floues des actionneurs	149
4.38	Architecture obtenue (SIL1 : $A_{avg} = 0.916904, C = 8800$) : Réseau de fiabilité . .	150
4.39	Architecture obtenue (SIL1 : $A_{avg} = 0.916904, C = 8800$) : schéma de connexion .	150
4.40	Architecture obtenue (SIL1 : $A_{avg}=0.912206, C=8800$) : Réseau de fiabilité	150

4.41	Architecture obtenue (SIL1 : $A_{avg}=0.912206, C=8800$) : schéma de connexion	. . . 150
4.42	Architecture obtenue (SIL2 : $A_{avg}=0.990437, C=18300$) : Réseau de fiabilité	. . . 151
4.43	Architecture floue obtenue (SIL2 : $A_{avg}=0.990437, C=18300$) : schéma de connexion	151
4.44	Architecture obtenue (SIL2 : $A_{avg}=0.994834, C=20400$) : Réseau de fiabilité	. . . 152
4.45	Architecture obtenue (SIL2 : $A_{avg}=0.994834, C=20400$) : schéma de connexion	. . 152
4.46	Architecture obtenue (SIL3 : $A_{avg} = 0.99906, C = 20400$) : Réseau de fiabilité	. . 153
4.47	Architecture obtenue (SIL3 : $A_{avg} = 0.99906, C = 20400$) : schéma de connexion	. 153

Liste des tableaux

1.1	Recueil des plus graves accidents industriels survenus dans le monde entre 1960 et 2007 [Ministère de l'écologie, 2007].	13
1.2	Table d'estimation du risque [Guiochet, 2003]	16
1.3	Niveaux d'intégrité de sécurité (SIL) : Sollicitation faible [IEC61508, 1998]	22
1.4	Niveaux d'intégrité de sécurité (SIL) : Sollicitation élevée ou continue [IEC61508, 1998]	24
2.1	Principales normes et co-normes triangulaires [Wu, 1998].	42
2.2	Paramètres flous des probabilités d'occurrence des événements de base	56
3.1	Paramètres flous des probabilités de défaillance des composants du SIS	101
3.2	Coupes minimales de l'arbre de défaillance	103
3.3	Mesures de possibilité (II) et de nécessité (N) des SIL obtenus	104
3.4	Facteur d'importance flou (FPIM)	106
3.5	Mesures de possibilité (II) et de nécessité (N) après modifications selon le facteur (FPIM)	106
3.6	Facteur d'incertitude flou (FPUM)	106
3.7	Mesures de possibilité (II) et de nécessité (N) après modifications selon le facteur <i>FPUM</i>	107
4.1	Principales méthodes d'allocation pondérées	114
4.2	Résumé de la terminologie utilisée en AG	117
4.3	Coûts et probabilités de défaillance des composants du SIS	139
4.4	Résultats pour l'obtention d'un SIL1	140
4.5	Résultats pour l'obtention d'un SIL 2	142
4.6	Coûts et probabilités de défaillance des composants du SIS n°2	144
4.7	Paramètres flous des probabilités de défaillance des capteurs	147
4.8	Paramètres flous des probabilités de défaillance des unités de traitement	148
4.9	Paramètres flous des probabilités de défaillance des actionneurs	148
4.10	Approche floue : résultats pour l'obtention d'un SIL1	149
4.11	Approche floue : résultats pour l'obtention d'un SIL2	151
4.12	Approche floue : résultats pour l'obtention d'un SIL3	153

Production personnelle

Revues internationales

- *A fuzzy probabilistic approach for determining Safety Integrity Level*

Sallak M., Simon C., Aubry. J-F.

Accepté à *IEEE Transactions on Fuzzy Systems*, À paraître en 2007.

Revues nationales

- *Aide à la décision dans la réduction de l'incertitude des SIL : une approche floue/possibiliste*

Sallak M., Aubry. J-F., Simon C.

e-STA, Revue des Sciences et Technologies de l'Automatique, vol. 3, n° 3, 2006.

Conférences internationales avec actes et comités de lecture

- *Evaluating Safety Integrity Level in presence of uncertainty*

Sallak M., Aubry J-F., Simon C.

International Conference on Safety and Reliability, KonBin 2006, Journal of KonBin 2006, vol. 1, pp. 411-419, Pologne

- *On the use of a new possibilist importance measure to reduce Safety Integrity Level uncertainty*

Sallak M., Aubry J-F., Simon C.

International Conference on Safety and Reliability, KonBin 2006, Journal of KonBin 2006, Pologne, vol. 2, pp. 77-85

- *Aide à la décision dans la réduction de l'incertitude des SIL : une approche floue/possibiliste*

Sallak M., Aubry J-F., Simon C.

Conférence Internationale Francophone d'Automatique, CIFA 2006, Bordeaux, France, CDROM papier n°183

- *Optimal design of safety instrumented systems*

Sallak M., Simon C., Aubry J-F.

ACD, Workshop on Advanced Control and Diagnosis, Nancy, 2006, CDROM papier n°30

- *Optimal design of safety instrumented systems : A graph reliability approach*

Sallak M., Simon C., Aubry J-F.

7^{ème} édition du congrès international pluridisciplinaire, Qualita 2007, Tanger, Maroc

- *SIL allocation of SIS by aggregation of experts opinions*

Simon C., Sallak M., Aubry J-F.

ESREL, Safety and Reliability Conference, Stavanger, Norvège, 2007

Conférences nationales avec actes et comités de lecture

- ***Impact de l'imprécision des taux de défaillances dans les Arbres de Défaillances et facteurs d'importance flous***
Sallak M., Aubry J-F., Simon C.
Journées Doctorales Modélisation, Analyse et Conduite des Systèmes dynamiques, JD-MACS 2005, Lyon, France, 2005, CDROM papier n°83
- ***Allocation de SIL par agrégation d'avis d'experts***
Simon C., Sallak M., Aubry J-F.
15ème Colloque National de Maîtrise des Risques et Sécurité de Fonctionnement, Lambda-mu 15, Lille, 2006

Acronymes

A Disponibilité
ANSI American National Standards Institute
AG Algorithme Génétique
BdF Bloc diagramme de Fiabilité
GCEF Graphe des Classes d'États Flous
E/E/PE Electriques/Electroniques/Electroniques programmables de sécurité
ESD Entité Source de Danger
ECD Entité Cible de Danger
EVD Événement Dangereux
EVI Événement Initiateur
EVR Événement Redouté
IEC International Electrotechnical Commission
INS Immédiates, nécessaires et suffisantes
ISA Instrument Society of America
ISO International Organization for Standardization
MMF Modèle de Markov Flou
OHSAS Occupational Health and Safety Assessment Series
PFD Probability of Failure on Demand
PFH Probability of Failure per Hour
FPIM Fuzzy Probist Importance Measure
FPUM Fuzzy Probist Uncertainty Measure
RdP Réseau de Petri
RdPF Réseau de Petri Flou
RdPTF Réseau de Petri Temporel Flou
SA Situation d'Accident
SD Situation Dangereuse
SE Situation d'Exposition
SIS Safety Integrity System
SIF Safety Integrity Function
SIL Safety Integrity Level
SRECS Systèmes de Commande Électriques Relatifs à la Sécurité de machines

Introduction

1 Problématique

Dans une société de plus en plus sensible aux notions de sécurité et de développement durable, les entreprises ont le souci d'éviter les dangers pouvant induire incendies, explosions et autres rejets de matières dangereuses, sources de dommages pour les personnes, l'environnement ou les biens. Ainsi, font-ils l'inventaire des risques, de leur nature, établissent-ils les mécanismes qui peuvent conduire au déclenchement d'accidents et estiment-ils leurs conséquences et la fréquence à laquelle ils sont susceptibles de survenir. Les risques traduisent à la fois la gravité du dommage et la fréquence d'occurrence de l'événement dangereux. Les échelles de gravité comportent plusieurs niveaux qui sont évalués en fonction des conséquences de l'événement dangereux sur les personnes, l'environnement et les biens. De même, un certain nombre de niveaux est fixé pour la probabilité ou la fréquence d'occurrence de l'événement dangereux. Une fois le niveau de risque évalué, l'entreprise cherchera à le ramener sous un seuil acceptable.

Les moyens à mettre en œuvre pour réduire les risques sont nombreux et variés. La conception du procédé, le choix des équipements participent en premier lieu à la réduction du risque. On peut aussi agir sur le système de contrôle commande du procédé, en prévoyant par exemple des redondances et des solutions de repli en cas de dysfonctionnement. Ces approches ne sont pas toujours suffisantes. Pour réduire encore les risques, il faut prévoir des systèmes de sécurité. Ceux-ci entrent en action lorsque le procédé se trouve dans des conditions anormales de fonctionnement et qu'une situation dangereuse risque de se développer. Il existe différents types de sécurité : les sécurités passives (confinement, rétention, etc.), les sécurités actives présentes sur le procédé (soupapes, etc.) et enfin les sécurités actives instrumentées. Ces dernières appartiennent au Systèmes Instrumentés de Sécurité (SIS, Safety Instrumented Systems).

Les SIS ont pour objectif de mettre le procédé en position de repli de sécurité lorsqu'il évolue vers une voie comportant un risque réel (explosion, feu, etc.), c'est-à-dire un état stable ne présentant pas de risque pour les personnes, l'environnement ou les biens. Pour définir les SIS, il faut identifier et évaluer les risques contre lesquels il faut se protéger. Les principales normes et standards en termes de sécurité (IEC 61508 [IEC61508, 1998], IEC 61511 [IEC61511, 2000], etc.) peuvent être utilisés pour les concevoir. Les méthodes que proposent ces normes sont basées sur une estimation de la réduction nécessaire du risque que doit permettre d'atteindre le SIS.

Les normes IEC 61508 [IEC61508, 1998] (à vocation générale) et IEC 61511 [IEC61511, 2000] (orientée vers les industries des procédés) définissent les niveaux d'intégrité de sécurité (SIL, Safety Integrity Levels) qui fixent le niveau de réduction du risque que doit atteindre le SIS. Il existe 4 niveaux possibles, notés SIL1 à SIL4. Chacun d'eux dépend de la gravité et de la fréquence d'occurrence du risque. Il est évident que si un risque est très important, il nécessite des

parades très efficaces. Il se verra automatiquement attribuer un SIL élevé (3 voire même 4). Ces deux normes définissent un critère important pour caractériser les SIS : la probabilité moyenne de défaillance sur demande (PFD_{avg} : Average Probability of Failure on Demand) pour les SIS faiblement sollicités (moins d'une sollicitation par an) et la probabilité de défaillance par heure (PFH : Probability of Failure per Hour) pour les SIS fortement sollicités ou agissant en mode continu.

Les méthodes usuelles de calcul du PFD_{avg} des SIS sont des méthodes probabilistes [Stavrianidis and Bhimavarapu, 1998; IEC61508, 1998; Beckman, 2000; IEC61511, 2000; Summers, 2002; Goble and Cheddie, 2006; Peguet and Boyer, 2006; Innal *et al.*, 2006]. Ces méthodes issues des études traditionnelles de sûreté de fonctionnement où les données de fiabilité relatives aux composants (taux de défaillance, taux de réparation, etc.) peuvent être connues avec précision et validées par le retour d'expérience. Nous pouvons cependant nous interroger sur leur adaptation à des systèmes hautement fiables pour lesquels les défaillances sont très rares. Dans ce cas, le retour d'expérience est insuffisant pour valider avec précision les taux de défaillance. En outre, en raison de l'évolution constante de l'environnement et de la complexité des installations industrielles, les conditions d'utilisation des composants des SIS utilisés dans les installations peuvent changer pour un même composant. L'idéal est de disposer d'une quantité d'information suffisante concernant les défaillances des composants pour pouvoir estimer avec précision leurs taux de défaillance. Or, dans les installations industrielles, les SIS sont des systèmes rarement sollicités. De ce fait, les composants des SIS ne sont pas suffisamment utilisés pour fournir des données de défaillances suffisantes [Villemeur, 1998; Wang *et al.*, 2004].

Les bases de données de fiabilité génériques sont souvent utilisées pour fournir les taux de défaillance des composants des systèmes relatifs à la sécurité [OREDA, 2002; Exida, 2005; Hauge *et al.*, 2006; Goble and Cheddie, 2006]. Des valeurs ponctuelles des taux de défaillance sont généralement utilisées pour estimer les probabilités de défaillance des composants du SIS. Cependant, la collecte de données de fiabilité et l'utilisation de ces données pour estimer les taux de défaillance d'autres composants du même type introduisent des incertitudes. Selon [Kletz, 1999], les données de fiabilité recueillies pour un composant peuvent changer d'un facteur de 3 ou 4, et un facteur de 10 n'est pas à exclure. Dans [MIL217B, 1974; Villemeur, 1998], on propose d'utiliser des coefficients d'influence pour tenir compte des conditions réelles d'utilisation des valeurs génériques des taux de défaillance proposées par les bases de données de fiabilité. Wang *et al.* [Wang *et al.*, 2004] ont discuté de l'impact de l'incertitude des taux de défaillance des composants du SIS sur la détermination du SIL des Fonctions Instrumentées de Sécurité (SIF) que doit réaliser le SIS. Mais, ils ont juste souligné que d'autres travaux doivent être menés pour étudier et justifier l'incertitude du SIL dans ce type de problème.

Certaines bases de données de fiabilité [IEEE, 1984; CCPS, 1991; CCPS, 2002] fournissent les bornes inférieures et supérieures, les valeurs moyennes et les facteurs d'erreurs des taux de défaillance des composants. Des méthodes autres que les approches probabilistes classiques (ensembles flous, théorie de l'évidence, etc.) peuvent utiliser avantageusement ces valeurs pour prendre en compte l'imprécision (qu'on appelle aussi incertitude de type épistémique [Helton and Oberkampf, 2004; Baea *et al.*, 2004]) liée à ces taux de défaillance et estimer les probabilités de défaillance des composants ainsi que les PFD_{avg} des SIS.

2 Objectifs

Cette thèse aborde d'abord la problématique de la prise en compte des incertitudes relatives aux données de fiabilité des composants pour l'évaluation de la sûreté de fonctionnement des systèmes complexes en général et les SIS en particulier. Dans un premier temps, nous avons comme objectif de modéliser les imprécisions des taux de défaillance des composants par des nombres flous, puis d'évaluer les probabilités floues de défaillance des composants à partir de ces taux de défaillance imprécis. Dans un second temps, nous proposons l'utilisation du formalisme des arbres de défaillance flous [Tanaka *et al.*, 1983; Singer, 1990; Soman and Misra, 1993] pour évaluer le PFD_{avg} du SIS et le SIL du SIS à partir des probabilités de défaillance floues. Dans nos travaux, nous parlerons indifféremment du SIL d'un SIS ou du SIL d'une fonction instrumentée de sécurité (SIF, Safety Instrumented Function), car nous supposons que chaque SIS ne peut réaliser qu'une seule SIF. La probabilité d'occurrence de l'événement sommet qui représente la probabilité de défaillance du SIS lors de sa sollicitation est obtenue à partir des probabilités de défaillance des composants du SIS en utilisant des opérations arithmétiques floues des composants. L'introduction de deux nouveaux facteurs d'importance flous inspirés des facteurs d'importance classiques [Birnbaum, 1969; Lambert, 1975] permettra de mettre en évidence les composants contribuant le plus à la défaillance du SIS et ceux contribuant le plus à l'incertitude entachant le SIL et de la réduire efficacement.

Le second objectif de la thèse est d'introduire une méthodologie d'aide à la conception des SIS en optimisant le choix des composants et la structure des SIS pour le respect du SIL exigé. L'intérêt de cette nouvelle approche basée sur les réseaux de fiabilité et les algorithmes génétiques réside dans le fait qu'elle peut être appliquée à la conception optimale de tout type de structure complexe pouvant être modélisée par des réseaux de fiabilité. Cette méthode sera étendue par la suite à la conception des SIS avec des composants ayant des taux de défaillance imprécis. En effet, dans le cadre de la conception des SIS, les fiabilistes assignent aux SIS à réaliser des objectifs de sûreté de fonctionnement (PFD_{avg} ou disponibilité moyenne A_{avg}) dès la phase d'expression du besoin en réduction de risque, afin d'une part d'aider le concepteur à rationaliser ses choix de composants et d'autre part de garantir à l'exploitant les objectifs de sûreté de fonctionnement exigés. Il est devenu donc primordial de trouver une stratégie d'allocation des objectifs de sûreté de fonctionnement qui permette d'établir le meilleur compromis entre le coût de conception du SIS et les objectifs de sûreté de fonctionnement qui seront exprimés sous la forme du SIL requis. Dans nos travaux, le problème de conception des SIS est ramené à un problème de minimisation du coût global du SIS sous contrainte du SIL exigé qui est exprimé en fonction de la disponibilité moyenne du SIS. La conception optimale du SIS se fait donc par le choix d'un certain nombre et type de composants dans chaque sous-système et des connexions entre ces différents composants du SIS qui garantissent un coût global minimal et le respect du SIL exigé. Ce problème d'allocation fait partie des problèmes d'allocation de redondance et donc des problèmes d'allocation de fiabilité en général. Cependant, comme l'a souligné Elegbede [Elegbede, 2000], les problèmes d'allocation de redondance ont traité principalement les systèmes à structure série parallèle et parallèle série. Cette restriction est due au fait que pour ce type de systèmes nous disposons d'une expression explicite de la fiabilité totale qui facilite le processus d'optimisation, ce qui n'est pas le cas pour des structures complexes. C'est pourquoi nous proposons une nouvelle approche qui est appliquée aux structures complexes. Cette approche sera étendue par la suite pour prendre en compte les informations imprécises sur les données de fiabilité des composants utilisés.

3 Organisation du rapport de thèse

Dans le premier chapitre, nous présentons le cadre global de nos développements : l'étude de la sûreté de fonctionnement des systèmes dédiés aux applications de sécurité. Nous étudions les problèmes soulevés par cette étude dans le contexte industriel et nous formulons l'objet de notre travail. Ce chapitre est composé de cinq parties. Nous présentons d'abord la problématique de la sécurité dans les installations industrielles. Nous proposons ensuite une synthèse des notions de sécurité fonctionnelle et des systèmes relatifs aux applications de sécurité. Nous présentons les principales normes de sécurité utilisées pour la conception de ces systèmes. Nous insistons en particulier sur les difficultés auxquelles les industriels sont confrontés lors du déploiement des deux normes de sécurité IEC 61508 [IEC61508, 1998] et IEC 61511 [IEC61511, 2000]. Nous exposons la problématique de l'évaluation de la sûreté de fonctionnement des systèmes instrumentés de sécurité (SIS). Enfin, dans une conclusion partielle, nous proposons les orientations de travail que nous avons choisies.

Dans le deuxième chapitre, nous présentons un état de l'art des méthodes de sûreté de fonctionnement et notamment lorsque les connaissances sont imparfaites (taux de défaillance, taux de réparation, intervalles de test, etc.). Nous nous concentrons principalement sur les méthodes de fiabilité PROBIST floue [Cai *et al.*, 1991] qui reposent sur les deux hypothèses fondamentales :

- Hypothèse probabiliste floue : le fonctionnement du système est complètement caractérisé par des mesures de probabilités floues.
- Hypothèse d'états binaires : le système ne peut avoir que deux états de fonctionnement ; l'état de défaillance ou l'état de fonctionnement normal.

Ce chapitre permet, au travers de la définition des principales théories utilisées pour la représentation des connaissances imparfaites (théorie des probabilités, théorie des ensembles flous, théorie des possibilités et théorie de l'évidence), d'introduire et d'expliquer leur utilisation (en particulier la théorie des ensembles flous) dans les méthodes d'évaluation de la sûreté de fonctionnement des systèmes complexes. Nous étudierons la méthode des arbres de défaillance flous [Tanaka *et al.*, 1983] que nous avons choisis pour évaluer les SIL dans nos travaux, et qui permet, en outre, de prendre en compte les incertitudes relatives aux paramètres de fiabilité et de caractériser l'incertitude des résultats obtenus. Puis, nous étudierons les bases théoriques concernant l'étude des chaînes de Markov floues introduites par Bellman et Zadeh [Bellman and Zadeh, 1970] et les chaînes de Markov à transitions floues [Zadeh, 1998]. Enfin, nous concluons par les Réseaux de Petri flous (RdPF) [Cardoso *et al.*, 1990; Valette *et al.*, 1989; Chen *et al.*, 1990; Lipp and Gunter, 1994] qui sont adaptées à la représentation des connaissances incertaines par des variables floues associées aux places, aux transitions, aux jetons, et au marquage.

Le troisième chapitre est consacré à la présentation d'une approche probabiliste floue pour déterminer le SIL d'un SIS à partir des paramètres de fiabilité imprécis de ses composants. L'introduction de deux nouveaux facteurs d'importance permettra de mettre en évidence les composants contribuant le plus à l'incertitude entachant le SIL obtenu et de la réduire efficacement. En effet :

- Le premier facteur que nous proposons permet d'identifier les composants critiques du système du point de vue de la fiabilité et/ou de la disponibilité en présence d'incertitudes.
- Le deuxième facteur d'incertitude permet d'identifier les composants dont l'incertitude de la probabilité de défaillance contribue significativement à l'incertitude entachant la probabilité de défaillance du système complet.

Nous présenterons une application tirée du document technique ISA-TR84.00.02-2002 [ISA-TR84.00.02-2002, 2002] qui illustre l'utilisation de la méthodologie proposée et met en évidence l'intérêt des résultats obtenus avec cette approche. Dans cette application, une Fonction Instrumentée de Sécurité (SIF) de SIL i doit être implémentée dans un SIS pour réduire le taux de rejet d'un réservoir sous pression. Notre objectif est de vérifier si le SIS proposé au concepteur est capable de satisfaire à l'exigence du SIL i requis pour réaliser la SIF de réduction du rejet des gaz.

Le quatrième chapitre concerne la conception optimale des SIS dont on exige un SIL donné. Dans cet objectif, nous proposons une approche basée sur les réseaux de fiabilité [Berge, 1958; Kaufmann *et al.*, 1975; Kaufmann, 1972; Kaufmann, 1968] et les algorithmes génétiques pour la conception optimale des SIS. Nous ramenons le problème de conception des SIS à un problème de minimisation du coût global du SIS sous contrainte du SIL exigé qui est exprimé en fonction de la disponibilité moyenne du SIS. Pour introduire la problématique de conception optimale des SIS, nous commencerons par donner un bref état de l'art des méthodes d'allocation de redondance pour le respect de contraintes de coût et de fiabilité. Ensuite, deux méthodes de modélisation des SIS seront introduites : les Blocs diagrammes de Fiabilité (BdF) [Kleinerman and Weiss, 1954; Blanton, 1957; Villemeur, 1998] et les réseaux de fiabilité [Kaufmann, 1972; Kaufmann, 1968; Kaufmann *et al.*, 1975]. Pour procéder à l'optimisation de la structure des SIS, nous allons d'abord les modéliser par des réseaux de fiabilité. Nous utiliserons ensuite un algorithme génétique personnalisé pour l'optimisation de la structure du réseau de fiabilité et l'obtention des configurations optimales des SIS. Une application illustrera la méthodologie proposée et donnera les configurations optimales obtenues pour les SIL exigés sous forme de BdF dans le but de faciliter leur implémentation. En outre, nous proposerons un autre modèle d'aide à la conception optimale des SIS avec des composants à taux de défaillance imprécis modélisés par des nombres flous. Dans cette approche, nous donnerons aussi les configurations optimales obtenues pour les SIL exigés.

Ces travaux de thèse sont clos par une conclusion générale repositionnant le cadre d'étude, les apports des travaux menés et les perspectives intéressantes qu'ouvre le travail.

1

Evaluation de la sûreté de fonctionnement des systèmes dédiés aux applications de sécurité

Nous présentons dans ce chapitre le cadre global de nos développements : l'étude de la sûreté de fonctionnement des systèmes dédiés aux applications de sécurité. Nous étudions les problèmes soulevés par cette étude dans le contexte industriel et nous formulons l'objet de notre travail. Le chapitre est composé de cinq parties.

Nous présentons d'abord la problématique de la sécurité dans les entreprises. Nous proposons ensuite une synthèse des notions de sécurité fonctionnelle et des systèmes relatifs aux applications de sécurité. Nous présentons les principales normes de sécurité utilisées pour la conception de ces systèmes. Nous exposons la problématique de l'évaluation de la sûreté de fonctionnement des systèmes instrumentés de sécurité (SIS). Enfin, pour conclure, nous proposons les orientations de travail que nous avons choisies.

1.1 Problématique de la sécurité

1.1.1 Introduction

Malgré les efforts déployés par les entreprises, les accidents (incendies, explosions, etc.) restent nombreux en France et à l'étranger [Ministère de l'écologie, 2007]. Le tableau 1.1 présente un recueil des plus importants accidents industriels survenus dans le monde entre 1960 et 2007. L'ampleur et la fréquence de ces accidents ont suscité de nombreux efforts sur les études de risques afin de mieux les prévenir, les prévoir et les gérer.

Dans les études de sécurité, l'étude de dangers reste l'outil d'information privilégié. Il doit permettre d'identifier les dangers que peut présenter l'exploitation de l'installation et les moyens de réduire la probabilité et les effets des potentiels de dangers.

L'étude de dangers comprend en général quatre points essentiels :

- La description des installations et de leur environnement : elle comporte une description générale de l'installation et une description spécifique des équipements permettant la compréhension des scénarios d'accidents pour les différentes phases d'exploitation (normales, dégradées, de maintenance, de démarrage, d'arrêt, etc.).
- L'identification et la caractérisation des potentiels de dangers : l'analyse des accidents potentiels liés aux installations est établie par la conjonction d'événements élémentaires. En effet, les accidents majeurs résultent le plus souvent de la combinaison d'événements élémentaires, généralement peu graves en eux-mêmes, par exemple la survenue simultanée de deux pannes ou la conjonction d'une défaillance technique et d'une défaillance humaine. Il convient donc que l'étude de dangers apporte la preuve que les conjonctions d'événements simples ont bien été prises en compte dans l'identification des causes d'accidents majeurs. Ces conjonctions d'événements simples constituent les scénarios à envisager pour l'étude de dangers.
- L'analyse des accidents passés : cette analyse recense les événements relatifs à des accidents ou des incidents survenus sur le site et sur d'autres sites mettant en œuvre des installations, des produits et des procédés comparables.
- L'analyse de risques : L'analyse des installations porte sur toutes leurs conditions d'exploitation (phases transitoires et d'arrêt incluses). Elle est complétée par les aspects liés à l'accidentologie pour permettre de définir des scénarios d'accidents et d'en évaluer les conséquences.

L'étude de dangers doit aboutir à un ensemble de mesures de maîtrise des risques mises en œuvre à l'intérieur de l'installation, qui réduisent le risque à l'intérieur et à l'extérieur de l'installation à un niveau jugé "acceptable" par l'exploitant de l'installation.

Dans ce chapitre, nous nous intéressons aux fonctions liées à la sécurité et plus particulièrement aux systèmes E/E/PE (électriques/électroniques/électroniques programmables) qui sont utilisés pour exécuter ces fonctions. Nous étudions par la suite les normes de sécurité qui décrivent le cycle de vie de ces systèmes depuis la rédaction du cahier des charges, en passant

Date	Localisation	Type d'accident	Victimes et dégâts
1966	Feyzin-France	Incendie d'une industrie de stockage pétrochimique	18 morts
1974	Flixborough-Grande Bretagne	Explosion sur un site industriel	28 morts
1976	Seveso-Italie	Fuite de dioxine d'une usine chimique	37 000 personnes touchées
1984	Bhopal-Inde	Fuite d'un gaz toxique	2 500 morts et 250 000 blessés
1984	Mexico-Mexique	Explosion d'une citerne de pétrole liquéfié	plus de 500 morts et 7000 blessés
1986	Tchernobyl-Ukraine	Explosion d'une centrale nucléaire	plus de 15000 morts
1992	La Mède-France	Explosion dans une raffinerie	6 morts et 7 blessés
2001	Toulouse-France	Explosion d'un site industriel	30 morts et plus de 2 000 blessés

TAB. 1.1 – Recueil des plus graves accidents industriels survenus dans le monde entre 1960 et 2007 [Ministère de l'écologie, 2007].

par la conception, l'installation, l'exploitation et la maintenance, jusqu'à la mise hors service. Avant d'entrer plus avant dans la thématique de la sécurité fonctionnelle et des systèmes dédiés aux applications de sécurité, nous proposons dans les lignes qui suivent de nous attarder sur les définitions relatives aux notions clés de danger, de risque et de sécurité.

1.1.2 Notion de danger

1.1.2.1 Définitions

Selon [Desroches, 1995] et la norme de sécurité IEC 61508 [IEC61508, 1998], *le danger désigne une nuisance potentielle pouvant porter atteinte aux biens (détérioration ou destruction), à l'environnement, ou aux personnes. Les dangers peuvent avoir une incidence directe sur les personnes, par des blessures physiques ou des troubles de la santé, ou indirecte, au travers de dégâts subis par les biens ou l'environnement.*

Et selon le référentiel OHSAS 18001 [OHSAS18001, 1999] un danger est *une source ou une situation pouvant nuire par blessure ou atteinte à la santé, dommage à la propriété et à l'environnement du lieu de travail ou une combinaison de ces éléments.*

Soulignons que de nombreux termes sont employés, selon les normes ou les auteurs, autour de la notion de danger et la rendent ambiguë. De plus, les dictionnaires associent souvent le terme danger au terme risque. En effet, plusieurs dictionnaires proposent le terme risque comme synonyme du terme danger, ce qui explique le fait qu'un grand nombre de personnes utilisent indifféremment ces termes. Même les documents et les textes officiels confondent danger et risque.

1.1.2.2 Phénomènes dangereux et situations dangereuses

Un phénomène dangereux désigne en général une source potentielle de dommage. On regroupe sous cette appellation l'ensemble des sources et des facteurs pouvant contribuer à la création du dommage. Ainsi, un bord coupant est un élément dangereux, mais cela ne provoquera pas obligatoirement un dommage.

Une situation dangereuse désigne une situation dans laquelle des personnes, des biens ou l'environnement sont exposés à un ou plusieurs phénomènes dangereux.

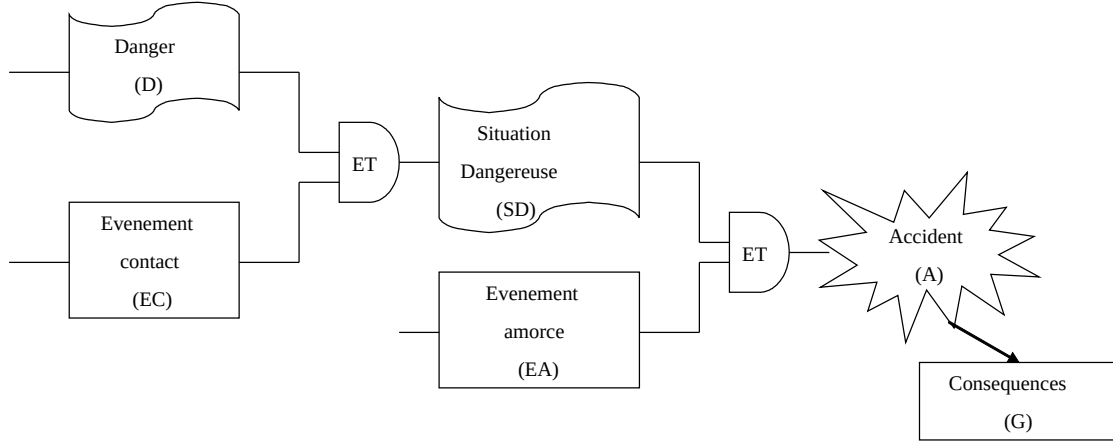


FIG. 1.1 – Scénario d'accident [Desroches, 2005]

Selon [Desroches, 2005], la situation dangereuse (ou à risque), notée (SD), résulte de la conjonction d'un danger (D) et d'un événement contact (EC) qui met le système en présence ou au contact du danger.

Cette configuration de l'état du système en présence de danger est modélisée par l'expression :

$$SD = D \cap EC \quad (1.1)$$

La situation à risque est définie par la nature et le potentiel de dangerosité et la vraisemblance de ce potentiel. De même l'événement redouté final ou accident (A) résulte de la conjonction de la situation dangereuse et d'un événement amorce (EA) qui déclenche la dangerosité sur le ou les éléments vulnérables du système. Cette configuration accidentelle est modélisée par l'expression :

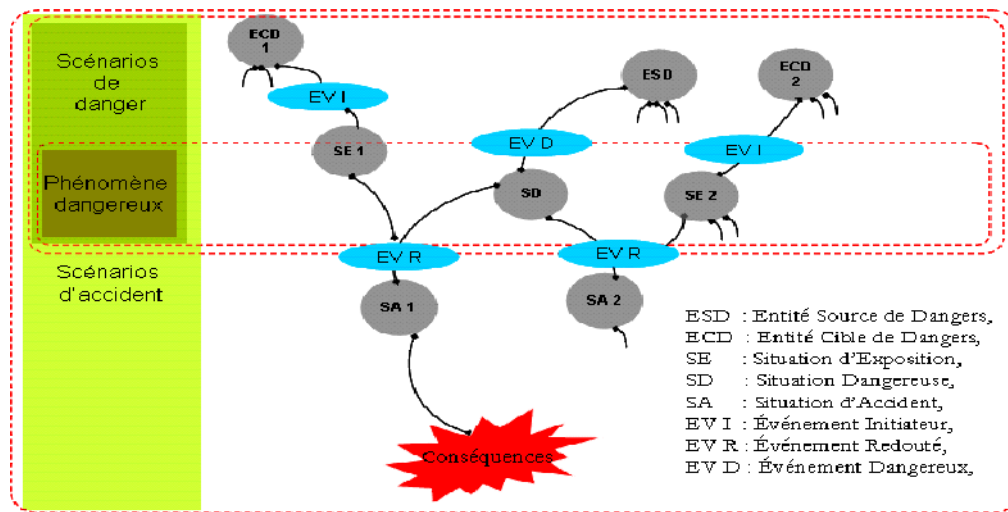
$$A = SD \cap EA = D \cap EC \cap EA \quad (1.2)$$

La gravité des conséquences directes et indirectes de l'accident, notée $G(A)$ correspond au montant des dommages en terme de perte ou préjudice mesurable. Le scénario d'accident qui en résulte est visualisé sur le diagramme de la figure 1.1.

Selon Mazouni *et al.* [Mazouni *et al.*, 2007], une situation dangereuse se caractérise par le développement d'un phénomène dangereux dès l'apparition d'un événement dangereux. Un phénomène dangereux est une libération de tout ou partie d'un potentiel de danger. Cette concrétisation produit des effets (dispersion d'un nuage de gaz toxique, dérapage d'une voiture, etc.). Ces effets n'entraînent des dommages que si des entités cibles de danger s'exposent, dans un espace de danger pendant une durée suffisante au phénomène dangereux. Le diagramme état/transition du processus accidentel est représenté sur la figure 1.2.

1.1.3 Notion de risque

La perception des dommages potentiels liés à une situation dangereuse se rapporte à la notion de risque. Le terme risque a plusieurs significations. De même, les risques peuvent être de nature très variée et beaucoup de classifications ont été proposées.

FIG. 1.2 – Diagramme état/transition du processus accidentel [Mazouni *et al.*, 2007]

Les définitions du risque à deux dimensions sont assez proches. Selon [Villemeur, 1998], le risque est *une mesure d'un danger associant une mesure de l'occurrence d'un événement indésirable et une mesure de ses effets ou conséquences*.

Et selon le référentiel OHSAS 18001 [OHSAS18001, 1999], *un risque est la combinaison de la probabilité et de la (des) conséquence(s) de la survenue*.

Cependant, il existe des définitions légèrement plus complexes que celle de [Villemeur, 1998] et du référentiel OHSAS 18001 [OHSAS18001, 1999] dans lesquelles apparaît une troisième dimension : l'acceptabilité du risque, seuil en dessous duquel on accepte l'existence du danger bien que sa gravité et sa probabilité d'occurrence ne soient pas nulles. Par ailleurs, il faut noter que Périlhon [Périlhon, 1998] propose une méthode logique pour mettre en évidence la majorité des risques d'une installation : la méthode organisée systémique d'analyse des risques ou MOSAR. Elle fait appel à une décomposition de l'installation en sous-systèmes et une recherche systématique des dangers présentés par chacun d'entre eux, ces sous-systèmes sont remis en relation pour faire apparaître des scénarios de risques majeurs. La démarche peut se poursuivre par une analyse détaillée de type sûreté de fonctionnement. Elle se termine sur la construction des plans d'intervention.

Dans la suite de nos travaux, le terme risque est lié sans ambiguïté aux risques encourus dans la conduite des systèmes et sont abordés par les études de sûreté de fonctionnement.

1.1.3.1 Mesure de risque

Qualitativement, le risque se caractérise par :

- L'ampleur des dommages, suite à un événement redouté, selon un critère de gravité (critique, marginal, mineur, insignifiant, etc.). Ce critère tient compte de l'appréciation des conséquences en terme de pertes humaines (morts, blessures) ou en termes économiques (coût liés aux dégradations, etc.).

Fréquence d'événement	Gravité du dommage				
	Catastrophique	Majeure	Mineure	Minime	Négligeable
Fréquent	H	H	H	H	I
Probable	H	H	H	I	I
Occasionnel	H	H	I	I	L
Rare	H	I	I	L	T
Improbable	I	I	L	T	T
Invraisemblable	I	L	T	T	T

TAB. 1.2 – Table d'estimation du risque [Guiochet, 2003]

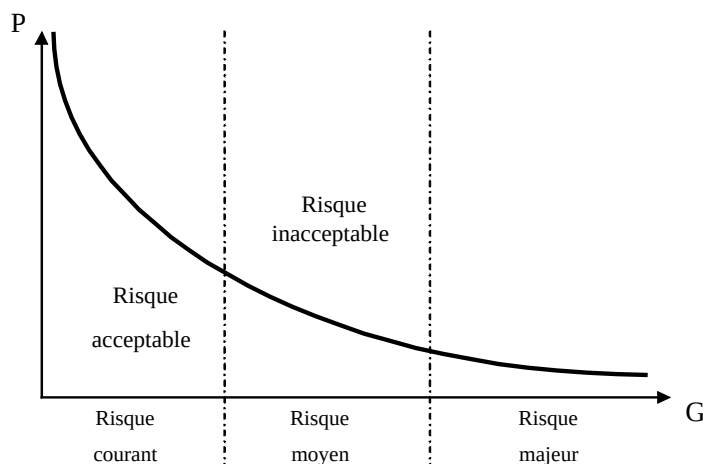


FIG. 1.3 – Courbe de Farmer [Farmer, 1967]

- Le caractère incertain lié à l'apparition d'un événement redouté provoquant le dommage, depuis une situation dangereuse déterminée (fréquent, rare, improbable, etc.).

Le tableau 1.2 représente un exemple de table d'estimation du risque. La lettre H représente le risque élevé, la lettre I représente le risque intermédiaire, la lettre L représente le risque faible et la lettre T représente le risque insignifiant. Les couples (fréquence, gravité) résultant en un risque élevé (H) sont mis en valeur car ils représentent les risques les moins tolérables. Il est clair qu'un risque de niveau H fera sans aucun doute l'objet d'un travail de réduction.

En général, le risque se rapporte au couple (gravité, probabilité), le plus souvent cela quantifie le produit de la gravité d'un accident par sa probabilité d'occurrence. Farmer a montré que la gravité et la probabilité d'apparition étaient reliées de façon à peu près linéaire [Farmer, 1967]. En effet, si nous représentons les accidents sur une grille (gravité (G), probabilité (P)), nous obtenons la courbe représentée dans la figure 1.3. Tous les risques classés sous la courbe sont considérés comme acceptables, à l'inverse, ils sont considérés comme inacceptables s'ils sont placés au-dessus de la courbe.

Selon Gouriveau [Gouriveau, 2003], le risque peut être défini par l'association d'événements causes et conséquences d'une situation donnée. Les événements causes peuvent être caractérisés

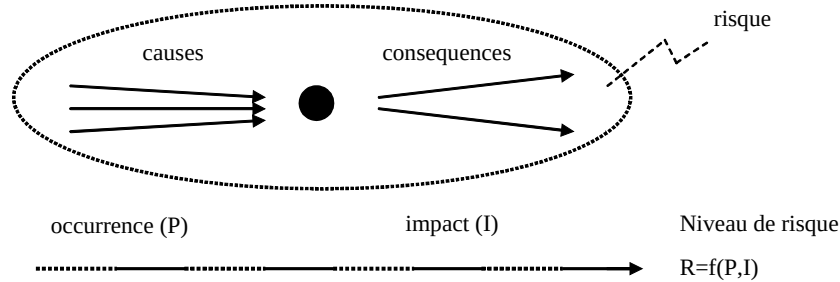


FIG. 1.4 – Caractérisation du risque [Gouriveau, 2003]

par leur occurrence (P) et les événements effets par leur impact (I) (cf. figure 1.4). La corrélation de ces grandeurs permet de construire un indicateur de risque $R = f(\text{occurrence}, \text{impact})$.

1.1.3.2 Classification des risques

Dans la littérature, on trouve plusieurs classifications des risques. Selon [Tanzi and Delmer, 2003], l'analyse des risques permet de les classer en cinq grandes familles :

- les risques naturels : inondation, feu de forêt, avalanche, tempête, séisme, etc. ;
- les risques technologiques : d'origine anthropique, ils regroupent les risques industriels, nucléaires, biologiques, ruptures de barrage, etc. ; les risques de transports collectifs (personnes, matières dangereuses) sont aussi considérés comme des risques technologiques.
- les risques de la vie quotidienne : accidents domestiques, accidents de la route, etc. ;
- les risques liés aux conflits.

Une des classifications les plus répandue est de classer les risques en deux catégories ; les risques naturels et les risques liés à l'activité humaine.

Selon cette classification, les risques peuvent être naturels dans le sens où ils ont trait à un événement sans cause humaine directe avérée. Les causes directes supposées ou indirectes ne doivent pas modifier cette distinction.

Les risques liés à l'activité humaine recouvrent un ensemble de catégories de risques divers :

- les risques techniques, technologiques, industriels, nucléaires ;
- les risques liés aux transports ;
- les risques sanitaires ;
- les risques économiques, financiers, managériaux ;
- les risques médiatiques.
- les risques professionnels.

Finalement, nous pouvons caractériser l'opération de classification des risques par les deux types de classement suivants :

- Classement "subjectif" fait par des individus à partir de l'idée qu'ils se font du risque en se fondant sur leur expérience et leurs connaissances ou "objectif" fait à partir de données

statistiques, d'enquêtes, etc. ;

- Classement "qualitatif" caractérisé par l'établissement d'un système d'ordre comparatif ou "quantitatif" basé sur le calcul de probabilités.

1.1.3.3 Risque acceptable

La notion de risque est essentielle pour caractériser la confiance attribuée à un système. En effet, si nous admettons souvent comme potentiels des dommages sévères, seule leur faible probabilité d'occurrence nous les font accepter [Lamy *et al.*, 2006; Simon *et al.*, 2007]. Par exemple, nous continuons à prendre l'avion malgré les accidents possibles du fait que la probabilité d'un écrasement conduisant aux décès des passagers est extrêmement faible. Nous établissons généralement cet arbitrage en fonction des risques que nous encourront par ailleurs, comme ceux induits par des phénomènes naturels : tremblements de terre, avalanches, inondations, etc.

Selon le référentiel OHSAS 18001 [OHSAS18001, 1999], *le risque acceptable est un risque qui a été réduit à un niveau tolérable pour un organisme en regard de ses obligations légales et de sa propre politique de santé et de sécurité au travail.*

Selon le Guide 51 ISO/IEC [ISO, 1999], le risque acceptable est *un risque accepté dans un contexte donné basé sur des valeurs courantes de notre société.*

Notons que l'acceptabilité concerne le risque et non la gravité du dommage ou sa probabilité d'occurrence considérées séparément. Ces définitions soulignent également le fait que l'acceptabilité dépend de valeurs courantes de notre société souvent fondées sur des données associées à des phénomènes naturels. Ainsi, nous acceptons de prendre le risque de mourir en prenant l'avion si la probabilité de ce décès par cette cause est identique voire inférieure à la probabilité de décès induit par un séisme ou une crise cardiaque (pour un corps sain) [Beugin, 2006].

Les définitions précisent par ailleurs que l'acceptabilité est fonction du contexte. Ce contexte peut tout d'abord caractériser l'état d'un savoir sur des pratiques ou sur une technologie de mise en œuvre.

1.1.3.4 Risque majeur

D'une manière générale, le risque majeur se caractérise par ses nombreuses victimes, un coût de dégâts matériels, des impacts sur l'environnement [Tanzi and Delmer, 2003]. Le risque majeur est caractérisé par deux critères qui définissent sa fréquence et sa gravité :

- une faible fréquence.
- une énorme gravité (nombreux morts et blessés).

1.1.3.5 Risque industriel

Le risque industriel se caractérise par un accident se produisant sur un site industriel et pouvant entraîner des conséquences graves pour le personnel, les populations, les biens, l'environnement ou le milieu naturel.

Les actions de prévention des risques industriels majeurs s'appuient sur la connaissance des phénomènes redoutés notamment au travers de l'étude de dangers. Cette prévention s'articule

autour de quatre axes principaux.

- La maîtrise du risque à la source ;
- La maîtrise de l’urbanisation ;
- L’organisation des secours ;
- L’information préventive et la concertation.

1.1.4 Notion de sécurité

La sécurité est souvent définie par son contraire : elle serait l’absence de danger, de risque, d’accident ou de sinistre.

Selon [Desroches *et al.*, 2003], la sécurité concerne *la non occurrence d’événements pouvant diminuer ou porter atteinte à l’intégrité du système, pendant toute la durée de l’activité du système, que celle-ci soit réussie, dégradée ou ait échouée.*

Et suivant le guide ISO/CEI 73 [ISO, 2002] élaboré par l’ISO (organisation internationale de normalisation) sur la terminologie du management du risque, la sécurité est *l’absence de risque inacceptable, de blessure ou d’atteinte à la santé des personnes, directement ou indirectement, résultant d’un dommage au matériel ou à l’environnement.*

Signalons que dans le cadre des installations industrielles, on parle de sécurité des installations vis-à-vis des accidents et de sûreté vis-à-vis des attaques externes volontaires (type malveillance ou attentat), des intrusions malveillantes et de la malveillance interne.

1.1.5 Principes généraux de protection

Nous pouvons distinguer les mesures de sécurité par leur mode d’action : les sécurités passives et les sécurités actives.

1.1.5.1 Sécurités passives

La sécurité passive désigne tous les éléments mis en jeu afin de réduire les conséquences d’un accident lorsque celui-ci n’a pu être évité. Elle agit par sa seule présence, sans intervention humaine ni besoin en énergie (exemple : bâtiment de confinement, cuvette de rétention, etc.). Cependant, il ne faut pas réduire la sécurité passive à la limitation des conséquences des accidents (l’isolation électrique est une mesure passive et préventive).

1.1.5.2 Sécurités actives

La sécurité active désigne tous les éléments mis en jeu afin d’éviter les accidents. Elle nécessite une action, une énergie et un entretien (exemple : détecteur, vannes, etc.).

La sécurité d’une installation repose sur l’utilisation de ces deux modes d’action. Une préférence est donnée au mode passif quand il est techniquement possible. Des critères de qualité sont exigés pour le mode actif, notamment la tolérance à la première défaillance : doublement de l’organe de sécurité (redondance). La sécurité fonctionnelle reste l’un des moyens les plus importants pour la prise en compte des risques. D’autres moyens de réduction ou d’élimination

des risques, tels que la sécurité intégrée dans la conception, sont également d'une importance essentielle.

1.2 Sécurité fonctionnelle

1.2.1 Définitions

Selon la norme IEC 61061 [IEC61061, 1998], la sécurité fonctionnelle est le *sous-ensemble de la sécurité globale se rapportant à la machine et au système de commande de la machine qui dépend du fonctionnement correct des systèmes électriques de commande relatifs à la sécurité, des systèmes relatifs à la sécurité basés sur une autre technologie et des dispositifs externes de réduction de risque.*

Suivant la norme IEC 61508 [IEC61508, 1998], la sécurité fonctionnelle est le *sous-ensemble de la sécurité globale qui dépend du bon fonctionnement d'un système ou d'un équipement en réponse à ses entrées.*

La sécurité fonctionnelle veille donc à contrôler l'absence de risques inacceptables qui pourraient :

- Engendrer des blessures ;
- Porter atteinte, directement ou indirectement, à la santé des personnes ;
- Dégrader l'environnement ;
- Altérer la propriété.

La sécurité fonctionnelle couvre les produits ou systèmes mettant en œuvre des solutions de protection fondées sur diverses technologies :

- Mécanique, hydraulique, pneumatique, électrique, électronique, électronique programmable, optique, etc.
- Ou toute combinaison de ces technologies.

1.2.2 Systèmes relatifs aux applications de sécurité

Un système E/E/PE (électrique/électronique/électronique programmable de sécurité) relatif aux applications de sécurité comprend tous les éléments du système nécessaires pour remplir la fonction de sécurité. C'est-à-dire, depuis le capteur, en passant par la logique de contrôle et les systèmes de communication, jusqu'à l'actionneur final, tout en incluant les actions critiques de l'opérateur.

Les systèmes de sécurité sont définis en termes d'absence de risque inacceptable de blessure ou de préjudice à la santé des personnes. Les dommages aux personnes peuvent être directs ou indirects, comme des dommages aux biens ou à l'environnement par exemple. Certains systèmes peuvent être principalement conçus pour se prémunir contre des pannes ayant des implications économiques majeures. Ceci signifie que dans l'esprit, à objectifs techniques comparables ou identiques, il n'y a pas de différence entre un système de sécurité et un système de contrôle commande. L'IEC 61508 [IEC61508, 1998] et l'IEC 61511 [IEC61511, 2000] peuvent donc être

utilisées pour développer n'importe quel système E/E/PE comportant des fonctions critiques, telles que la protection des équipements, des biens ou de la productivité.

1.3 Référentiel normatif

1.3.1 Introduction

La sécurité fonctionnelle a depuis longtemps retenu l'attention des industriels. Pour mener à bien leur démarche sécurité, ils peuvent s'appuyer sur des normes. La norme internationale de sécurité IEC 61508 [IEC61508, 1998] est une des dernières normes dédiées à la sécurité fonctionnelle. Elle est devenue une norme française en 1999. Les normes filles que cette norme de base a générées, sont plus récentes et restent encore assez peu connues des acteurs de la sécurité dans certains secteurs industriels français. La norme ANSI/ISA S84.01 [ANSI/ISA-S84.01-1996, 1996] s'identifie à la norme dérivée IEC 61511 [IEC61511, 2000] applicable au secteur de l'industrie des procédés. Cet ensemble normatif s'impose comme la référence pour le développement, la mise en œuvre et l'exploitation des systèmes relatifs aux applications de sécurité.

Ces normes de sécurité fonctionnelle introduisent une approche probabiliste qui vient compléter l'approche déterministe classique. Elles préconisent l'utilisation de certaines méthodes (équations simplifiées, arbres de défaillance, approches markoviennes) sans pour autant les rendre obligatoire.

1.3.2 Norme IEC 61508

1.3.2.1 Domaine d'application

La norme IEC 61508 [IEC61508, 1998] est une norme internationale qui porte plus particulièrement sur les systèmes E/E/PE, c'est-à-dire les systèmes électriques/électroniques/électroniques programmables de sécurité. La norme propose une approche opérationnelle pour mettre en place un système de sécurité E/E/PE, en partant de l'étude des exigences de sécurité (avec une définition du périmètre couvert, une analyse et une évaluation du risque) et en prenant en compte toutes les étapes du cycle de vie du système E/E/PE. Un des intérêts de cette norme est d'être générique et donc d'être applicable dans tous les secteurs où la sécurité peut être traitée avec des systèmes E/E/PE : industries manufacturières, industries des process continus, pharmaceutiques, nucléaires, ferroviaires, etc.

1.3.2.2 Structure générale de la norme

Cette norme comprend 7 parties (cf. figure 1.5), afin de couvrir les multiples aspects des systèmes E/E/PE :

- 61508-1 : Prescriptions générales.
- 61508-2 : Prescriptions propres aux systèmes E/E/PE.
- 61508-3 : Prescriptions relatives au logiciel.
- 61508-4 : Définitions et abréviations.
- 61508-5 : Exemples de méthodes pour déterminer le niveau d'intégrité de la sécurité.
- 61508-6 : Guides pour l'application des parties 2 et 3 de la norme.
- 61508-7 : Tour d'horizon des techniques et des mesures.

Signalons que la partie 6 est seulement informative. Les méthodes qui y sont données ne sont pas obligatoires.

1.3.2.3 Objectifs de la norme

L'IEC 61508 [IEC61508, 1998] a pour but de :

- Fournir le potentiel des technologies E/E/PE pour améliorer à la fois les performances économiques et de sécurité,
- Permettre des développements technologiques dans un cadre global de sécurité,
- Fournir une approche système, techniquement saine, suffisamment flexible pour le futur,
- Fournir une approche basée sur le risque pour déterminer les performances des systèmes concernés par la sécurité,
- Fournir une norme générique pouvant être utilisée par l'industrie, mais qui peut également servir à développer des normes sectorielles (par exemple : machines, usines chimiques, ferroviaires ou médicales) ou des normes produit (par exemple : variateurs de vitesse),
- Fournir les moyens aux utilisateurs et aux autorités de réglementation d'acquiescer la confiance dans les technologies basées sur l'électronique programmable,
- Fournir des exigences basées sur des principes communs pour faciliter :
 - une compétence améliorée de la chaîne d'approvisionnement des fournisseurs de sous systèmes et de composants à des secteurs variés,
 - des améliorations de la communication et des exigences (c'est-à-dire de clarifier ce qui doit être spécifié),
 - le développement de techniques et de mesures pouvant être utilisées par tous les secteurs, augmentant de ce fait la disponibilité des ressources,
 - le développement des services d'évaluation de la conformité si nécessaire.

1.3.2.4 Evaluation du niveau d'intégrité de sécurité (SIL)

La norme IEC 61508 [IEC61508, 1998] fixe le niveau d'intégrité de sécurité (SIL) qui doit être atteint par un SIS qui réalise la Fonction Instrumentée de Sécurité (SIF). Elle donne le SIL en fonction de sa probabilité de défaillance moyenne (PFD_{avg}) sur demande pour les SIS faiblement sollicités (moins d'une sollicitation par an) (cf. tableau 1.3) ou en fonction de probabilité de défaillance par heure (PFH) pour les SIS fortement sollicités ou agissant en mode continu (cf. tableau 1.4). L'allocation du SIL se fait par des méthodes qualitatives et semi quantitatives, alors que l'évaluation du PFD_{avg} des SIS qui doivent satisfaire au SIL exigé se fait par des méthodes quantitatives.

SIL	Probabilité moyenne de défaillance à la sollicitation (PFD_{avg})
1	$[10^{-2}, 10^{-1}[$
2	$[10^{-3}, 10^{-2}[$
3	$[10^{-4}, 10^{-3}[$
4	$[10^{-5}, 10^{-4}[$

TAB. 1.3 – Niveaux d'intégrité de sécurité (SIL) : Sollicitation faible [IEC61508, 1998]

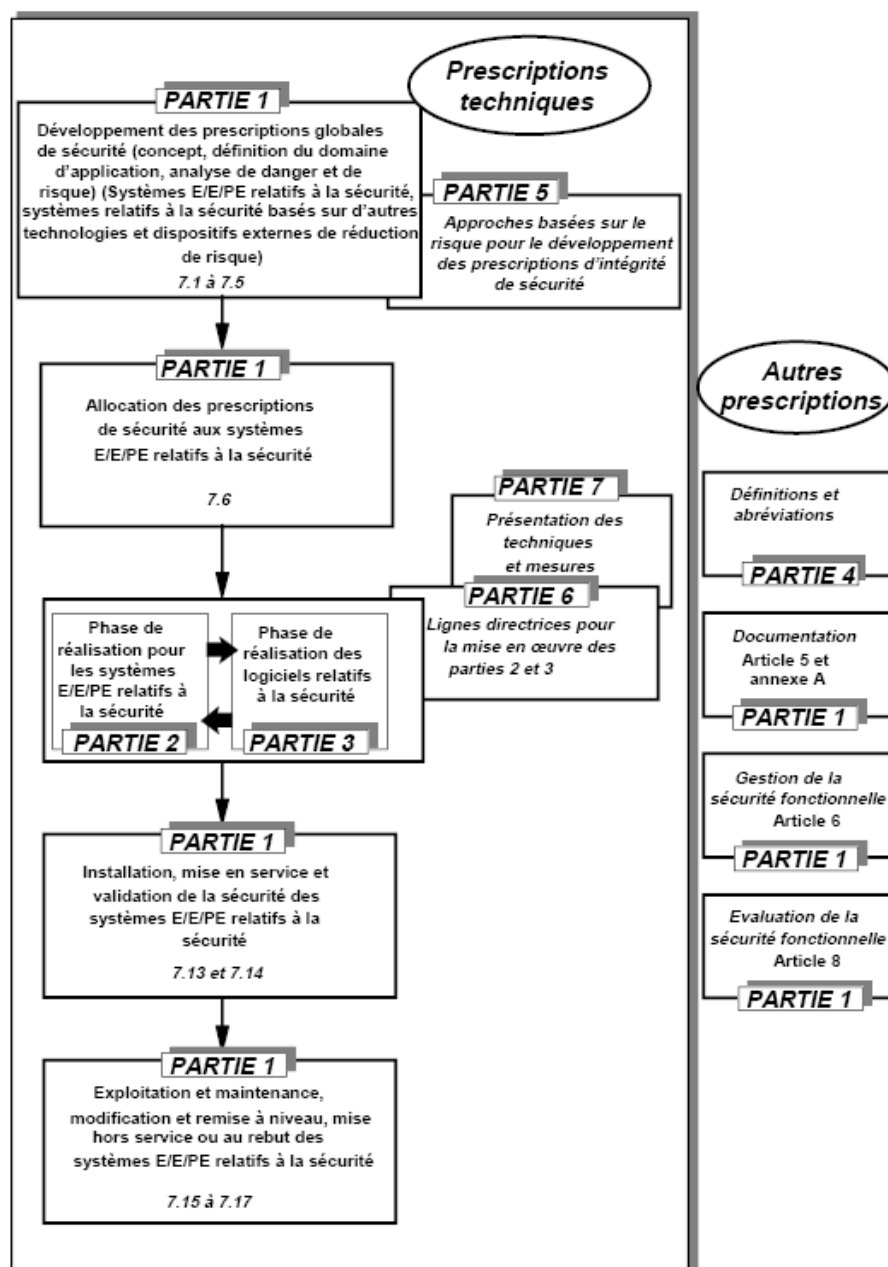


FIG. 1.5 – Structure générale de la norme IEC 61508 [IEC61508, 1998]

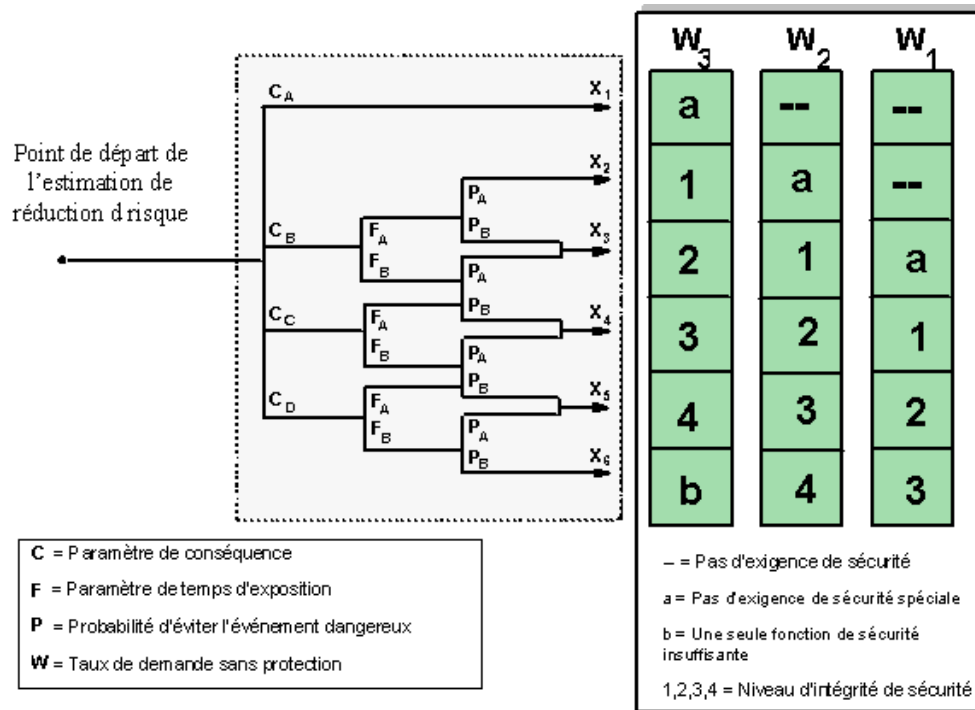


FIG. 1.6 – Graphe de risque [IEC61508, 1998]

SIL	Fréquence des défaillances dangereuses par heure (N)
1	$[10^{-6}, 10^{-5}[$
2	$[10^{-7}, 10^{-6}[$
3	$[10^{-8}, 10^{-7}[$
4	$[10^{-9}, 10^{-8}[$

TAB. 1.4 – Niveaux d'intégrité de sécurité (SIL) : Sollicitation élevée ou continue [IEC61508, 1998]

1. Les méthodes qualitatives [IEC61508, 1998; Stavrianidis and Bhimavarapu, 1998; McCormick, 1981; DIN19250, 1990] : il s'agit de méthodes qui permettent d'allouer le SIL à partir de la connaissance des risques associés au procédé. La méthode la plus utilisée est le graphe de risque [Simon *et al.*, 2006; Simon *et al.*, 2007] (cf. figure 1.6).

Le graphe de risque s'appuie sur l'équation suivante :

$$R = f * C \quad (1.3)$$

où R est le risque en l'absence de systèmes relatifs à la sécurité, f est la fréquence de l'événement dangereux en l'absence de systèmes relatifs à la sécurité et C est la conséquence de l'événement dangereux.

La fréquence de l'événement dangereux f est supposée être le résultat des trois facteurs suivants (cf. figure 1.6) :

- Fréquence et durée d'exposition dans une zone dangereuse ;
- La possibilité d'éviter l'événement dangereux ;
- La probabilité que l'événement dangereux se produise en l'absence de systèmes relatifs à la sécurité. C'est ce qu'on appelle la probabilité d'occurrence non souhaitée.

Finalement, nous obtenons les quatre paramètres de risque suivants :

- Conséquence de l'événement dangereux (C) ;
- Fréquence et durée d'exposition au danger (F) ;
- Possibilité d'éviter l'événement dangereux (P) ;
- Probabilité de l'occurrence non souhaitée (W).

En combinant les paramètres de risque décrits ci-dessus, on obtient le graphe de risque présentée à la figure 1.6. Ce graphe est tiré de la norme IEC 61508 [IEC61508, 1998].

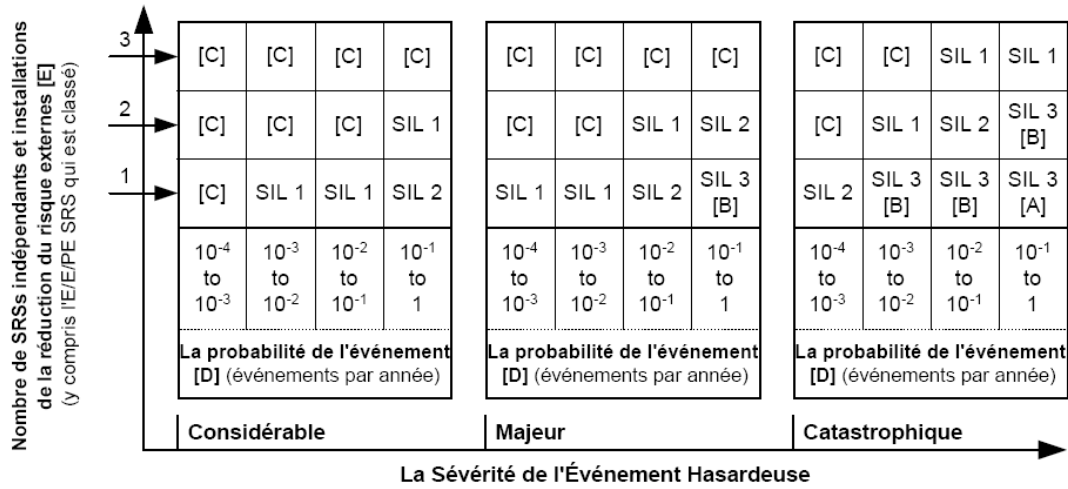
Il faut souligner que les paramètres de risque ne sont donnés que d'une manière indicative et que d'autre paramètres doivent ou peuvent être choisis selon l'application étudiée.

Ce graphe s'explique de la manière suivante. L'utilisation des paramètres de risque C , F et P aboutit à un certain nombre de sorties ($X_1, X_2, \dots, X_n$). Chaque sortie est consignée dans une des trois échelles (W_1, W_2 et W_3). Chaque échelle indique le niveau de SIL nécessaire auquel doit satisfaire le système relatif à la sécurité pris en considération. La mise en correspondance avec W_1, W_2 ou W_3 permet de réaliser la contribution d'autres mesures de réduction du risque. Le décalage dans les échelles W_1, W_2 et W_3 est nécessaire pour avoir trois niveaux différents de réduction des risques à partir d'autres mesures. Cette échelle est composée de l'échelle W_3 , qui fournit la réduction minimale du risque grâce à d'autres mesures, c'est-à-dire la plus forte probabilité de l'apparition d'un événement non désiré. L'échelle W_2 caractérise une contribution moyenne et l'échelle W_1 caractérise une contribution maximale. La sortie finale du graphe de risque donne le niveau de SIL du SIS (c'est-à-dire 1, 2, 3 ou 4) et correspond à une mesure de la réduction nécessaire du risque pour le système.

2. Les méthodes semi quantitatives [IEC61508, 1998; CCPS, 1993; McCormick, 1981; Stavrianidis, 1995; Bhimavarapu *et al.*, 1997] : la méthode la plus répandue est la matrice de risque (cf. figure 1.7). Cette matrice donne le niveau de SIL en fonction de la gravité de risque et de sa fréquence d'occurrence.

La matrice de risque est présentée dans la norme IEC 61508 [IEC61508, 1998] comme une méthode qui considère trois questions pour arriver à l'estimation du SIL exigé :

- Le paramètre du risque de la conséquence.
- Le paramètre du risque de la fréquence.
- Le nombre de paramètre de fonctions de protection indépendantes.



- [A] Un SIL 3 E/E/PE le système sécurité-apparenté ne fournit pas réduction du risque suffisante à ce niveau du risque. Les mesures de la réduction du risque supplémentaires sont exigées.
- [B] Un SIL 3 E/E/PE le système sécurité-apparenté ne peut pas fournir réduction du risque suffisante à ce niveau du risque. Le hasard et risque l'analyse est exigée pour déterminer si les mesures de la réduction du risque supplémentaires sont nécessaires.
- [C] Un E/E/PE indépendant que le système sécurité-apparenté n'est pas exigé probablement.
- [D] La probabilité de l'événement est la probabilité que l'événement hasardeux se produit sans toute sécurité systèmes apparentés ou installations de la réduction du risque externes.
- [E] SRS = système sécurité-apparenté. Probabilité de l'événement et le nombre total de couches de la protection indépendantes sont définis par rapport à la candidature spécifique.

FIG. 1.7 – Matrice de risque [IEC61508, 1998]

3. Les méthodes quantitatives [ISA-TR84.00.02-2002, 2002; Bhimavarapu *et al.*, 1997; of European Communities, 1992; Siu, 1996; Stavrianidis, 1995] : il s'agit des méthodes qui permettent de calculer la probabilité de défaillance moyenne des SIS à partir des paramètres de fiabilité de leurs composants.

Les méthodes les plus répandues sont :

- Les équations simplifiées [ISA, 2002a] qui donnent le PFD_{avg} du SIS en fonction de l'architecture des composants (1oo1, 1oo2, etc.) et des paramètres de fiabilité utilisés (taux de défaillance des composants λ , taux de défaillance dangereuses détectées DC et le facteur β qui caractérise les défaillances de cause commune).

Nous avons jugé bon de rappeler certains principes et hypothèses de base qui justifient l'utilisation des équations données par la norme :

- L'évaluation probabiliste des boucles de sécurité s'applique à des composants ayant des défaillances aléatoires et modélisées par une distribution exponentielle. Les taux de défaillance sont supposés être constants et indépendants du temps.
- Les pannes sont classées en quatre catégories. Nous distinguons les défaillances sûres et les défaillances dangereuses, chacune de ces catégories étant divisée en défaillances détectées et non détectées.

- Il faut tenir compte du taux de couverture de diagnostic qui exprime la proportion des défaillances dangereuses détectées.
- La probabilité moyenne de défaillance dangereuse (PFD_{avg}) à la sollicitation est due :
 - A l'apparition de défaillances dangereuses détectées lors de tests en ligne.
 - A la présence de défaillances dangereuses non détectées pendant la phase de fonctionnement et nécessitant une intervention. Pendant la phase de réparation, la fonction de sécurité n'est pas assurée.
- Les défaillances dangereuses totales sont égales aux défaillances dangereuses détectées et non détectées.
- Le temps de détection des tests en ligne est considéré comme nul.
- On introduit également des défaillances de causes communes pour les architectures redondantes.
- Les arbres de défaillance [ISA-TR84.00.02-2002, 2002].
- Les approches markoviennes [ISA, 2002b].

1.3.2.5 Quelques commentaires

Malgré tout, le déploiement de l'IEC 61508 [IEC61508, 1998] sectorielles se heurte à des difficultés diverses :

- Comme l'ont souligné plusieurs chercheurs dans le domaine de la sûreté de fonctionnement [Peguet and Boyer, 2006; Innal *et al.*, 2005; Innal *et al.*, 2006], il est nécessaire d'utiliser des méthodes de sûreté de fonctionnement classiques telles que les approches markoviennes ou les arbres de défaillance pour évaluer la fiabilité des SIS (le PFD_{avg} et le SIL), plutôt qu'utiliser les équations simplifiées données dans la partie 6 de la norme IEC 61508 [IEC61508, 1998].

En effet, la communauté des fiabilistes s'est rendue compte que certaines équations citées dans la partie 6 de la norme IEC 61508 [IEC61508, 1998] ne sont valables que sous plusieurs hypothèses qui ne sont pas citées dans la norme. En outre, ces formules ne sont valables que pour certains types d'architecture k parmi n .

- La disparité des institutions réglementant la sécurité des installations, les intérêts contradictoires des industriels et le manque de formation concernant cette norme sont autant d'obstacles à une diffusion large.
- L'approche proposée est inhabituelle pour les utilisateurs, généralement confrontés à des normes déterministes définissant des architectures et des prescriptions techniques très strictes. L'IEC 61508 leur laisse au contraire beaucoup d'initiative.
- L'IEC 61508 ne couvre pas les précautions qui peuvent se révéler nécessaires pour empêcher des personnes sans autorisation d'endommager et/ou d'affecter la sécurité fonctionnelle réa-

lisée par les systèmes E/E/PE concernés par la sécurité, notamment les intrusions dans les réseaux.

- Les définitions de quelques termes clés qui restent très ambigus.

On peut citer par exemple le PFD_{avg} qu'on peut assimiler à une indisponibilité moyenne du SIS sur une période donnée.

- Rappelons également que les données d'entrées spécifiées dans la norme sont bien souvent difficiles à obtenir et sont souvent des approximations (taux de couverture de diagnostic (DC), modes communs de défaillances, etc.).

1.3.3 Norme IEC 61511

L'IEC 61508 [IEC61508, 1998] est une norme assez complexe à appréhender parce qu'elle est très générale. C'est pour cela que ses concepteurs ont développé des normes sectorielles s'appliquant à des secteurs bien précis (cf. figure 1.8). Ces principales normes sont :

- L'IEC 61511 [IEC61511, 2000] : Cette norme concerne les process industriels. Dans les sections suivantes, nous détaillerons les principes de cette norme.
- L'IEC 62061 [IEC62061, 2005] : la norme IEC 62061 [IEC62061, 2005] est une norme spécifique au secteur des machines dans le cadre de l'IEC 61508 [IEC61508, 1998]. Elle est destinée à faciliter la spécification du fonctionnement des systèmes de commande électriques relatifs à la sécurité par rapport aux dangers significatifs des machines. Elle donne un cadre spécifique au secteur des machines pour la sécurité fonctionnelle d'un SRECS (système de commandes électriques relatif à la sécurité de machines). Elle couvre uniquement les aspects du cycle de vie de sécurité relatifs à l'allocation des exigences de sécurité jusqu'à la validation de la sécurité.

On peut conclure que cette norme est destinée à être utilisée par les concepteurs de machines, les fabricants et les intégrateurs de systèmes de commande et autres, impliqués dans la spécification, la conception et la validation d'un SRECS.

- L'IEC 61513 [IEC61513, 2001] : la norme IEC 61513 [IEC61513, 2001] établit les prescriptions relatives aux systèmes et équipements d'instrumentation et de contrôle commande (systèmes d'IC) utilisés pour accomplir les fonctions importantes pour la sûreté des centrales nucléaires. Notons que les systèmes d'IC importants pour la sûreté peuvent être réalisés à l'aide de composants traditionnels câblés, de composants informatiques ou d'une combinaison des deux.

Cette norme souligne en particulier les relations entre :

- les objectifs de sûreté de la centrale nucléaire et les exigences relatives à l'architecture globale des systèmes d'IC importants pour la sûreté ;
- l'architecture globale des systèmes d'IC et les exigences applicables à chacun des systèmes d'IC importants pour la sûreté.

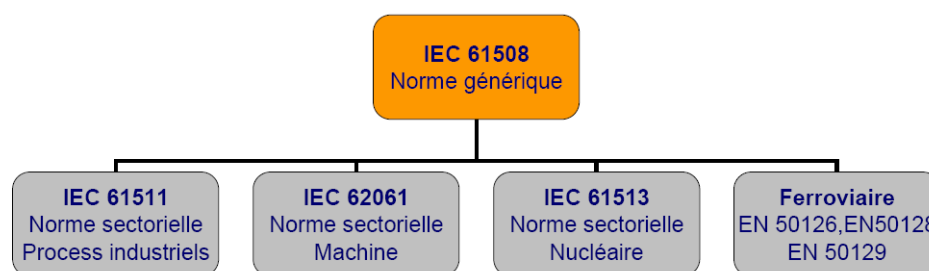


FIG. 1.8 – L’IEC 61508 et ses normes sectorielles

- L’EN 50126/8/9 [EN50126, 1999; EN50128, 2001; EN50129, 1998] : ces trois normes concernent les applications ferroviaires. La norme EN 50126 [EN50126, 1999] est destinée à la spécification et la démonstration de la fiabilité, de la disponibilité, de la maintenabilité et de la sécurité. La norme EN 50128 [EN50128, 2001] concerne les logiciels pour systèmes de contrôle/commande et de protection ferroviaire. La norme EN 50129 [EN50129, 1998] est destinée aux systèmes électroniques de sécurité pour la signalisation.

Dans ces normes sectorielles, une distinction est faite entre l’application du système E/E/PE de sécurité (qui dépend du secteur) et les spécifications détaillées de conception (qui sont indépendantes du secteur). Les spécifications indépendantes du secteur font référence à des parties et paragraphes de l’IEC 61508 et évitent les répétitions (cf. fig. 1.9 et fig. 1.10). Les utilisateurs qui veulent mettre en œuvre une norme sectorielle ont donc besoin de l’IEC 61508.

1.3.3.1 Domaine d’application

La norme IEC 61511 [IEC61511, 2000] concerne les SIS qui sont basés sur l’utilisation d’une technologie E/E/PE. Elle permet de définir des exigences relatives aux spécifications, à la conception, à l’installation, à l’exploitation et à l’entretien d’un SIS, de telle manière qu’il puisse être mis en œuvre en toute confiance, et ainsi établir et/ou maintenir les processus dans un état de sécurité convenable. Dans le cas où d’autres technologies sont utilisées pour les unités logiques, il convient aussi d’appliquer les principes fondamentaux de cette norme. Cette norme concerne également les capteurs et les éléments terminaux des SIS, quelle que soit la technologie utilisée. Cette norme est spécifique à la production industrielle par processus dans le cadre de l’IEC 61508 [IEC61508, 1998]. Nous pouvons ainsi conclure que l’IEC 61511 est destinée aux intégrateurs et aux utilisateurs, alors que l’IEC 61508 [IEC61508, 1998] est une norme générique difficile à mettre en œuvre et donc destinée surtout aux fabricants et fournisseurs de systèmes E/E/PE.

1.3.3.2 Structure générale de la norme

La norme IEC 61511 [IEC61511, 2000] comprend les parties suivantes (cf. figure 1.11) :

- Partie 1 : Cadre, définitions, exigences pour le système, le matériel et le logiciel.
- Partie 2 : Lignes directrices pour l’application de l’IEC 61511-1.
- Partie 3 : Guide pour la détermination des niveaux d’intégrité de sécurité requis.

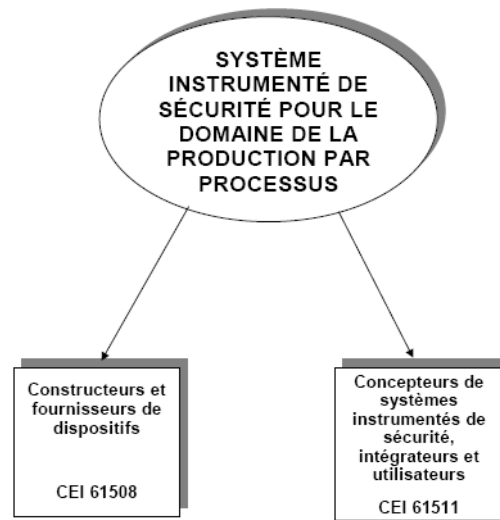


FIG. 1.9 – Utilisateurs de l'IEC 61511 et l'IEC 61508

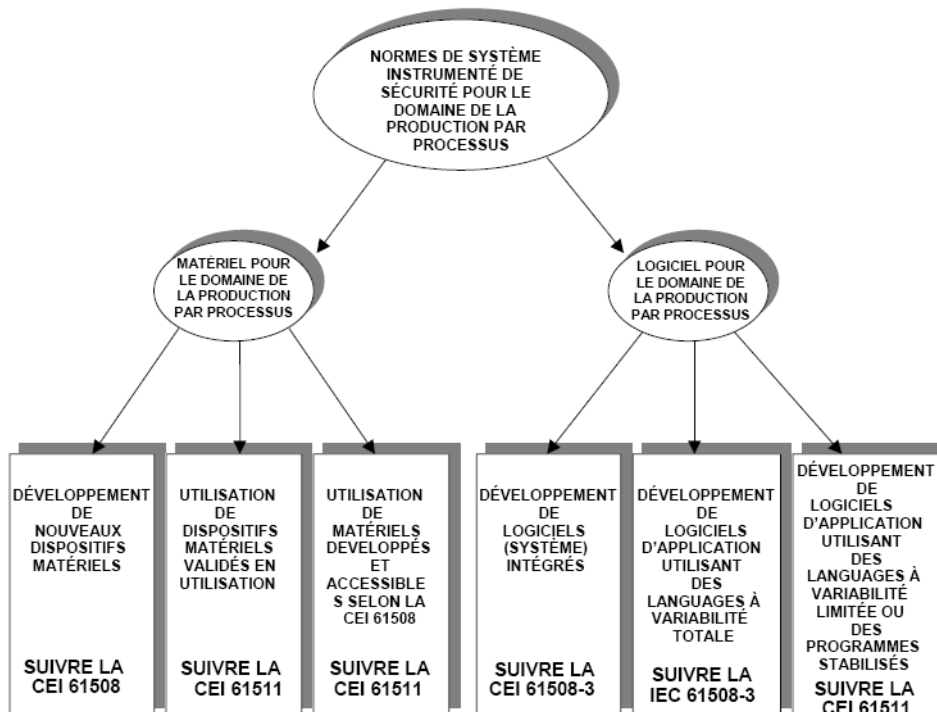


FIG. 1.10 – Relations entre l'IEC 61511 et l'IEC 61508

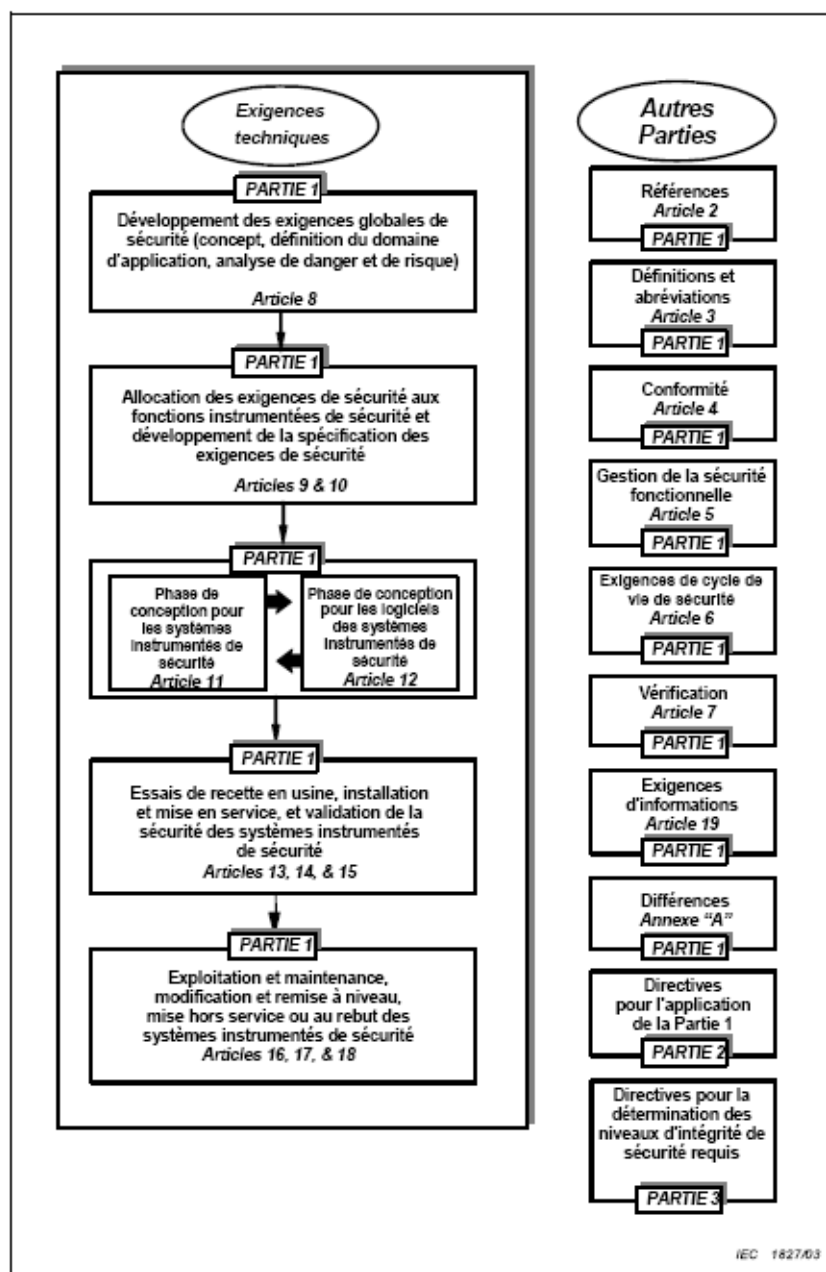


FIG. 1.11 – Structure générale de la norme IEC 61511 [IEC61511, 2000]

1.3.3.3 Objectifs de la norme

Cette norme présente une approche relative aux activités liées au cycle de vie de sécurité, pour satisfaire à ces normes minimales. Cette approche a été adoptée afin de développer une politique technique rationnelle et cohérente. Dans la plupart des cas, la meilleure sécurité est obtenue par une conception de processus de sécurité intrinsèques, chaque fois que cela est possible, combinée, au besoin, avec d'autres systèmes de protection, fondés sur différentes technologies (chimique, mécanique, hydraulique, pneumatique, électrique, électronique, électronique programmable) et qui couvrent tous les risques résiduels identifiés. Pour faciliter cette approche, cette norme nécessite :

- De réaliser une évaluation des dangers et des risques pour identifier les exigences globales de sécurité ;
- D'effectuer une allocation des exigences de sécurité au(x) système(s) instrumenté(s) de sécurité ;
- D'être inscrite dans un cadre applicable à toutes les méthodes instrumentées qui permettent d'obtenir la sécurité fonctionnelle ;
- De détailler l'utilisation de certaines activités, telles que la gestion de la sécurité, qui peuvent être applicables à toute méthode permettant d'obtenir la sécurité fonctionnelle.

1.3.3.4 Quelques commentaires

En conclusion, nous pouvons dire que cette norme :

- Prend en compte toutes les phases du cycle de vie de sécurité, depuis le concept initial, en passant par la conception, la mise en œuvre, l'exploitation et la maintenance, jusqu'au déclassement ;
- S'harmonise avec les normes spécifiques de processus industriels existantes.

1.4 Problématique de l'évaluation de la sûreté de fonctionnement des SIS

Les SIS sont des systèmes qui ont pour objectif de mettre le procédé en position de repli de sécurité (c'est-à-dire un état stable ne présentant pas de risque pour les personnes, l'environnement ou les biens) lorsque celui-ci évolue vers une voie comportant un risque réel (explosion, feu, etc.). Les normes IEC 61508 [IEC61508, 1998] et IEC 61511 [IEC61511, 2000] définissent les niveaux d'intégrité de sécurité que doit atteindre le SIS en fonction de leur $PF D_{avg}$ ou leur PFH . Dans ce mémoire, nous nous plaçons dans le contexte des SIS faiblement sollicités.

Les méthodes usuelles de calcul du $PF D_{avg}$ des SIS sont des méthodes probabilistes [Stavrianidis and Bhimavarapu, 1998; IEC61508, 1998; Beckman, 2000; IEC61511, 2000; Summers, 2002; Goble and Cheddie, 2006; Peguet and Boyer, 2006; Innal *et al.*, 2006]. Ces méthodes issues des études traditionnelles de sûreté de fonctionnement où les données de fiabilité relatives aux composants (taux de défaillance, taux de réparation, etc.) peuvent être connues avec précision et validées par le retour d'expérience. Nous pouvons cependant nous interroger sur leur adaptation à des systèmes hautement fiables pour lesquels les défaillances sont très rares comme les SIS rarement sollicités. Dans ce cas, le retour d'expérience est insuffisant pour valider avec précision les taux de défaillance. En outre, en raison de l'évolution constante de l'environnement et de la complexité des installations industrielles, les conditions d'utilisation des composants des

SIS utilisés dans les installations peuvent changer pour un même composant. L'idéal est de disposer d'une quantité d'information suffisante concernant les défaillances des composants pour pouvoir estimer avec précision leur taux de défaillance. Or, dans les installations industrielles, les SIS sont des systèmes rarement sollicités. De ce fait, les composants des SIS ne sont pas suffisamment utilisés pour fournir des données de défaillances suffisantes [Wang *et al.*, 2004; Villemeur, 1998].

Des bases de données de fiabilité génériques sont souvent utilisées pour fournir les taux de défaillance des composants des systèmes relatifs à la sécurité [OREDA, 2002; Exida, 2005; Hauge *et al.*, 2006; Goble and Cheddie, 2006]. Des valeurs ponctuelles des taux de défaillance sont généralement utilisés pour estimer les probabilités de défaillance des composants du SIS. Cependant, la collecte de données de fiabilité et l'utilisation de ces données pour estimer les taux de défaillance d'autres composants du même type introduit des incertitudes. Selon [Kletz, 1999], les données de fiabilité recueillies pour un composant peuvent changer d'un facteur de 3 ou 4, et un facteur de 10 n'est pas à exclure. Dans [MIL217B, 1974; Villemeur, 1998], les auteurs proposent d'utiliser des coefficients d'influence pour tenir compte des conditions réelles d'utilisation des valeurs génériques des taux de défaillance proposées par les bases de données de fiabilité. Wang *et al.* [Wang *et al.*, 2004] ont discuté de l'impact de l'incertitude des taux de défaillance des composants du SIS sur la détermination du SIL des Fonctions Instrumentées de Sécurité que doit réaliser le SIS. Mais, ils ont juste souligné que d'autres travaux doivent être menés pour étudier et justifier l'incertitude du SIL dans ce type de problème.

Certaines bases de données de fiabilité [IEEE, 1984; CCPS, 1991; CCPS, 2002] fournissent les bornes inférieures et supérieures, les valeurs moyennes médianes et les facteurs d'erreurs des taux de défaillance des composants. Des méthodes autres que les approches probabilistes classiques (ensembles flous, théorie de l'évidence, etc.) peuvent utiliser avantageusement ces valeurs pour prendre en compte l'imprécision (qu'on appelle aussi incertitude de type épistémique) liée à ces taux de défaillance et estimer les probabilités de défaillance des composants ainsi que les PFD_{avg} des SIS.

1.5 Conclusion

Ce chapitre a souligné la problématique de la gestion de risque dans les installations industrielles. Nous avons d'abord repris les notions clés comme la sécurité, le risque, et le danger et nous avons cherché à faire le lien entre ces différents termes. Ensuite, nous avons introduit la notion de sécurité fonctionnelle. Enfin, nous avons évoqué les difficultés liées à l'évaluation de la sûreté de fonctionnement de ces systèmes sécuritaires à partir de paramètres de fiabilité imparfaits issus d'un faible retour d'expérience ou en conditions d'exploitations imprécises.

2

Evaluation de la sûreté de fonctionnement des systèmes en présence d'informations imparfaites

Ce chapitre présente un état de l'art des méthodes de sûreté de fonctionnement qui prennent en compte les données imprécises ou incertaines de fiabilité des composants pour l'évaluation de la fiabilité des systèmes. Nous nous concentrons principalement sur les méthodes de fiabilité PROBIST floue où le fonctionnement du système est caractérisé par des mesures de probabilités floues et où il ne peut y avoir que deux états de fonctionnement du système.

Nous nous sommes fixé comme objectif de permettre, au travers de la définition des principales théories de représentation des connaissances imparfaites (théorie des probabilités, théorie des ensembles flous, théorie des possibilités et théorie de l'évidence), d'analyser leur utilisation pour l'évaluation de la sûreté de fonctionnement des systèmes en présence d'informations imparfaites (imprécises ou incertaines).

2.1 Introduction

Dans les études de sûreté de fonctionnement, les connaissances dont nous disposons concernant les données de fiabilité des composants sont généralement imparfaites [Villemeur, 1998; Gouriveau, 2003; Wang *et al.*, 2004; Sallak *et al.*, 2007a]. La validité des résultats des études dépend de la prise en compte totale ou partielle de l'imperfection des connaissances utilisées. Cela exige des méthodes qui permettent la modélisation et la manipulation de ces imperfections.

Dans la littérature, nous distinguons principalement deux sortes d'imperfection des connaissances [Kruse *et al.*, 1991; Dubois and Prade, 1994; Grabisch *et al.*, 1994; Bouchon-Meunier, 1995] : l'incertitude et l'imprécision. Les connaissances sont incertaines quand nous avons un doute sur leur validité. Si nous éprouvons une difficulté à les exprimer clairement, elles sont alors imprécises. Bouchon-Meunier considère que les imperfections dans les connaissances peuvent être divisées en trois types principaux [Bouchon-Meunier, 1995] :

- Les incertitudes qui représentent un doute sur la validité d'une connaissance.
- Les imprécisions qui correspondent à une difficulté dans l'énoncé ou dans l'obtention de la connaissance. Ces imprécisions sont aussi appelées incertitudes de type épistémique [Helton and Oberkampf, 2004; Baea *et al.*, 2004].
- Les incomplétudes qui sont des absences de connaissances ou des connaissances partielles sur certaines caractéristiques du système.

En anglais, nous employons souvent le terme "Uncertainty" pour désigner les connaissances imparfaites en général, alors que le terme "Imprecision" utilisé pour désigner les connaissances imprécises est rarement cité. Ces types d'imperfections sont souvent intimement mêlées, mais n'ont cependant pas présenté la même importance dans les préoccupations scientifiques.

Dans ce chapitre, nous allons étudier la prise en compte de l'imperfection des données de fiabilité des composants pour l'évaluation de la sûreté de fonctionnement des systèmes. Dans le domaine de la sûreté de fonctionnement, la fiabilité représente l'aptitude d'une entité à accomplir une fonction requise E , dans des conditions données, pendant une durée donnée [Villemeur, 1998]. Elle est généralement mesurée par la probabilité qu'une entité accomplisse la fonction requise E , pendant un intervalle de temps $[0, t]$. La théorie de fiabilité repose sur deux hypothèses fondamentales :

- Hypothèse probabiliste : le fonctionnement du système est complètement caractérisé par des mesures de probabilité.
- Hypothèse d'états binaires : le système ne peut avoir que deux états de fonctionnement ; l'état de défaillance et l'état de fonctionnement normal.

Ce formalisme correspond à la fiabilité PROBIST. Cai *et al.* [Cai *et al.*, 1991; Cai, 1993] ont suggéré trois autres types de fiabilité nommées respectivement, PROFUST (hypothèse probabiliste et hypothèse d'états flous), POSBIST (hypothèse possibiliste et hypothèse d'états binaires) et POSFUST (hypothèse possibiliste et hypothèse d'états flous).

Dans nos travaux, nous nous intéressons aux systèmes qui sont caractérisés par des mesures de probabilités imprécises pour modéliser les taux de défaillance imprécis des composants. En outre, et par souci de simplification, ces systèmes ne doivent avoir que deux états de fonctionnement

(défaillance et fonctionnement normal). C'est pourquoi, nous proposons dans ce chapitre, un état de l'art des méthodes de sûreté de fonctionnement utilisées dans le cadre de la fiabilité PROBIST floue.

2.2 Représentation des connaissances imparfaites

Pour étudier la fiabilité des systèmes, nous disposons d'informations qui sont généralement imparfaites. Plusieurs théories ont été développées pour permettre la modélisation et la manipulation de ces imperfections (incertitudes et imprécisions).

En ce qui concerne l'incertain, il a d'abord été abordé par la notion de probabilité dès le *XVII^{me}* siècle par Pascal et Fermat [Kolmogorov, 1960]. La théorie des probabilités fournit une structure mathématique pour l'étude des phénomènes qui présentent des incertitudes aléatoires. Le problème de l'imprécision a été traité par le calcul d'erreurs, restreint aux imprécisions de caractère numérique. En 1965, Lotfi Zadeh [Zadeh, 1965], professeur à l'université de Berkley en Californie, a introduit la notion de sous-ensemble flou (en anglais "Fuzzy set") dans une généralisation de la théorie classique des ensembles, admettant des situations intermédiaires entre le tout et le rien. Les développements de cette notion fournissent des moyens de représenter les connaissances imprécises et, ils établissent une interface entre les données décrites symboliquement et numériquement. Lotfi Zadeh a ensuite introduit la théorie des possibilités, à partir de 1978 [Zadeh, 1978], qui a été développée par Dubois et Prade [Dubois and Prade, 1988]. Elle permet de traiter les incertitudes sur les connaissances. L'association de la théorie des possibilités à la théorie des sous-ensembles flous permet le traitement des connaissances à la fois imprécises et incertaines. La théorie des fonctions de croyances permet aussi de traiter ces deux types d'imperfections [Dempster, 1967; Smets and Kennes, 1994]. Elle est basée sur la modélisation et la quantification de la crédibilité attribuée à des faits. Elle définit le degré avec lequel un événement est crédible ou plausible.

2.2.1 Théorie des probabilités

La théorie des probabilités bénéficie de quatre siècles de travaux et repose donc sur des fondements mathématiques et une expérience solides. Elle constitue le plus ancien formalisme permettant de traiter les incertitudes dans les connaissances imparfaites.

2.2.1.1 Définitions et propriétés

Définition

La probabilité d'occurrence d'un événement A d'un ensemble de référence Ω est décrite par la mesure de probabilité :

$$P : S(\Omega) \mapsto [0, 1] \quad (2.1)$$

Où $S(\Omega)$ est l'ensemble des parties de Ω ; et qui satisfait les axiomes suivants :

$$\forall A \in S(\Omega) \quad P(\emptyset) = 0 \leq P(A) \leq P(\Omega) = 1 \quad (2.2)$$

$$\forall A, B \in S(\Omega), \text{ Si } A \cap B = \emptyset \quad P(A \cup B) = P(A) + P(B) \quad (2.3)$$

L'axiome d'additivité (2.3) établit que, si deux événements sont incompatibles ($A \cap B = \emptyset$), alors la probabilité qu'au moins l'un d'entre eux ait lieu, est la somme de leurs probabilités individuelles.

On peut déduire aussi des deux axiomes (2.2) et (2.3) les deux propriétés suivantes :

Propriété 1

Si la probabilité d'un événement A est connue alors la probabilité de l'événement contraire $\bar{A}$ est exactement déterminée :

$$\forall A \in S(\Omega) \quad P(A) + P(\bar{A}) = 1 \quad (2.4)$$

Propriété 2

On peut caractériser la probabilité d'événements ω_i sur Ω en utilisant la distribution de probabilité :

$$p : \Omega \mapsto [0, 1] \quad (2.5)$$

$$\text{Où : } \sum_{\omega_i \in \Omega} p(\omega_i) = 1.$$

La distribution de probabilité p est définie à partir de la mesure de probabilité P par $p(\omega_i) = P(\{\omega_i\})$.

On obtient alors :

$$\forall A \in S(\Omega) \quad P(A) = \sum_{\omega_i \in A} p(\omega_i) \quad (2.6)$$

La mesure de probabilité P a une interprétation soit fréquentiste soit subjectiviste. Dans l'interprétation fréquentiste, utilisée dans le cas des informations statistiques, $P(A)$ représente la limite de la fréquence d'occurrence de l'événement A dans une séquence infinie d'expériences indépendantes [Ruegg, 1985]. Dans l'interprétation subjectiviste, $P(A)$ représente la croyance qu'un expert accorde à l'occurrence de A [Shafer, 1976].

2.2.1.2 Cadre bayésien

Dans le cadre bayésien, la règle de Bayes permet de combiner plusieurs distributions de probabilités. Elle permet d'estimer les probabilités d'occurrences d'événements futurs (probabilités a posteriori) à partir des probabilités d'occurrences d'événements similaires passés (probabilités a priori).

Si $A_1, A_2, \dots, A_N$ forment un ensemble d'événements incompatibles et recouvrant l'ensemble de référence, la probabilité a posteriori d'un événement A_i connaissant l'information n_j peut être déterminée par :

$$P(A_i|n_j) = \frac{P(A_i).P(n_j|A_i)}{\sum_{k=1}^N P(A_k).P(n_j|A_k)} \quad (2.7)$$

où $P(A_i)$ est la probabilité a priori de l'événement A_i , $P(n_j|A_i)$ représente la probabilité d'observer l'information n_j lorsque l'événement A_i est réalisée et $P(A_i|n_j)$ est la probabilité a posteriori de l'événement A_i connaissant l'information n_j .

Les lois $P(A_i)$ et $P(n_j|A_i)$ sont en pratique rarement connues. Elles sont souvent estimées à partir des données. Les probabilités $P(A_i)$ sont déterminées par l'expérience ou par une analyse d'exemples. Les probabilités conditionnelles $P(n_j|A_i)$ sont estimées par des lois statistiques. Ces dernières peuvent être paramétriques, et dans ce cas une forme est choisie pour $P(n_j|A_i)$ et ses paramètres sont estimés (par maximum de vraisemblance par exemple), ou non paramétriques (par exemple les fenêtres de Parzen [Parzen, 1962; Chappelle, 2005] ou la méthode des k plus proches voisins [Quinquis and Radoi, 1998]).

2.2.1.3 Limitations de la théorie des probabilités

Lorsque certains indices observés permettent d'augmenter la confiance en une hypothèse, l'axiome d'additivité (cf. 2.4) nécessite de diminuer d'autant la confiance correspondant à l'hypothèse contraire. Néanmoins, dans certains cas, des informations peuvent très bien favoriser une hypothèse sans pour autant discréditer l'hypothèse contraire. La théorie des probabilités semble donc peu adaptée à des situations où la connaissance d'un événement comme la connaissance de son contraire sont très limitées [Zouhal, 1997]. Signalons que plusieurs exemples sont disponibles dans la littérature pour mettre en évidence le fait que la théorie des probabilités est trop rigide pour exprimer le cas de l'ignorance totale [Zimmermann, 1983; Zouhal, 1997]. Elle modélise ce cas par un ensemble d'événements ω_i mutuellement disjoints et équiprobables :

$$\forall \omega_i \in \Omega, \quad p(\omega_i) = \frac{1}{|\Omega|} \quad (2.8)$$

Où : $|\Omega|$ est le nombre d'événements de l'ensemble de référence Ω .

2.2.1.4 Conclusion

La théorie des probabilités constitue un outil efficace pour le traitement des incertitudes aléatoires et les cas où nous disposons d'une bonne connaissance des événements et de leurs événements contraires. Elle ne peut cependant pas traiter les imprécisions qui sont une autre forme d'imperfection des connaissances [Gouriveau, 2003; Baudrit, 2005]. Nous allons introduire la notion de sous-ensemble flous dans la section suivante qui permet de traiter l'aspect imprécis et vague des connaissances imparfaites.

2.2.2 Théorie des sous-ensembles flous

Les sous-ensembles flous ont été introduits en 1965 par Zadeh [Zadeh, 1965] comme une généralisation des ensembles classiques aux ensembles dont les frontières sont mal définies. Ils permettent ainsi une transition graduelle de l'état de non appartenance à l'état d'appartenance totale.

2.2.2.1 Définitions et concepts de base

Définition 1

Un sous-ensemble usuel E d'un ensemble de référence Ω peut être défini par sa fonction caractéristique :

$$\mu_E : \Omega \mapsto \{0, 1\} \quad (2.9)$$

Un élément u de Ω est un élément de l'ensemble E si et seulement si $\mu_E(u) = 1$. Un élément v n'appartient pas à E si et seulement si $\mu_E(v) = 0$.

Définition 2

Un sous-ensemble flou A est lui, caractérisé par sa fonction d'appartenance μ_A ; Zadeh [Zadeh, 1965] en donne la définition suivante :

Un sous-ensemble flou A sur un référentiel Ω est caractérisé par une fonction d'appartenance μ_A qui associe à chaque élément x de Ω un nombre réel dans l'intervalle $[0, 1]$:

$$\mu_A : \Omega \mapsto [0, 1] \quad (2.10)$$

La valeur $\mu_A(x)$ représente le degré d'appartenance de x à A . Le degré d'appartenance d'un élément u de Ω à A n'appartient plus à la paire $\{0, 1\}$ comme pour un ensemble classique, mais à l'intervalle réel $[0, 1]$. Si $\mu_A(x) = 0$, x n'appartient pas du tout à A , si $\mu_A(x) = 1$, il lui appartient totalement. Si $0 < \mu_A(x) < 1$ alors l'appartenance de x à Ω est plus ou moins complète.

Les éléments qui appartiennent complètement à l'ensemble flou $A : \{x \in \Omega / \mu_A(x) = 1\}$, constituent le noyau de A . Ceux pour lesquels le degré d'appartenance n'est pas nul constituent le support de A . Si le noyau de A est non vide, A est dit normalisé.

2.2.2.2 Opérations simples sur les sous-ensembles flous

La théorie des sous-ensembles flous propose plusieurs opérateurs ensembliste. Les principaux opérateurs et relations ensemblistes flous sont présentés ci-dessous.

Inclusion

On dit que A est inclus dans B , et on note $A \subseteq B$, si et seulement si :

$$\forall x \in \Omega \quad \mu_A(x) \leq \mu_B(x) \quad (2.11)$$

Egalité

On dit que A et B sont égaux, et on note $A = B$, si et seulement si :

$$\forall x \in \Omega \quad \mu_A(x) = \mu_B(x) \quad (2.12)$$

S'il existe au moins un x de Ω tel que l'égalité $\mu_A(x) = \mu_B(x)$ n'est pas satisfaite, on dit que A et B sont inégaux.

Complémentation

On dit que A et B sont complémentaires, et on note $A = \overline{B}$ ou $\overline{A} = B$, si et seulement si :

$$\forall x \in \Omega \quad \mu_B(x) = 1 - \mu_A(x) \quad (2.13)$$

Intersection

On définit l'intersection de A et B , et on note $A \cap B$, par le plus grand ensemble flou de Ω contenu à la fois dans A et B , c'est à dire :

$$\forall x \in \Omega \quad \mu_{A \cap B}(x) = \min(\mu_A(x), \mu_B(x)) \quad (2.14)$$

Union

On définit l'union ou la réunion de A et B , et on note $A \cup B$, par le plus petit ensemble flou de Ω qui contient à la fois A et B , c'est à dire :

$$\forall x \in \Omega \quad \mu_{A \cup B}(x) = \max(\mu_A(x), \mu_B(x)) \quad (2.15)$$

Toutes ces opérations sont des extensions des opérations ensemblistes usuelles avec lesquelles elles coïncident si les ensembles considérés sont des ensembles usuels. Les extensions de ces opérations ensemblistes aux ensembles flous ne sont pas uniques [Wu, 1998].

Norme triangulaire

Une norme triangulaire (ou t-norme) $\top$ est une fonction de $[0, 1] \times [0, 1] \mapsto [0, 1]$ qui vérifie pour tous x, y, z et t de $[0, 1]$:

- $\top(x, y) = \top(y, x)$ (commutativité)
- $\top(x, \top(y, z)) = \top(\top(x, y), z)$ (associativité)
- Si $x \leq z$ et $y \leq t$ alors $\top(x, y) \leq \top(z, t)$ (monotonie)
- $\top(x, 1) = x$ (élément neutre 1)

Co-norme triangulaire

Une co-norme triangulaire (ou t-conorme) $\perp$ est une fonction de $[0, 1] \times [0, 1] \mapsto [0, 1]$ qui vérifie pour tous x, y, z et t de $[0, 1]$:

- $\perp(x, y) = \perp(y, x)$ (commutativité)
- $\perp(x, \perp(y, z)) = \perp(\perp(x, y), z)$ (associativité)
- Si $x \leq z$ et $y \leq t$ alors $\perp(x, y) \leq \perp(z, t)$ (monotonie)
- $\perp(x, 0) = x$ (élément neutre 0)

Nous remarquons d'après les propriétés ci-dessus que les notions de t-normes et de t-conormes sont duales l'une de l'autre. Nous pouvons passer de l'une à l'autre par un opérateur de négation, par exemple n défini par $n(x) = 1 - x$.

Nom	Norme	Co-norme	Propriétés
Zadeh	$\min(x,y)$	$\max(x,y)$	de Morgan ; distributivité et idempotence ; ni tiers exclus, ni non contradiction.
Probabiliste	$x.y$	$x+y-x.y$	de Morgan ; non distributivité et non idempotence ; tiers exclu, non contradiction.
Lukasiewicz	$\max(x+y-1,0)$	$\min(x+y,1)$	tiers exclu ; non-contradiction.
Weber	x si $y=1$; y si $x=1$; 0 sinon	x si $y=0$; y si $x=0$; 1 sinon	de Morgan ; non distributivité et non idempotence ; tiers exclu, non contradiction.

TAB. 2.1 – Principales normes et co-normes triangulaires [Wu, 1998].

Nous pouvons vérifier que les fonctions $\min$ et $\max$ utilisées par Zadeh sont respectivement une norme et une co-norme triangulaire. Le couple de normes triangulaires ($\min, \max$) et la complémentation sont les fonctions les plus utilisées car elles sont les uniques fonctions qui munissent l'ensemble des sous-ensembles flous de Ω de toutes les propriétés habituelles, exceptés la loi de non-contradiction ($A \cap \bar{A} = \emptyset$) et le principe du tiers-exclus ($A \cup \bar{A} = \Omega$). En effet, les fonctions $\min$ et $\max$ sont les seules définitions possibles de l'intersection et de l'union qui présentent une structure de treillis vectoriel (ensemble ordonné où toute paire d'éléments a une borne supérieure et une borne inférieure) sur les parties floues de Ω . Plus généralement, quand l'appartenance est graduelle, le respect de la propriété d'idempotence ($A \cup A = A$ et $A \cap A = A$) entraîne la perte des lois de non-contradiction et du tiers exclus (et vice-versa).

Nous présentons dans le tableau 2.1, les principales normes et co-normes triangulaires ainsi que leurs propriétés. Elles sont ordonnées ainsi :

$$\forall x, y \in [0, 1] \quad \top_{Weber}(x, y) \leq \top(x, y) \leq \min(x, y) \quad (2.16)$$

$$\forall x, y \in [0, 1] \quad \max(x, y) \leq \perp(x, y) \leq \perp_{Weber}(x, y) \quad (2.17)$$

2.2.2.3 Nombres flous du type $L - R$

Les définitions des opérations floues décrites (unions et intersections de sous ensembles flous) peuvent être relativement difficiles à mettre en œuvre. Nous nous intéressons à des familles de formes particulières de fonctions d'appartenance qui conduisent à des calculs simples pour ces opérations.

Définition 1

Soit x une variable réelle continue de fonction d'appartenance $\mu(x) \in [0, 1]$, et qui satisfait aux conditions suivantes :

- $\mu(x)$ est continue par morceaux ;
- $\mu(x)$ est convexe ;

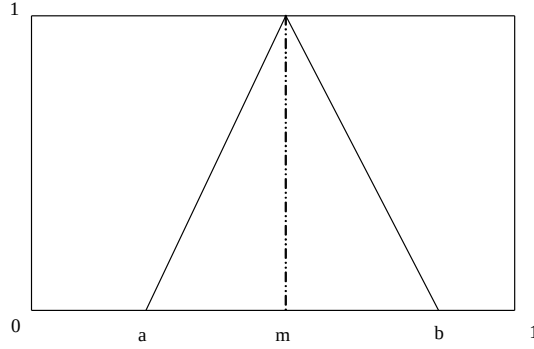


FIG. 2.1 – Nombre flou du type L-R

- $\mu(x)$ est normale (il existe au moins une valeur x_0 telle que $\mu(x_0) = 1$).

L'ensemble flou dont la fonction d'appartenance satisfait à ces conditions est appelé nombre flou.

Définition 2

Soient trois paramètres réels (m, a, b) , a et b étant strictement positifs, et deux fonctions, notées L et R , définies sur l'ensemble $\mathbb{R}^+$ des réels positifs, à valeurs dans $[0, 1]$, semi continues supérieurement, telles que :

$$L(0) = R(0) = 1, \quad L(1) = 0 \text{ ou } L(x) > 0 \quad \forall x \quad \text{avec} \quad \lim_{x \rightarrow \infty} L(x) = 0, \quad R(1) = 0 \quad (2.18)$$

$$\text{ou } R(x) > 0 \quad \forall x \text{ avec } \lim_{x \rightarrow \infty} R(x) = 0$$

Un nombre flou M est de type $L - R$ si sa fonction d'appartenance f_M est définie par :

$$f_M(x) = \begin{cases} R((x-m)/b) & \text{si } x > m \\ L((m-x)/a) & \text{si } x \leq m \end{cases} \quad (2.19)$$

Nous notons $M = (m, a, b)_{LR}$ un nombre flou de type $L - R$. m est sa valeur modale avec $f_M(m) = 1$, a est la largeur de son support à gauche de m , encore appelé étalement gauche, b celle de son support à droite de m , encore appelé étalement droit, sur l'axe des réels (cf. figure 2.1). L et R sont les deux fonctions qui déterminent sa fonction d'appartenance respectivement à gauche et à droite de m .

Nous pouvons, par exemple, utiliser les fonctions suivantes pour définir L et R :

- $g(x) = \max(0, 1 - x^2)$
- $g(x) = \max(0, 1 - x)^2$
- $g(x) = \exp(-x)$
- $g(x) = \max(0, 1 - x)$ qui donne une fonction d'appartenance linéaire par morceaux (cf. figure 2.1).

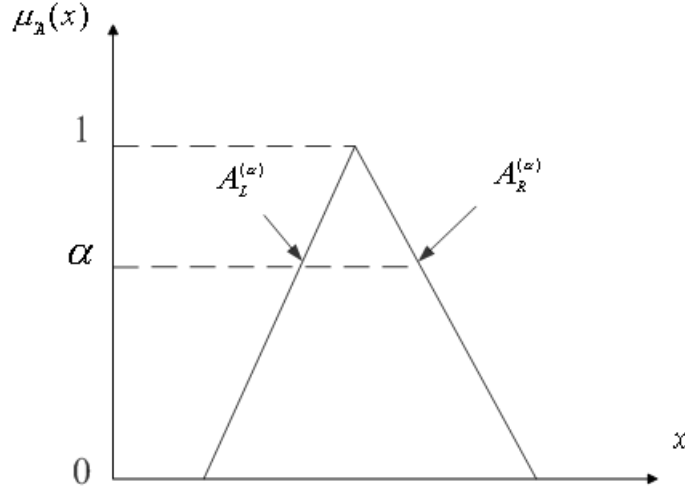


FIG. 2.2 – α -coupes d'un nombre flou

Propriétés

Étant donné deux nombres flous du même type $L - R$, $M = (m, a, b)_{LR}$ et $N = (n, c, d)_{LR}$, les opérations floues donnent les nombres flous suivants :

- $-M = (-m, a, b)_{LR}$, de type $R - L$,
- $M - N = (m - n, a + c, b + d)_{LR}$ si $L = R$, de type $L - L$,
- $M \times N$ n'est, par contre, généralement pas du type $L - R$, mais il est possible d'en donner une valeur approchée de type $L - R$ lorsque M et N ont un support inclus dans $\mathbb{R}^+$, que $m - a$ et $b - m$ sont petits devant m , $n - c$ et $d - n$ sont petits devant n :
 $M \times N = (mn, mc + na, md + nb)_{LR}$

2.2.2.4 Notion de α -coupes

Les opérations arithmétiques utilisées pour manipuler les nombres flous requièrent beaucoup de ressources. Kaufman et Gupta [Kaufman and Gupta, 1991] ont montré que ces efforts de calculs sont largement simplifiés par la décomposition des fonctions d'appartenance des nombres flous en α -coupes ($0 \leq \alpha \leq 1$). En effet, si nous considérons un nombre flou $\tilde{A}$ de fonction d'appartenance $\mu_{\tilde{A}}(x)$ (cf. Fig. 2.2), nous pouvons obtenir plusieurs intervalles emboîtés en utilisant la méthode des α -coupes. $A_L^{(\alpha)}$ et $A_R^{(\alpha)}$ représentent respectivement les limites droites et gauches de la fonction d'appartenance $\mu_{\tilde{A}}(x)$ à chaque coupe de niveau α .

Les opérations arithmétiques appliquées à deux nombres flous $\tilde{A}$ et $\tilde{B}$ donnent les expressions suivantes :

$$\tilde{C} = \tilde{A} + \tilde{B} \rightarrow [C_L^{(\alpha)}, C_R^{(\alpha)}] = [A_L^{(\alpha)} + B_L^{(\alpha)}, A_R^{(\alpha)} + B_R^{(\alpha)}] \quad (2.20)$$

$$\tilde{C} = \tilde{A} - \tilde{B} \rightarrow [C_L^{(\alpha)}, C_R^{(\alpha)}] = [A_L^{(\alpha)} - B_L^{(\alpha)}, A_R^{(\alpha)} - B_R^{(\alpha)}] \quad (2.21)$$

$$\tilde{C} = \tilde{A} \cdot \tilde{B} \rightarrow [C_L^{(\alpha)}, C_R^{(\alpha)}]$$

$$\tilde{C} = [\min(A_L^{(\alpha)} \cdot B_L^{(\alpha)}, A_R^{(\alpha)} \cdot B_L^{(\alpha)}, A_L^{(\alpha)} \cdot B_R^{(\alpha)}, A_R^{(\alpha)} \cdot B_R^{(\alpha)}), \max(A_L^{(\alpha)} \cdot B_L^{(\alpha)}, A_R^{(\alpha)} \cdot B_L^{(\alpha)}, A_L^{(\alpha)} \cdot B_R^{(\alpha)}, A_R^{(\alpha)} \cdot B_R^{(\alpha)})] \quad (2.22)$$

$$\tilde{C} = \tilde{A} \div \tilde{B} \rightarrow [C_L^{(\alpha)}, C_R^{(\alpha)}] = [A_L^{(\alpha)}, A_R^{(\alpha)}] \times [\frac{1}{B_R^{(\alpha)}}, \frac{1}{B_L^{(\alpha)}}] \quad (2.23)$$

2.2.3 Théorie des possibilités

La théorie des sous-ensembles flous permet de traiter les imprécisions et constitue une interface entre descriptions symboliques et descriptions numériques. Cependant, elle ne permet de traiter que les imprécisions (incertitudes de type épistémique) dans les connaissances. C'est pourquoi la théorie des possibilités a été introduite en 1978 par Zadeh [Zadeh, 1978], puis développée par Dubois et Prade [Prade and Dubois, 1985], dans le but d'obtenir un cadre permettant de traiter à la fois les concepts d'imprécisions et d'incertitudes. La théorie des possibilités fournit donc un moyen pour formaliser des incertitudes subjectives sur des événements, c'est-à-dire un moyen de dire dans quelle mesure la réalisation d'un événement est possible et dans quelle mesure nous sommes certain, sans toutefois avoir à sa disposition l'évaluation de la probabilité de cette réalisation [Sandri, 1991; Dubois *et al.*, 1993].

2.2.3.1 Mesure et distribution de possibilités

Une mesure de possibilité Π attribue à chaque sous-ensemble A de l'ensemble de référence Ω un réel dans $[0,1]$ évaluant à quel point l'événement A est possible [Prade and Dubois, 1985]. Ainsi, Π est une fonction de $S(\Omega)$ l'ensemble des parties de Ω à valeurs dans $[0,1]$, telle que :

$$\Pi(\emptyset) = 0 \quad (2.24)$$

$$\Pi(\Omega) = 1 \quad (2.25)$$

$$\forall A_i \in S(\Omega) \quad \Pi(\bigcup_i A_i) = \sup_i \Pi(A_i) \quad (2.26)$$

Dans le cas de deux sous-ensembles A et B , la dernière égalité se réduit à :

$$\Pi(A \cup B) = \max(\Pi(A), \Pi(B)) \quad (2.27)$$

Elle exprime que la réalisation de l'un des deux événements A ou B , pris indifféremment, est affectée du même coefficient de possibilité que la réalisation de l'événement le plus possible. Une mesure de possibilité n'est donc pas additive, mais suradditive.

Un événement est tout à fait possible si la mesure de sa possibilité est égale à 1, et impossible si elle est égale à 0. Une mesure de possibilité permet de déterminer le degré avec lequel l'union d'événements, dont on sait à quel point ils sont possibles, sera elle-même un événement possible.

La mesure de possibilité associée à l'intersection de parties de Ω est un réel majoré par la plus petite des mesures attribuées à chacune des parties.

Si enfin nous étudions un événement A et son contraire, l'un au moins est tout à fait possible et :

$$\max(\Pi(A), \Pi(\overline{A})) = 1 \quad (2.28)$$

$$\Pi(A) + \Pi(\overline{A}) \geq 1 \quad (2.29)$$

Une mesure de possibilité Π est totalement définie si toute partie de Ω se voit attribuer un réel dans $[0,1]$. Ainsi, si $|\Omega| = n$, il faut déterminer 2^n coefficients pour connaître complètement Π . Pour définir cette mesure plus simplement, il suffit d'indiquer les coefficients attribués aux singletons de Ω (un sous-ensemble quelconque de Ω pouvant être vu comme l'union de tels singletons), ramenant le problème à la détermination de n coefficients. Une telle définition repose sur la donnée d'une autre fonction, une distribution de possibilité π , qui attribue à tout singleton de Ω un réel dans $[0,1]$, et qui vérifie de plus (cas où Ω contient un événement certain) :

$$\sup_{x \in \Omega} \pi(x) = 1 \quad (2.30)$$

L'ensemble non flou $\{x \in \Omega / \pi(x) \geq \alpha\}$ est appelé α -coupe de la distribution de possibilité π . En particulier, la 0-coupe est le support de π , et la 1-coupe est le noyau de π .

2.2.3.2 Mesure de nécessité

Une mesure de possibilité Π donne une information sur l'occurrence d'un événement A relatif à l'ensemble de référence Ω , mais elle ne suffit pas pour décrire l'incertitude existante sur cet événement [Prade and Dubois, 1985]. Par exemple, si $\Pi(A) = 1$, il est tout à fait possible que A soit réalisé mais on peut avoir en même temps :

- $\Pi(\overline{A}) = 1$, qui exprime une indétermination complète sur la réalisation de A ;
- $\Pi(\overline{A}) = 0$, qui met en évidence le fait que A seul peut être réalisé.

Pour compléter l'information sur A , une mesure de nécessité N permet d'indiquer le degré avec lequel la réalisation de A est certaine. N est la grandeur duale de la mesure de possibilité. Cette mesure attribue à tout A un réel dans $[0,1]$. La nécessité d'un événement correspond à l'impossibilité de l'événement contraire :

$$\forall A \subset \Omega \quad N(A) = 1 - \Pi(\overline{A}) \quad (2.31)$$

Une mesure de nécessité vérifie les propriétés suivantes :

$$N(\emptyset) = 0 \quad (2.32)$$

$$N(\Omega) = 1 \quad (2.33)$$

$$\forall A_i \in S(\Omega) \quad N(\bigcap_i A_i) = \inf_i \Pi(A_i) \quad (2.34)$$

Dans le cas de deux sous-ensembles A et B , la dernière égalité se réduit à :

$$N(A \cap B) = \min(N(A), N(B)) \quad (2.35)$$

Les degrés de possibilité et de nécessité sont liés par les relations suivantes :

$$N(A) > 0 \Rightarrow \Pi(A) = 1 \quad (2.36)$$

$$\Pi(A) < 1 \Rightarrow N(A) = 0 \quad (2.37)$$

Ces deux relations indiquent que tout événement dont nous sommes au moins un peu certain est tout à fait possible, et que nous ne pouvons avoir la moindre certitude sur un événement qui n'est pas complètement possible.

2.2.4 Théorie des fonctions de croyance

La théorie des possibilités et la théorie des probabilités apparaissent comme cas particuliers d'une théorie plus vaste, la théorie des fonctions de croyance. Initiée par les travaux de Dempster sur les bornes inférieures et supérieures d'une famille de distribution de probabilités induites par une fonction multivaluée [Dempster, 1967], elle a été développée par Shafer [Shafer, 1976]. Cette théorie consiste en une quantification de la crédibilité attribuée à des connaissances imparfaites. Il existe plusieurs modèles distincts de la théorie de croyance : la théorie des probabilités imprécises ou des probabilités inférieures et supérieures [Walley, 1991], le modèle initial de Dempster développé par Kohlas *et al.* [Kohlas, 1997], le modèle des croyances transférables [Smets and Kennes, 1994], etc.

2.2.4.1 Fonction de masse (ou de croyance)

Soit Ω un ensemble de N hypothèses H_i exclusives et exhaustives, appelé cadre de discernement. 2^Ω désigne l'ensemble des 2^N sous-ensembles A_i de Ω . Une fonction de masse élémentaire m (qu'on appelle en anglais basic probability assignment ou mass function) est définie par :

$$m : 2^\Omega \rightarrow [0, 1] \quad (2.38)$$

$$m(\emptyset) = 0 \quad (2.39)$$

$$\sum_{2^\Omega} m(A_i) = 1 \quad (2.40)$$

Les éléments focaux sont les sous-ensembles de Ω ayant une masse non nulle. Lorsque ces éléments focaux se réduisent aux seuls singletons de Ω , les masses coïncident avec les probabilités. L'apport de la théorie des croyances est donc de permettre l'évaluation conjointe d'ensembles quelconques de ces singletons H_i . Dès lors, les événements considérés ne sont plus nécessairement exclusifs. On remarque aussi qu'une masse $m(A_i)$ est représentative de la vraisemblance attribuable à l'un des éléments du sous-ensemble A_i , sans aucun discernement possible entre les différents éléments de A_i . En particulier, $m(\Omega)$ désigne le degré d'incertitude ou d'ignorance totale.

2.2.4.2 Fonctions de crédibilité et de plausibilité

Une fonction de crédibilité Cr peut être définie sur les mêmes ensembles par :

$$Cr(\emptyset) = 0 \quad (2.41)$$

$$Cr(\Omega) = 1 \quad (2.42)$$

$$Cr(A_j) = \sum_{A_i \subseteq A_j} m(A_i) \quad (2.43)$$

La crédibilité de A_j est la somme des éléments focaux A_i qui entraînent A_j . Elle mesure donc à quel point les informations données par une source soutiennent A_j . Les fonctions de masse élémentaire et de crédibilité sont définies et utilisables de façon indépendante. Il existe cependant une bijection entre l'ensemble des fonctions de masse élémentaire et l'ensemble des fonctions de crédibilité, qui associe à chaque jeu de masses sur 2^Ω un jeu de crédibilités sur le même ensemble.

La fonction de plausibilité peut également être introduite à partir des fonctions de masse élémentaire :

$$Pl(A_j) = \sum_{A_i \cap A_j \neq \emptyset} m(A_i) \quad (2.44)$$

Cette fonction mesure à quel point les informations données par une source ne contredisent pas A_j . En fait, il a été démontré [Shafer, 1976] que la connaissance d'une des trois fonctions (m, Cr, Pl) sur 2^Ω était suffisante pour en déduire les deux autres. De façon intuitive, la crédibilité peut être interprétée comme une mesure de vraisemblance minimale d'un événement, et la plausibilité comme une mesure de vraisemblance maximale (croyance en l'événement ajoutée à l'incertain sur sa réalisation) :

$$\forall A \subset \Omega : Cr(A) \leq P(A) \leq Pl(A) \quad (2.45)$$

2.2.5 Conclusion

La théorie des probabilités a été introduite pour ne représenter que l'incertitude contenue dans les connaissances imparfaites. Il se peut cependant que des informations comprennent simultanément deux types d'imperfections : les incertitudes et les imprécisions, c'est par exemple le cas avec la phrase : *Il est probable que le taux de défaillance du capteur soit proche de $10^{-2}/ans$* . L'imprécision concerne le contenu de l'information, tandis que l'incertitude concerne sa véracité [Prade and Dubois, 1985]. La théorie des possibilités associée aux sous-ensembles flous et la théorie des croyances ont été introduites pour la représentation et le traitement de l'information à la fois imprécise et incertaine.

Deux conformations des éléments focaux permettent de retrouver les théories des probabilités et des possibilités comme cas particuliers de la théorie des croyances :

- Lorsque les éléments focaux A_i sont emboîtés, *i.e.* tels que $A_1 \subset A_2 \subset \dots \subset A_n$, la fonction de plausibilité a les propriétés d'une mesure de possibilité et la fonction de crédibilité a les propriétés d'une mesure de nécessité ;
- Lorsque les éléments focaux sont des singletons, toute partie A de Ω est telle que $Cr(A) = Pl(A)$ et cette valeur commune est la probabilité de l'événement A .

Après avoir introduit les concepts de base des principales théories utilisées pour la représentation des connaissances imparfaites, nous allons présenter, dans la suite, un état de l'art des principales méthodes de sûreté de fonctionnement (arbres de défaillance, approches markoviennes et réseaux de Petri) basées sur l'utilisation de ces théories et qui permettent de prendre

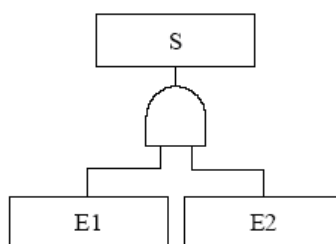


FIG. 2.3 – Arbre de défaillance

en compte les paramètres de fiabilité imparfaits des composants dans l'évaluation de la sûreté de fonctionnement des systèmes.

2.3 Méthodes de sûreté de fonctionnement en présence d'informations imparfaites

2.3.1 Arbres de défaillance

2.3.1.1 Principes, construction et analyse

La méthode des arbres de défaillance a été élaborée en 1961 [Watson, 1961] par la compagnie américaine Bell Telephone et fut exploitée pour évaluer et améliorer la fiabilité des systèmes de tir de missiles. En 1965, Boeing et l'université de Washington ont organisé le premier congrès traitant de l'utilisation des arbres de défaillance dans les études de fiabilité [Mearns, 1965]. Depuis, cette méthode a été largement utilisée dans de nombreux domaines industriels (chimique, nucléaire, aéronautique, etc.). Cette méthode est aussi connue sous le nom de méthode des arbres de causes (MAC) et méthode des arbres de défauts.

Dans cette partie, nous présentons les principes de construction des arbres de défaillance et ceux relatifs à leur analyse quantitative. Une conclusion résumant les avantages et les inconvénients de cette méthode est également présentée [Aubry, 2003].

Principes

Pour construire un arbre de défaillance, toutes les combinaisons possibles d'événements entraînant la réalisation de l'événement non désiré sont recensées. Ensuite, nous présentons graphiquement ces combinaisons au moyen d'une structure arborescente dont l'événement non désiré est le tronc. Nous construisons l'arbre de manière déductive, jusqu'à l'obtention des causes premières pouvant entraîner l'événement redouté, ce sont les événements de base. L'arbre de défaillance est la combinaison des événements de base menant à l'événement redouté, les combinaisons étant construite à l'aide de portes logiques. Dans la figure 2.3, l'événement redouté S est la combinaison des deux événements de base E1 et E2.

Construction

La détermination de l'événement non désiré dépend de la nature de l'analyse que l'on souhaite réaliser sur le système. Ainsi, pour une analyse de sécurité, l'événement non désiré est l'événement engendrant une catastrophe, tandis que pour l'analyse de fiabilité, l'événement non

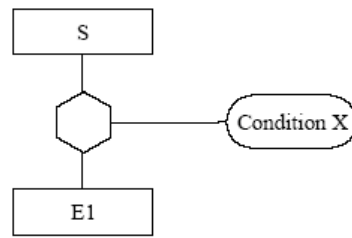


FIG. 2.4 – Opérateur de condition

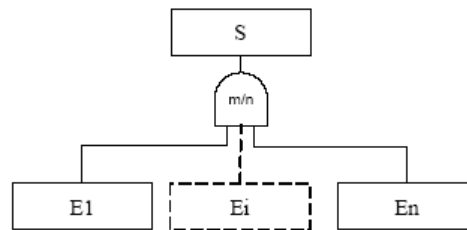


FIG. 2.5 – Opérateur " m/n "

désiré est que le système tombe en panne.

Les opérateurs les plus souvent utilisés sont les opérateurs logiques "ET" et "OU". Néanmoins, d'autres opérateurs existent, tel que l'opérateur de condition (cf. figure 2.4), ou encore l'opérateur m/n (cf. figure 2.5).

Trois types d'événements constituent un arbre de défaillance : les événements de base, l'événement non désiré et les événements intermédiaires, combinaison d'événements de base menant vers l'événement non désiré. L'élaboration d'un arbre de défaillance se fait généralement en recherchant les causes immédiates, nécessaires et suffisantes (INS) [Villemeur, 1998]. Cette recherche doit être faite de la manière la plus rigoureuse possible puisqu'elle conditionne la suite de l'étude. En effet, elle permet de détecter les événements intermédiaires (cf. figure 2.6), que l'on peut ensuite classer entre :

- Événement de base.
- Événement défaillance de composant.
- Événement défaillance du système.

De la sorte, il est possible d'établir un arbre reliant l'événement non désiré uniquement à des événements de base. La figure 2.7 résume le processus d'élaboration d'un arbre de défaillance.

Analyse quantitative

Pour analyser l'arbre de défaillance, on peut rechercher les coupes minimales. Une coupe est un ensemble d'événements entraînant l'événement indésirable. Une coupe est minimale si elle ne contient pas d'autres coupes.

Nous calculons la probabilité de l'événement non désiré S à partir de celles des coupes mini-

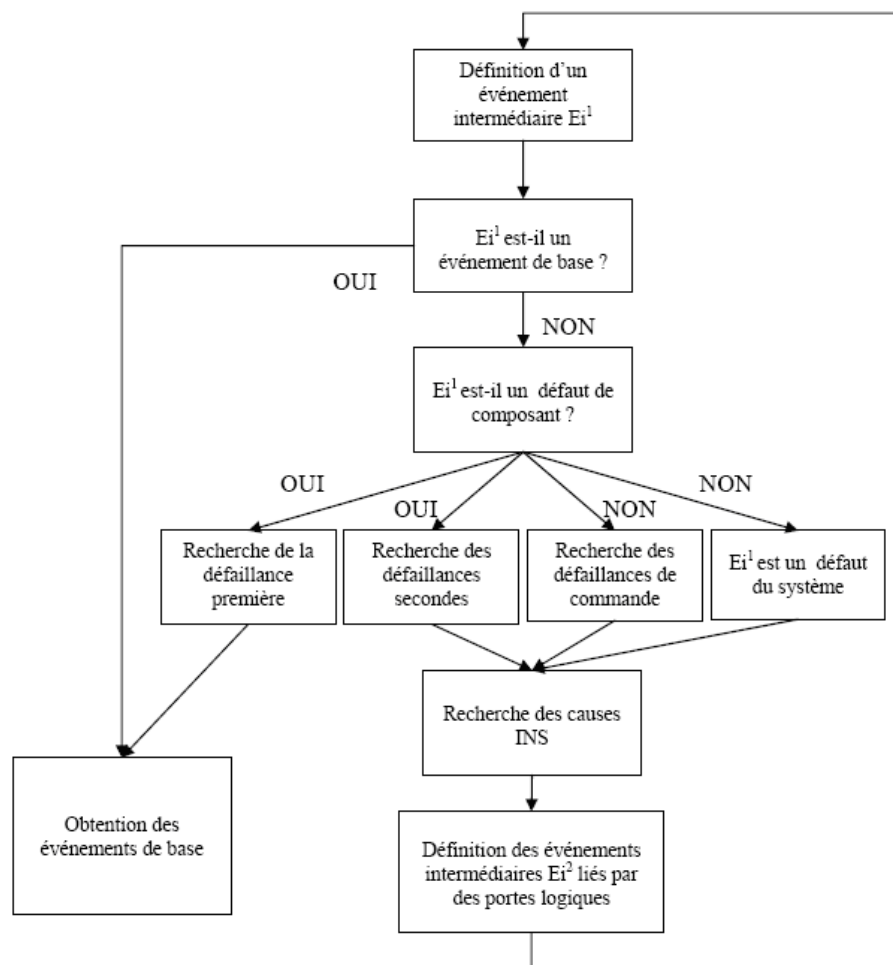


FIG. 2.6 – Processus d'analyse des événements intermédiaires [Villemeur, 1998]

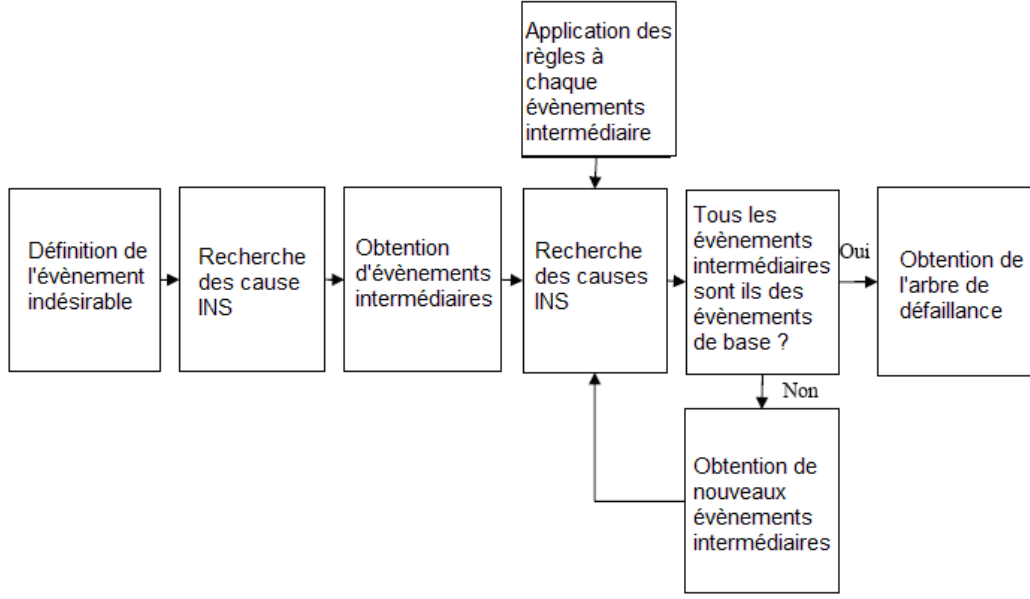


FIG. 2.7 – Processus d'élaboration d'un arbre de défaillance [Villemeur, 1998]

males C_i par la relation :

$$P[S] = P[C_1 + C_2 + \dots + C_n] \quad (2.46)$$

Où "+" représente l'opérateur "OU" et n le nombre de coupes minimales de l'arbre.

Deux coupes sont corrélées si elles contiennent des éléments en commun. On doit donc utiliser le théorème de Sylvester-Poincaré pour effectuer le calcul de $P[S]$:

$$P[S] = \sum_{i=1}^n P[C_i] - \sum_{j=2}^n \sum_{i=1}^{j-1} P[C_i.C_j] + \sum_{j=3}^n \sum_{k=2}^{j-1} \sum_{i=1}^{k-1} P[C_i.C_j.C_k] - \dots + (-1)^n P[C_1.C_2 \dots C_n] \quad (2.47)$$

L'équation 2.47 est particulièrement fastidieuse à calculer avec l'augmentation du nombre de coupes. Aussi, plusieurs méthodes permettent de calculer efficacement la probabilité d'occurrence de l'évènement sommet de l'arbre comme l'algorithme d'Abraham [Abraham, 1979] et la méthode des diagrammes de décision binaires (BDD) [Rauzy *et al.*, 2003].

Conclusion

La méthode des arbres de défaillance est aujourd'hui l'une des méthodes les plus utilisées dans les analyses de sûreté de fonctionnement des systèmes. Elle a pour objectif le recensement des causes entraînant l'apparition des événements indésirables et de calculer leur probabilités d'occurrence. Elle constitue donc un moyen de représentation de la logique des défaillances et d'identification des points faibles de la conception. Toutefois, cette méthode est peu adaptée pour l'étude des systèmes élémentaires présentant des défaillances de causes communes et des taux de défaillance dépendants du temps.

2.3.1.2 Arbres de défaillances flous

Les premiers travaux d'analyse floue des arbres de défaillance appartiennent à Tanaka *et al.* [Tanaka *et al.*, 1983]. Ces travaux sont basés sur la représentation de la probabilité d'occurrence des événements de base par des nombres flous trapézoïdaux et l'utilisation du principe d'extension de Zadeh [Zadeh, 1968] pour calculer la probabilité d'occurrence de l'événement sommet de l'arbre. Singer [Singer, 1990] a analysé les arbres de défaillance en représentant les probabilités d'occurrence des événements de base par des nombres flous du type L-R pour simplifier les opérations arithmétiques floues. Liang et Wang [Liang and Wang, 1993] ont représenté les probabilités d'occurrence des événements de base par des nombres flous triangulaires pour évaluer la possibilité d'occurrence de l'événement sommet de l'arbre. Soman et Misra [Soman and Misra, 1993] ont proposé une méthode connue sous le nom de l'identité de résolution basée sur la méthode des α -coupes pour traiter les arbres de défaillance comportant des événements répétés. Sawyer et Rao [Sawyer and Rao, 1994] ont utilisé la méthode des α -coupes pour étudier la fiabilité d'un système mécanique. Par la suite, un nombre important de travaux a été développé sur l'utilisation des arbres de défaillance flous [P. V. Suresh and Raj, 1996; Feng and Wu, 2001; Huang *et al.*, 2004; Sallak *et al.*, 2007a]. Les deux principales méthodes de traitement des arbres de défaillance flous sont la méthode de Singer [Singer, 1990] et la méthode des α -coupes [Soman and Misra, 1993].

2.3.1.3 Méthode de Singer

L'analyse floue des arbres de défaillance est basée sur le fait de considérer les probabilités d'occurrence des événements de base comme des nombres flous. Le principe d'extension flou, défini par Zadeh [Zadeh, 1978], est ensuite appliqué pour le calcul des fonctions d'appartenance des différents événements intermédiaires de l'arbre. Nous commençons par considérer les probabilités d'occurrence des événements de base comme des nombres flous du type L-R. Puis, nous utilisons les opérateurs définis précédemment pour ces nombres afin d'évaluer la probabilité floue d'occurrence de l'événement sommet [Singer, 1990; Sallak *et al.*, 2005]. Nous nous intéressons en particulier aux deux opérateurs ($\times$) (porte ET) et ($+$) (porte OU).

La forme floue d'une association d'opérateurs ET (cf. figure 2.8) est :

$$\widetilde{P}_Y = \prod_{i=0}^n p_i = AND(p_0, p_1, \dots, p_n) \quad (2.48)$$

où p_i est la probabilité floue d'occurrence de l'événement i et $\widetilde{P}_Y$ est la probabilité floue d'occurrence de l'événement sommet Y . En utilisant les opérations sur les nombres flous du type $L - R$, nous obtenons :

$$\widetilde{P}_Y = (m_Y, a_Y, b_Y) = (m_{p_0}, a_{p_0}, b_{p_0}) \times (m_{p_1}, a_{p_1}, b_{p_1}) \times \dots \times (m_{p_n}, a_{p_n}, b_{p_n}) \quad (2.49)$$

D'où :

$$\widetilde{P}_Y = (m_{p_{i-1}} m_{p_i}, m_{p_{i-1}} a_{p_i} + m_{p_i} a_{p_{i-1}}, m_{p_{i-1}} b_{p_i} + m_{p_i} b_{p_{i-1}}) |_{i=1, \dots, n} \quad (2.50)$$

Où m_{p_i} , a_{p_i} et b_{p_i} sont obtenues avec les formules de récurrences suivantes :

Pour $i = 1, 2, \dots, n$,

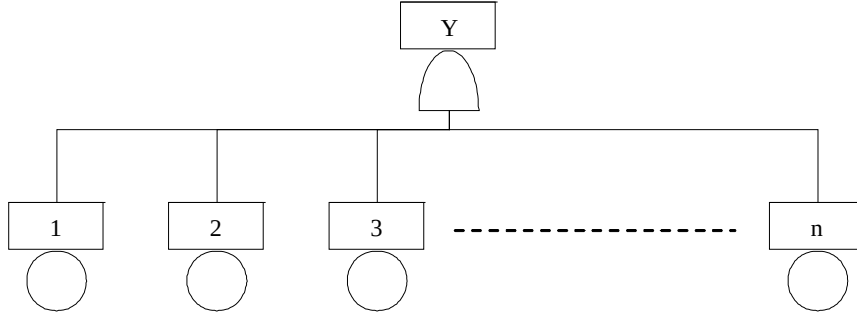


FIG. 2.8 – Arbre de défaillance avec un opérateur ET

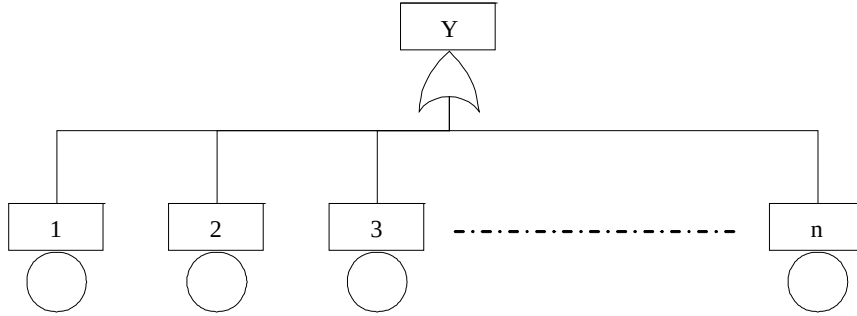


FIG. 2.9 – Arbre de défaillance avec un opérateur OU

$$m_{p_0} = m_{p_1}, \quad m_{p_1} = m_{p_1} m_{p_2}, \quad \dots, \quad (2.51)$$

$$a_{p_0} = a_{p_1}, \quad a_{p_1} = m_{p_1} a_{p_2}, \quad \dots, \quad (2.52)$$

$$b_{p_0} = b_{p_1}, \quad b_{p_1} = m_{p_1} b_{p_2}, \quad \dots, \quad (2.53)$$

De même, la forme floue d'une association d'opérateurs OU (cf. figure 2.9) est :

$$\widetilde{P}_Y = 1 - \prod_{i=0}^n (1 - p_i) = OR(p_0, p_1, \dots, p_n) \quad (2.54)$$

Avec :

$$\widetilde{P}_Y = (m_Y, a_Y, b_Y) \quad (2.55)$$

$$\widetilde{P}_Y = (1, 0, 0) - ((1, 0, 0) - (m_{p_0}, a_{p_0}, b_{p_0})) \times ((1, 0, 0) - (m_{p_1}, a_{p_1}, b_{p_1})) \times \dots \times ((1, 0, 0) - (m_{p_n}, a_{p_n}, b_{p_n})) \quad (2.56)$$

Nous obtenons finalement :

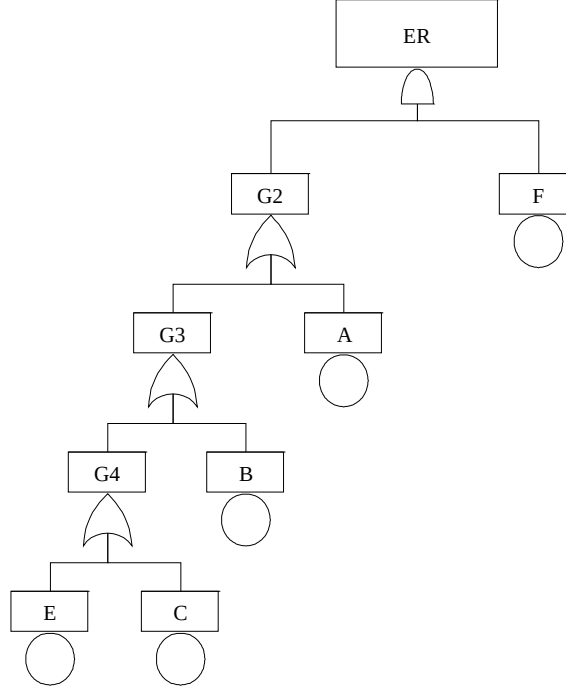


FIG. 2.10 – Arbre de défaillance simplifié

$$\widetilde{P}_Y = 1 - (m_{p_{i-1}}(1 - m_{p_i}), m_{p_{i-1}}a_{p_i} + (1 - m_{p_i})a_{p_{i-1}}, m_{p_{i-1}}b_{p_i} + (1 - m_{p_i})b_{p_{i-1}})|_{i=1,\dots,n} \quad (2.57)$$

Où m_{r_i} , a_{r_i} et b_{r_i} sont obtenues avec les formules de récurrences suivantes :

Pour $i = 1, 2, \dots, n$:

$$m_{p_0} = m_{p_1}, \quad m_{p_1} = (1 - m_{p_1})(1 - m_{p_2}), \quad \dots, \quad (2.58)$$

$$a_{p_0} = a_{p_1}, \quad a_{p_1} = (1 - m_{p_1})a_{p_2} + (1 - m_{p_2})a_{p_1}, \quad \dots, \quad (2.59)$$

$$b_{p_0} = b_{p_1}, \quad b_{p_1} = (1 - m_{p_1})b_{p_2} + (1 - m_{p_2})b_{p_1}, \quad \dots, \quad (2.60)$$

2.3.1.4 Exemple

Considérons un système dont l'arbre de défaillance simplifié est représenté sur la figure 2.10. L'objectif est de calculer la probabilité d'occurrence de l'événement redouté (ER).

TAB. 2.2 – Paramètres flous des probabilités d'occurrence des événements de base

Événements de base	a_i	m_i	b_i
Défaillance première E	0.008	0.01	0.012
Défaillance première C	0.001	0.002	0.0025
Défaillance première B	0.085	0.09	0.092
Défaillance première A	0.078	0.08	0.0805
Défaillance première F	0.0695	0.07	0.0722

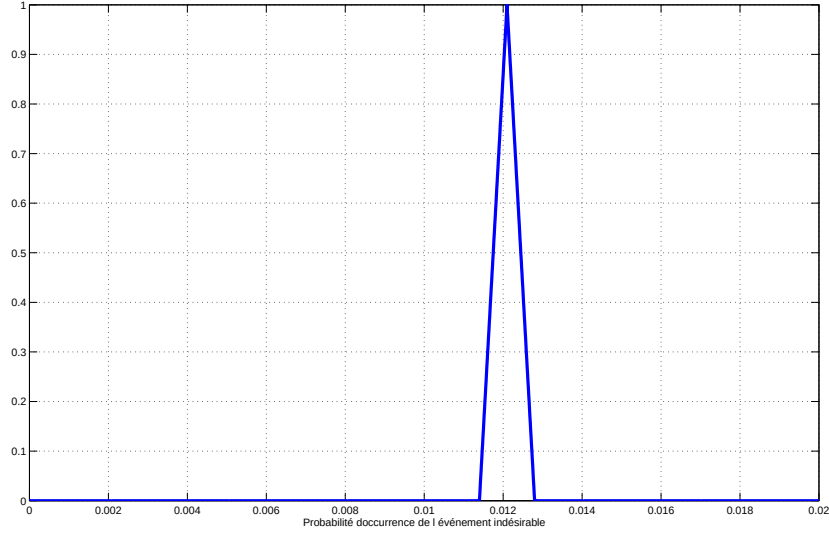


FIG. 2.11 – Probabilité d'occurrence de l'événement "surchauffe du fil AB"

Les probabilités floues d'occurrence des événements de base sont des nombres flous du type L-R caractérisés par les 3 paramètres m_i , a_i et b_i , tel que m_i est la valeur modale avec $\mu_{PFD}(m_i) = 1$, a_i est la limite à gauche de m_i et b_i est la limite à droite de m_i (cf. figure 2.1). Dans le tableau 2.2, nous donnons les valeurs des 3 paramètres m_i , a_i et b_i pour chaque événement de base.

En utilisant les relations de récurrences définies pour les portes ET et OU, nous obtenons la probabilité d'occurrence de l'événement redouté (ER) (cf. figure 2.11).

En conclusion, nous pouvons dire que la méthode de Singer apparaît comme étant une excellente approximation, mais qui limite le choix de la probabilité floue de l'événement de base au seul nombre flou de type L-R.

2.3.1.5 Méthode des α -coupes

Pour appliquer cette méthode, il faut considérer les probabilités d'occurrence des événements de base comme des nombres flous. Puis, il faut appliquer les opérateurs d'addition pour les portes "OU" de l'arbre et de multiplication pour les portes "ET" pour calculer la probabilité d'occurrence de l'événement sommet à chaque α -coupe (cf. section 2.2.2.4). Enfin, la distribution de la probabilité floue d'occurrence de l'événement sommet est construite à partir des probabilités

d'occurrence calculées à chaque α -coupe.

En outre, nous pouvons conclure que la méthode des α -coupes est plus précise que la méthode de Singer, mais elle nécessite une plus grande puissance de calcul [Kaufman and Gupta, 1991]. Elle a d'autre part l'avantage de permettre l'utilisation de n'importe quelle forme de probabilités floues pour l'événement de base.

2.3.1.6 Cas des événements répétés

Dans le cas où l'arbre de défaillance contient un ou plusieurs événements répétés, il faut d'abord recenser les coupes minimales de l'arbre, puis nous utilisons la méthode d'Abraham [Abraham, 1979] pour écrire un polynôme constitué des coupes minimales sans répétition de variables dans les différents monômes. Enfin, il faut appliquer l'opérateur d'addition flou pour évaluer la probabilité d'occurrence floue de l'événement sommet de l'arbre à partir des probabilités d'occurrence floues des monômes.

2.3.1.7 Dépendance des événements

L'arbre de défaillance suppose une totale indépendance entre les événements de base. Or ceux-ci peuvent être liés par des relations de dépendance quand la réalisation d'un événement favorise (dépendance stochastique) ou impose (forçage d'état) la réalisation d'autres événements de l'arbre.

La notion de dépendance dans les arbres de défaillance flous a été traitée de deux façons différentes proposées par Onisawa [Onisawa, 1988] et Misra et Weber [Misra and Weber, 1989]. Dans le premier cas, la dépendance est prise en compte en utilisant des relations causales floues. Dans le deuxième, la dépendance est prise en compte en utilisant un facteur d qui modélise le degré de dépendance entre les occurrences des événements.

Le facteur d peut être un nombre réel ou un nombre flou. La probabilité d'occurrence des deux événements A et B est donnée par la relation :

$$p(A \cap B) = p(A \setminus B) \cdot p(B) \quad (2.61)$$

En posant :

$$p(A \setminus B) = d + (1 - d) * p(A) \quad (2.62)$$

Nous obtenons :

$$p(A \cap B) = (d + (1 - d) * p(A)) \cdot p(B) = (1 - (1 - d)(1 - p(A)))p(B) \quad (2.63)$$

Où d est le degré de dépendance des occurrences des événements A et B .

Si nous considérons l'opérateur d'union, nous obtenons :

$$p(A \cup B) = p(A) + p(B) - p(A \cap B) = p(A) + p(B) - (1 - (1 - d)(1 - p(A)))p(B) \quad (2.64)$$

D'où :

$$p(A \cup B) = 1 - (1 - p(A))(1 - (1 - d)p(B)) \quad (2.65)$$

2.3.1.8 Remarque

Bien que la méthode des α -coupes soit la plus souvent présentée, il existe d'autres méthodes qui permettent de traiter les arbres de défaillance en présence d'incertitudes et d'imprécisions [Simon and Weber, 2007] comme la méthode de Guth [Guth, 1991] qui considère que la probabilité d'occurrence p d'un événement A est telle que :

$$p \in [bel(A), pls(A)] \quad (2.66)$$

La crédibilité bel représente la borne inférieure et la plausibilité pls représente la borne supérieure. Il définit ainsi les trois ensembles focaux Vrai (T), Faux (F) et Inconnu (T, F) :

$$m(T) = bel(A) \quad (2.67)$$

$$m(F) = 1 - pls(A) \quad (2.68)$$

$$m(T, F) = pls(A) - bel(A) \quad (2.69)$$

2.3.2 Approches markoviennes

La théorie des processus de Markov fournit un outil efficace pour le calcul des paramètres de sûreté de fonctionnement des systèmes en fonction des paramètres de sûreté de fonctionnement des composants élémentaires utilisés. Cette théorie est particulièrement adaptée à l'étude des systèmes réparables modélisables par des processus Markoviens et les systèmes dont les taux de défaillance varient au cours du temps.

Les bases théoriques concernant les approches markoviennes floues ont été introduites par Bellman et Zadeh [Bellman and Zadeh, 1970] qui ont été les premiers à étudier des processus décisionnels flous. Esogbue et Bellman [Esogbue and Bellman, 1984] se sont intéressés à plusieurs types de programmations dynamiques floues avec un nombre fini d'états et d'actions. Kurse *et al.* [Kurse *et al.*, 1987] ont introduit les chaînes de Markov floues basées sur l'utilisation des probabilités floues pour représenter les éléments des matrices de transition afin de calculer la puissance efficace d'un processeur. En 1998, Zadeh a proposé des algorithmes de Markov flous [Zadeh, 1998] utilisant des chaînes de Markov floues. Bhattacharyya [Bhattacharyya, 1999] a prouvé par la suite que, si une chaîne de Markov est irréductible alors la chaîne de Markov floue qui lui est associée est également irréductible. Il a utilisé ce résultat pour étudier des processus décisionnels flous. En 2004, Symeonaki [Symeonaki and Stamou, 2004] a défini des conditions de convergence des chaînes de Markov non homogènes à états flous. En même temps, Tanriovenaa *et al.* [Tanrioven *et al.*, 2004] ont réalisé plusieurs applications des chaînes de Markov floues, en se basant notamment sur les systèmes d'inférences et les probabilités de transitions floues.

L'idée directrice de cette section est de recenser les différentes méthodes markoviennes définies dans la littérature et plus particulièrement celles basées sur l'utilisation du formalisme flou. Nous énoncerons les théorèmes et les définitions relatives à ce formalisme, ainsi que quelques applications basées sur les systèmes d'inférences et la modélisation floue des systèmes Markoviens.

2.3.2.1 Processus stochastiques

Préliminaires

- Soit X un ensemble. Un ensemble Λ de parties de X est appelé une σ -algèbre, encore appelée une tribu) sur X si :

1. L'ensemble vide est dans Λ .
2. Λ est stable par passage au complémentaire, *i.e.* :

$$\forall A \in \Lambda, \quad A^c \in \Lambda \quad (2.70)$$

Où A^c est le complémentaire de A .

3. Λ est stable par union dénombrable, *i.e.* :

$$\forall n \in N, \quad A_n \in \Lambda \quad \text{alors} \quad \cup_{n \in N} A_n \in \Lambda \quad (2.71)$$

- Soit Ω un ensemble non vide, Λ une σ -algèbre des événements sur Ω et une fonction P de Λ dans $\mathbb{R}^+$ possédant les propriétés suivantes :

1. $P(\Omega) = 1$.
2. $A_1, A_2, \dots, A_n$ étant une suite finie ou dénombrable d'ensembles disjoints deux à deux de Λ :

$$P(A_1 \cup A_2 \cup \dots \cup A_n) = P(A_1) + P(A_2) + \dots + P(A_n) \quad (2.72)$$

La fonction P est appelé probabilité et le triplet (Ω, Λ, P) est appelé espace de probabilité.

- Soit (Ω, Λ, P) un espace de probabilité et v une application de Ω dans R . On dit que v est une variable aléatoire si :

$$\forall x \in R, \quad v^{-1}(]-\infty, x]) \in \Lambda \quad (2.73)$$

Une telle variable peut être continue ou discontinue. Par exemple, la durée de vie d'un composant est une variable aléatoire.

- Un process stochastique est une application d'un ensemble d'indices S dans l'ensemble des variables aléatoires définies sur (Ω, Λ, P) [Pages and Gondran, 1980]. Nous utiliserons souvent comme éléments de l'ensemble S le paramètre t désignant le temps. Nous considérons généralement deux modèles possibles :
 - si les ensembles Ω et S sont discrets, nous parlerons de chaîne stochastique,
 - si l'ensemble Ω est discret alors que S est continu, nous parlerons de processus stochastique.

2.3.2.2 Chaînes de Markov

Définitions et propriétés

Définition 1

Un système Markovien est un système qui transite de l'état i à l'état j avec une probabilité $p_{i,j}(t)$ qui ne dépend que des états i et j (système sans mémoire).

La matrice $P = (p_{i,j})$ est stochastique, chaque ligne i contient toutes les transitions possibles en partant de i . La somme des taux de transition de chaque ligne est égale à 1.

Propriété

Les vecteurs stochastiques $p(t)$ vérifient la formule suivante :

$$(p_0(t+1), p_1(t+1), \dots, p_n(t+1)) = (p_0(t), p_1(t), \dots, p_n(t)) \begin{pmatrix} p_{0,0} & p_{0,1} & \dots & p_{0,n} \\ p_{1,0} & p_{1,1} & \dots & p_{1,n} \\ \dots & \dots & \dots & \dots \\ p_{n,0} & p_{n,1} & \dots & p_{n,n} \end{pmatrix} \quad (2.74)$$

L'équation (2.74) peut s'écrire sous la forme compacte :

$$p(t+1) = p(t)P \quad (2.75)$$

Définition 2 (Equation de Chapman-Kolmogorov en temps discret)

Une chaîne de Markov homogène dans le temps est la donnée d'une suite de vecteurs stochastiques $\{p(t)\}$ et d'une matrice stochastique P telle que :

$$\forall t \in N \quad p(t+1) = p(t)P \quad (2.76)$$

La chaîne est dite homogène dans le temps parce que la matrice P ne dépend pas du temps.

Évolution temporelle

L'objectif est d'obtenir une expression générale qui donne l'état d'un système à un instant t en fonction de son état initial. En appliquant la formule de récurrence $p(t+1) = p(t)P$ aux différents instants $t = 0, 1, \dots, n$, nous obtenons les relations suivantes :

$$p(1) = p(0)P \quad (2.77)$$

$$p(2) = p(1)P = p(0)P^2 \quad (2.78)$$

Finalement, à un instant t donné, nous obtenons la formule générale suivante :

$$p(t) = p(0)P^t \quad (2.79)$$

Distribution limite

Dans la plupart des cas, la distribution $p(t)$ converge vers une distribution limite, notée $p(\infty)$, lorsque t tend vers l'infini.

Pour qu'une distribution limite existe et ne dépende pas de la distribution initiale $p(0)$, il suffit qu'au moins une puissance de la matrice de transition P n'ait que des composantes strictement positives. Cela signifie que le processus, étudié pendant un temps suffisamment long peut transiter entre deux états quelconques avec une probabilité strictement positive.

Propriété

Si une certaine puissance de la matrice de transition P n'a que des composantes strictement positives, alors :

- Il existe une distribution limite $p(\infty)$ dont toutes les composantes sont strictement positives et telle que $p(t) \mapsto p(\infty)$ lorsque $t \mapsto \infty$ qui est indépendante du vecteur stochastique initial $p(0)$.
- La suite des matrices P^t lorsque $t \mapsto \infty$ converge vers la matrice P^∞ dont toutes les lignes sont égales au vecteur $p(\infty)$.

Distribution stationnaire

Une distribution stochastique $p = (p_0, p_1, \dots, p_n)$ est dite stationnaire par rapport à la matrice stochastique P si et seulement si :

$$p = p.P \quad (2.80)$$

Cette condition s'écrit de manière équivalente $p(P - Id) = 0$. Elle traduit le fait que si la chaîne de Markov est initialisée avec la distribution $p(0) = p$, alors le vecteur stochastique $p(t)$ est constant pour $t = 0, 1, \dots, n$.

Propriété 1

Une chaîne de Markov finie admet au moins une distribution stationnaire, ce qui n'est plus nécessairement vrai si l'espace des états est infini.

Propriété 2

Si la distribution limite $p(\infty)$ d'une chaîne de Markov existe et ne dépend pas du vecteur stochastique initial $p(0)$, alors $p(\infty)$ est l'unique distribution stationnaire de cette chaîne.

Dans ce cas, le système est ergodique, ce qui signifie que le pourcentage de temps passé par le processus étudié dans un état donné est égal à la probabilité stationnaire de cet état. En d'autres termes, sur un intervalle de temps assez grand, la proportion de temps passé par le système dans un état j est égale à la probabilité de trouver ce système dans l'état j à la fin de cet intervalle. En effet, quand il y'a ergodicité, nous déduisons que les moyennes temporelles et les moyennes spatiales coïncident.

Propriété 3

Une Chaîne de Markov est dite irréductible quand tous les états communiquent. Nous pouvons donc passer d'un point à un autre soit directement soit par l'intermédiaire d'autres points.

Propriété 4

Une chaîne de Markov est dite périodique si sa période est supérieure à 1. Dans le cas contraire, elle est dite apériodique.

Propriété 5

Une chaîne de Markov est dite ergodique si elle est irréductible et apériodique.

2.3.2.3 Processus de Markov

Dans cette partie, le temps t est un nombre réel positif ou nul. Les concepts introduits dans le cas discret ont tous une analogie dans le cas continu.

Définitions et propriétés

Définition

Une matrice carrée $A = (a_{i,j})$ est appelée générateur stochastique infinitésimal si :

1. La somme des éléments d'une ligne quelconque est nulle.
2. Les éléments sur la diagonale sont négatifs ou nuls et les autres sont positifs ou nuls.

Son interprétation est la suivante : quand le processus entre dans l'état i , il y reste une durée aléatoire de distribution exponentielle de paramètre $-a(i, i)$, puis saute vers l'état j avec la probabilité $\frac{-a(i, i)}{a(i, j)}$.

Propriété (Equation de Chapman-Kolmogorov en temps continu)

La fonction $t \rightarrow p(t)$ est un processus de Markov si et seulement si le vecteur stochastique $p(t)$ vérifie l'équation différentielle :

$$\frac{d}{dt}p(t) = p(t)P \quad (2.81)$$

où P est un générateur stochastique infinitésimal.

Évolution temporaire

L'équation d'état ordinaire $\frac{d}{dt}p(t) = p(t)P$ admet pour solution le vecteur $p(t) = p(0)\exp(tP)$ avec :

$$\exp(tP) = 1 + tP + \frac{1}{2!}t^2P^2 + \frac{1}{3!}t^3P^3 + \dots \quad (2.82)$$

Si la matrice P est un générateur stochastique infinitésimal, alors la matrice $\exp(tP)$ est, à tout instant t , une matrice stochastique.

La résolution de l'équation d'état (2.81) nous ramène à un calcul des exponentielles des matrices de transition. Dans le cas continu, le calcul de l'exponentielle est basé sur le calcul des valeurs propres de P . Lorsque P est une matrice diagonale alors :

$$P = \text{Diag}(\lambda_0, \lambda_1, \dots, \lambda_n) \quad (2.83)$$

Dans ce cas, l'exponentielle de la matrice P est donnée par la relation :

$$\exp(tP) = \text{Diag}(\exp(t\lambda_0), \exp(t\lambda_1), \dots, \exp(t\lambda_n)) \quad (2.84)$$

Si P n'est pas une matrice diagonale, il existe une matrice de passage Q réversible telle que :

$$P = Q.D.Q^{-1} \quad (2.85)$$

où D est une matrice diagonale.

Dans ce cas, l'exponentielle de la matrice P est donnée par la relation :

$$\exp(tP) = Q.\exp(tD).Q^{-1} \quad (2.86)$$

Distribution limite

De même que dans le cas discret, pour qu'une distribution limite existe et ne dépende pas de la distribution initiale $p(0)$, il suffit qu'au moins une puissance de la matrice de transition P n'ait que des composantes strictement positives. Cela signifie que le processus, étudié pendant un temps suffisamment long, peut transiter entre deux états quelconques avec une probabilité strictement positive.

Propriété

Si une certaine puissance de la matrice de transition P n'a que des composantes strictement positives, alors :

- Il existe une distribution limite $p(\infty)$ dont toutes les composantes sont strictement positives et telle que $p(t) \mapsto p(\infty)$ lorsque $t \mapsto \infty$ qui est indépendante du vecteur stochastique initial $p(0)$.
- La suite des matrices P^t lorsque $t \mapsto \infty$ converge vers la matrice P^∞ dont toutes les lignes sont égales au vecteur $p(\infty)$. De plus, la distribution $p(\infty)$ est stationnaire *i.e.* $p(\infty)P = 0$.

Distribution stationnaire

La distribution stochastique stationnaire $p = (p_0, p_1, \dots, p_n)$ est donnée par :

$$p.P = 0 \tag{2.87}$$

Cette condition traduit le fait que si la chaîne de Markov est initialisée avec la distribution $p(0) = p$, alors le vecteur stochastique $p(t)$ est constant pour $t = 0, 1, \dots, n$.

2.3.2.4 Observations et limitations

Dans certaines situations, nous sommes confrontés au fait que certains états que peut prendre un système sont flous, c'est-à-dire que nous ne pouvons pas définir ces états d'une manière précise. En effet, l'approche Markovienne classique est basée sur le fait qu'à tout instant t le système est dans un état précis donné. Or, parfois le système a de grandes chances d'être dans un état i sans que nous soyons pour autant sûr qu'il est effectivement dans cet état. Dans ce cas l'introduction d'états flous auxquels appartient le système avec un certain degré s'impose. Il est alors nécessaire de définir des chaînes de Markov qui utilisent le formalisme flou. Dans la littérature [Kurse *et al.*, 1987; Bhattacharyya, 1999; Symeonaki and Stamou, 2004; Tanrioven *et al.*, 2004], nous avons distingué trois approches principales :

1. Chaînes de Markov à états non flous et probabilités de transition floues [Kurse *et al.*, 1987; Bhattacharyya, 1999].
2. Chaînes de Markov à états flous et non flous et probabilités de transitions précises [Kurse *et al.*, 1987; Bhattacharyya, 1999; Symeonaki and Stamou, 2004].
3. Modèle de Markov avec un système d'inférence flou (MMF) [Tanrioven *et al.*, 2004].

2.3.2.5 Chaînes de Markov à états non flous et probabilités de transition floues

Dans cette partie, on suppose que les états que peut prendre le système, sont définis d'une manière claire et précise. Les probabilités des matrices de transition seront représentées par des nombres flous [Kurse *et al.*, 1987; Bhattacharyya, 1999].

Définitions et propriétés

Définition 1

Un sous-ensemble (ou une distribution) flou, sur l'ensemble des indices S , est défini par l'application $p : S \mapsto [0, 1]$, qui associe à chaque état du système un réel appartenant à l'intervalle $[0, 1]$ qui exprime le degré d'appartenance du système à cet état.

Dans la suite, nous définissons le vecteur d'état flou p par : $p = (p_0, p_1, \dots, p_n)$, où p_i désigne la possibilité que le système soit dans l'état i . L'ensemble des parties de S est noté $P(S)$.

Définition 2

Une relation floue P est définie par un ensemble flou sur le produit cartésien $S \times S$. Cette relation floue sera représentée par la matrice $P = (p_{i,j})$, où $p_{i,j}$ désigne le terme $P(i, j)$ qui représente le degré de possibilité pour que le système transite de l'état i à l'état j .

Définition 3

A tout instant t , l'état d'un système est complètement décrit par le sous-ensemble (ou la distribution) flou $p(t) = (p_0(t), p_1(t), \dots, p_n(t))$. La loi de transition de la chaîne de Markov floue est donnée par la relation floue P :

$$p_j(t+1) = \max_{i \in S} \{p_i(t) \wedge p_{i,j}\}, \quad j \in S \quad (2.88)$$

où $p(0)$ désigne le vecteur d'état flou initial.

La loi de transition de la chaîne de Markov classique est définie par :

$$p_j(t+1) = \sum_{i \in S} p_i(t) p_{i,j}, \quad j \in S \quad (2.89)$$

A partir de cette équation, nous remarquons que les chaînes de Markov floues sont obtenues à partir des chaînes de Markov classiques en remplaçant les opérateurs des produits et des sommes par les opérateurs *min* et *max*, et en remplaçant les probabilités par des mesures de possibilité.

Évolution temporelle

L'état flou k d'une chaîne de Markov peut être caractérisé à tout instant $t = 1, \dots, n$ par la formule de récurrence :

$$p_k(t) = \max_{l \in S} \{p_l(0) \wedge p_{l,k}^t\} \quad k \in S \quad (2.90)$$

où :

- $p_k(t)$ désigne la probabilité pour que le système soit dans l'état k à l'instant t .
- $p_l(0)$ est la probabilité pour que le système soit dans l'état l à l'instant initial.
- $p_{l,k}^t$ représente le terme $p_{l,k}$ de la matrice de transition floue P^t .

– $S = \{1, \dots, n\}$ est l'ensemble des états.

L'équation (2.90) peut s'écrire en utilisant la notation matricielle sous la forme :

$$p(t) = p(0) \circ P^t \quad (2.91)$$

où :

- $p(t) = (p_0(t), p_1(t), \dots, p_n(t))$
- $p(0) = (p_0(0), p_1(0), \dots, p_n(0))$
-

$$P^t = \begin{pmatrix} p_{0,0}^t & p_{0,1}^t & \dots & p_{0,n}^t \\ p_{1,0}^t & p_{1,1}^t & \dots & p_{1,n}^t \\ \dots & \dots & \dots & \dots \\ p_{n,0}^t & p_{n,1}^t & \dots & p_{n,n}^t \end{pmatrix} \quad (2.92)$$

L'opérateur $\circ$ désigne la composition des deux opérateurs *min* et *max*, i.e. :

$$A \circ B = \max\{\min(A, B)\} \quad (2.93)$$

où A et B sont deux matrices floues quelconques.

La résolution de l'équation d'état flou (2.90) nécessite le calcul des puissances successives de la matrice de transition floue $P = (p_{i,j})$. Or, nous savons que pour tout instant $t = 1, \dots, n$, nous avons :

$$P^t = P.P^{t-1} \quad (2.94)$$

En formalisme flou, ce calcul de puissance se fait de manière itérative en utilisant la formule suivante pour $t = 1, \dots, n$:

$$\forall (i, j) \in N \times N \quad p_{i,j}^t = \max_{k \in S} \{p_{i,k} \wedge p_{k,j}^{t-1}\} \quad (2.95)$$

Avec :

$$p_{i,j}^1 = p_{i,j}, \quad p_{i,j}^0 = \delta_{i,j} \quad (2.96)$$

où $\delta_{i,j}$ est le symbole de Kronecker.

Convergence de la matrice de transition floue

Théorème 1

Les puissances de la matrice de transition floue soit convergent vers une matrice idempotente P^τ , où τ est un nombre fini, soit oscillent avec une période finie ν à partir d'une certaine puissance.

Théorème 2

Si les puissances de la matrice de transition convergent à partir d'un certain indice τ vers une solution non périodique, alors la chaîne de Markov est appelée chaîne non périodique et $P^* = P^\tau$ est appelée matrice de transition limite.

Dans le cas des chaînes de Markov classiques, les puissances des matrices de transition convergent, en général, vers une solution stationnaire en un nombre infini d'étapes (convergence

infinie), c'est-à-dire que la convergence a lieu quand t tend vers l'infini. Par contre, l'étude des chaînes de Markov à états non flous et probabilités de transition flous montre que les puissances des matrices de transition floues convergent en un nombre fini d'étapes (convergence finie). Cela est dû aux différences existant entre les propriétés ergodiques des chaînes de Markov à états non flous et probabilités de transition flous et les chaînes classiques.

2.3.2.6 Chaînes de Markov à états flous et non flous et probabilités de transition singulières

Dans cette partie, nous supposons que les états que peut prendre le système, peuvent être flous et non flous. Les probabilités des matrices de transition seront des valeurs singulières.

Notations et hypothèses

- $F = \{F_1, F_2, \dots, F_n\}$ est l'espace d'état flou du système, c'est-à-dire que F est l'ensemble des états flous du système.
- $S = \{1, \dots, m\}$ est l'espace d'état non flou du système.
- Pour $r = 1, \dots, n$: F_r est un sous-ensemble flou défini sur S auquel est associé la fonction d'appartenance $\mu_{F_r} : S \mapsto [0, 1]$.
- $\{F_1, F_2, \dots, F_n\}$ définissent une pseudo-partition floue sur S tel que $\sum_{r=1}^n \mu_{F_r}(i) = 1$, pour $i \in S$.
- $Q_{0,i}$ est le vecteur de probabilités initial selon lequel le système se trouve dans l'état non-flou i .
- Les matrices de transition associées aux états non-flous sont stationnaires (les probabilités de transition sont indépendantes du temps) et irréductibles.

Dans la plupart des cas, le nombre d'états flous est inférieur au nombre d'états non-flous m .

Préliminaires

Nous allons introduire quelques définitions relatives aux probabilités des événements flous qui seront utilisées par la suite.

Définition

Le sous-ensemble flou A de Ω est appelé événement flou de Ω . Si nous associons à cet événement flou la fonction d'appartenance $\mu_A : \Omega \mapsto [0, 1]$ alors, la probabilité de l'événement flou A donnée par Zadeh [Zadeh, 1968] est :

$$Prob(A) = \int_{\Omega} \mu_A(\omega) dP \quad (2.97)$$

Propriété 1 (Norme de Larsen)

Le produit de deux événements flous A et B est défini par :

$$A.B \leftrightarrow \mu_{A.B} = \mu_A \cdot \mu_B \quad (2.98)$$

Propriété 2

La probabilité conditionnelle de l'événement flou A par rapport à l'événement flou B est donnée par :

$$Prob(A \setminus B) = \frac{Prob(A.B)}{Prob(B)}, \quad Prob(B) > 0 \quad (2.99)$$

L'utilisation du produit $A.B$ au lieu de l'intersection $A \cap B$ permet d'avoir $Prob(A \setminus B) = Prob(A)$ si les événements flous A et B sont indépendants (comme dans le cas où A et B sont des événements non-flous indépendants). En effet, Zadeh [Zadeh, 1968] a démontré qu'il n'est pas possible d'obtenir cette équivalence si nous utilisons l'opérateur d'intersection $A \cap B$ au lieu du produit $A.B$.

Probabilités de transition des états flous

Soient X_t et X_t^f les états non-flous et flous du système à l'instant t et Q_{0i} la matrice des probabilités initiales relative à un état non-flou [Kurse *et al.*, 1987; Bhattacharyya, 1999; Symeonaki and Stamou, 2004]. La probabilité que le système transite de l'état flou F_r à l'état flou F_s est donnée par :

$$p_{F_r, F_s} = Prob(X_1^f = F_s \setminus X_0^f = F_r) = \frac{Prob(X_1^f = F_s, X_0^f = F_r)}{Prob(X_0^f = F_r)} \quad (2.100)$$

En utilisant la définition (2.97) donnée par Zadeh :

$$Prob(X_0^f = F_r) = \sum_{i=1}^m \mu_{F_r}(i) Prob(X_0 = i) = \sum_{i=1}^m \mu_{F_r}(i) Q_{0,i} \quad (2.101)$$

Nous avons aussi :

$$Prob(X_1^f = F_s, X_0^f = F_r) = \sum_{i=1}^m \sum_{j=1}^m Prob(X_1 = j, X_0 = i) \mu_{F_r, F_s}(i, j) \quad (2.102)$$

d'où :

$$Prob(X_1^f = F_s, X_0^f = F_r) = \sum_{i=1}^m \sum_{j=1}^m Prob(X_1 = j \setminus X_0 = i) Prob(X_0 = i) \mu_{F_r}(i) \mu_{F_s}(j) \quad (2.103)$$

par conséquent :

$$Prob(X_1^f = F_s, X_0^f = F_r) = \sum_{i=1}^m \sum_{j=1}^m p_{i,j} Q_{0,i} \mu_{F_r}(i) \mu_{F_s}(j) \quad (2.104)$$

d'où :

$$p_{F_r, F_s} = \frac{\sum_{i=1}^m \sum_{j=1}^m p_{i,j} Q_{0,i} \mu_{F_r}(i) \mu_{F_s}(j)}{\sum_{i=1}^m \mu_{F_r}(i) Q_{0,i}} \quad (2.105)$$

En éliminant les termes Q_{0i} , nous obtenons :

$$p_{F_r, F_s} = \frac{\sum_{i=1}^m \sum_{j=1}^m p_{i,j} \mu_{F_r}(i) \mu_{F_s}(j)}{\sum_{i=1}^m \mu_{F_r}(i)} \quad (2.106)$$

En posant $m_{F_r} = \sum_{i=1}^m \mu_{F_r}(i)$, nous obtenons finalement :

$$p_{F_r, F_s} = \frac{\sum_{i=1}^m \sum_{j=1}^m p_{i,j} \mu_{F_r}(i) \mu_{F_s}(j)}{m_{F_r}} \quad (2.107)$$

Equation d'état

Soient :

- Γ la matrice des fonctions d'appartenance des états flous :

$$\Gamma = \begin{pmatrix} \mu_{F_1}(1) & \mu_{F_2}(1) & \dots & \mu_{F_n}(1) \\ \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots \\ \mu_{F_1}(m) & \mu_{F_2}(m) & \dots & \mu_{F_n}(m) \end{pmatrix} \quad (2.108)$$

- P la matrice des probabilités de transition des états non-flous :

$$P = \begin{pmatrix} p_{1,1} & p_{1,2} & \dots & p_{1,m} \\ \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots \\ p_{m,1} & p_{m,2} & \dots & p_{m,m} \end{pmatrix} \quad (2.109)$$

- $P(F)$ la matrice des probabilités de transition des états flous :

$$P(F) = \begin{pmatrix} p_{F_1, F_1} & p_{F_1, F_2} & \dots & p_{F_1, F_n} \\ \dots & \dots & \dots & \dots \\ \dots & \dots & \dots & \dots \\ p_{F_n, F_1} & p_{F_n, F_2} & \dots & p_{F_n, F_n} \end{pmatrix} \quad (2.110)$$

L'équation d'état est donnée par :

$$P(F) = \Delta_1 \Gamma^t \Delta_2 P \Gamma \quad (2.111)$$

Où :

- $\Delta_1 = \text{diag}(\delta_{1,r})$ avec $\delta_{1,r} = (m_{F_r})^{-1}$.
- $\Delta_2 = \text{diag}(Q_{0,i})$.

Exemple

Nous considérons un système simple constitué de deux composants en redondance active [Tanrioven *et al.*, 2004]. Ce système ne peut être que dans l'un des quatre états suivants :

- État 1 : Les deux composants fonctionnent.
- État 2 : Le composant 1 est en panne, le composant 2 fonctionne.
- État 3 : Le composant 1 fonctionne, le composant 2 est en panne.
- État 4 : Les deux composants sont en panne.

Dans cet exemple, le composant qui tombe en panne ne peut pas être déterminé d'une manière certaine. Considérons que le système fonctionne avec un seul composant et que nous ne sachions pas quel composant est tombé en panne. Le système peut se trouver soit dans l'état 2 soit dans l'état 3. Nous définissons ainsi un état flou F_1 par sa fonction d'appartenance $\mu_{F_1}(\cdot)$ tel que :

$$\mu_{F_1}(2) = \widetilde{0.5} \quad \mu_{F_1}(3) = \widetilde{0.5} \quad \mu_{F_1}(1) = 0 \quad \mu_{F_1}(4) = 0 \quad (2.112)$$

De même, si nous ne sommes pas tout à fait sûr que les deux composants fonctionnent, nous pouvons définir un état flou F_2 par sa fonction d'appartenance $\mu_{F_2}(\cdot)$ tel que :

$$\mu_{F_2}(1) = \widetilde{0.8} \quad \mu_{F_2}(2) = \widetilde{0.1} \quad \mu_{F_2}(3) = \widetilde{0.1} \quad \mu_{F_2}(4) = 0 \quad (2.113)$$

Enfin, nous définissons un troisième état flou F_3 par sa fonction d'appartenance $\mu_{F_3}(\cdot)$ tel que :

$$\mu_{F_3}(1) = 0 \quad \mu_{F_3}(2) = \widetilde{0.05} \quad \mu_{F_3}(3) = \widetilde{0.05} \quad \mu_{F_3}(4) = \widetilde{0.9} \quad (2.114)$$

Nous cherchons à déterminer la matrice de transition des états flous $P(F)$ à partir de la matrice de transition des états non-flous P .

La matrice Γ des fonctions d'appartenance des états flous est donnée par :

$$\Gamma^t = \begin{pmatrix} \mu_{F_1}(1) & \mu_{F_1}(2) & \mu_{F_1}(3) & \mu_{F_1}(4) \\ \mu_{F_2}(1) & \mu_{F_2}(2) & \mu_{F_2}(3) & \mu_{F_2}(4) \\ \mu_{F_3}(1) & \mu_{F_3}(2) & \mu_{F_3}(3) & \mu_{F_3}(4) \end{pmatrix} \quad (2.115)$$

D'où :

$$\Gamma^t = \begin{pmatrix} 0 & \widetilde{0.5} & \widetilde{0.5} & 0 \\ \widetilde{0.8} & \widetilde{0.1} & \widetilde{0.1} & 0 \\ 0 & \widetilde{0.05} & \widetilde{0.05} & \widetilde{0.9} \end{pmatrix} \quad (2.116)$$

La matrice des probabilités initiales relative à un état non flou i est donnée par :

$$Q_{0,i} = (\ 0.1 \ 0.2 \ 0.5 \ 0.2 \) \quad (2.117)$$

D'après l'équation (2.111), nous avons :

$$\delta_{1,r} = (m_{F_r})^{-1} = \left(\sum_{i=1}^4 \mu_{F_r}(i) Q_{0,i} \right)^{-1} \quad r = 1, 2, 3. \quad (2.118)$$

D'où :

$$\delta_{1,1} = \left(\sum_{i=1}^4 \mu_{F_1}(i) Q_{0,i} \right)^{-1} = 2.857 \quad (2.119)$$

$$\delta_{1,2} = \left(\sum_{i=1}^4 \mu_{F_2}(i) Q_{0,i} \right)^{-1} = 6.667 \quad (2.120)$$

$$\delta_{1,3} = \left(\sum_{i=1}^4 \mu_{F_3}(i) Q_{0,i} \right)^{-1} = 4.651 \quad (2.121)$$

Finalement :

$$\Delta_1 = \text{diag}(\delta_{1,r}) = \begin{pmatrix} 2.857 & 0 & 0 \\ 0 & 6.667 & 0 \\ 0 & 0 & 4.651 \end{pmatrix} \quad (2.122)$$

Et :

$$\Delta_2 = \text{diag}(Q_{0,i}) = \begin{pmatrix} 0.1 & 0 & 0 & 0 \\ 0 & 0.2 & 0 & 0 \\ 0 & 0 & 0.5 & 0 \\ 0 & 0 & 0 & 0.2 \end{pmatrix} \quad (2.123)$$

La matrice des probabilités de transition des états non-flous est définie par :

$$P = \begin{pmatrix} p_{1,1} & p_{1,2} & p_{1,3} & p_{1,4} \\ p_{2,1} & p_{2,2} & p_{2,3} & p_{2,4} \\ p_{3,1} & p_{3,2} & p_{3,3} & p_{3,4} \\ p_{4,1} & p_{4,2} & p_{4,3} & p_{4,4} \end{pmatrix} \quad (2.124)$$

D'où, en considérant les taux de défaillance et de réparation donnés par l'expert :

$$P = \begin{pmatrix} 0.23 & 0 & 0.6 & 0.17 \\ 0.2 & 0.5 & 0.1 & 0.2 \\ 0.3 & 0.42 & 0.16 & 0.12 \\ 0.9 & 0.04 & 0.01 & 0.05 \end{pmatrix} \quad (2.125)$$

Puisque l'équation d'état du système est donnée par :

$$P(F) = \Delta_1 \Gamma^t \Delta_2 P \Gamma \quad (2.126)$$

Nous obtenons finalement la matrice de transition des états flous $P(F)$:

$$P(F) = \begin{pmatrix} 0.46 & 0.31 & 0.23 \\ 0.13 & 0.57 & 0.3 \\ 0.31 & 0.22 & 0.47 \end{pmatrix} \quad (2.127)$$

Propriété importante

Si P est irréductible, alors $P(F)$ est irréductible.

2.3.2.7 Modèle de Markov avec un système d'inférence flou (MMF)

Pour introduire les modèles de Markov basés sur les systèmes d'inférence flous, nous considérons l'exemple d'un système d'alimentation électrique simple [Tanrioven *et al.*, 2004]. Notre objectif étant d'étudier la fiabilité de ce système.

Nous adoptons les notations suivantes :

- $\lambda_{i,j}$ est le taux de transition entre les états i et j .
- T_t désigne la température à l'instant t .
- L_t désigne le niveau de sollicitation du système à l'instant t .

Nous adoptons aussi l'hypothèse suivante :

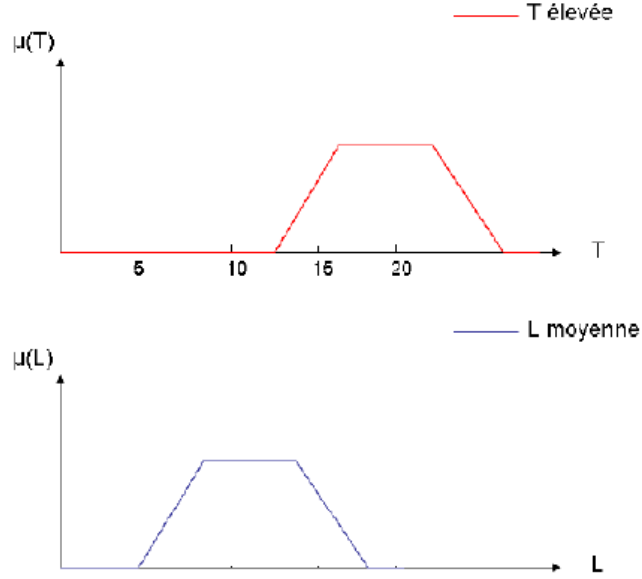


FIG. 2.12 – Fonctions d'appartenance des valeurs T élevée et L moyenne

- Les taux de transition $\lambda_{i,j}$ dépendent de T_t et L_t :

$$\lambda_{i,j} = f(T_t, L_t) \quad (2.128)$$

Le modèle de Markov flou est construit en plusieurs étapes :

- Première étape (Fuzzification) :

Dès que l'expert fournit les valeurs linguistiques des entrées à l'instant t (par exemple à $t = 10 \text{ jours}$: la température T_t est élevée et la sollicitation L_t du système d'alimentation est moyenne), nous procédons à la fuzzification des valeurs de ces entrées en les transformant en nombres flous (cf. figure 2.12).

- Deuxième étape (Inférence) :

Nous appliquons les règles fournies par l'expert sous forme de relations entre les variables linguistiques d'entrées et de sorties. Par exemple, l'une des règles peut être : *Si* (la température est moyenne) *et* (la sollicitation est élevée) *alors* (le taux de défaillance est élevé).

- Troisième étape (Défuzzification) :

L'inférence fournit une fonction d'appartenance floue résultante pour la variable de sortie. Il faut transformer cette information floue en une valeur précise. Nous pouvons utiliser plusieurs méthodes de défuzzification. La plus utilisée est le centroïde (centre de masse) de la fonction d'appartenance résultante de sortie.

- Quatrième étape :

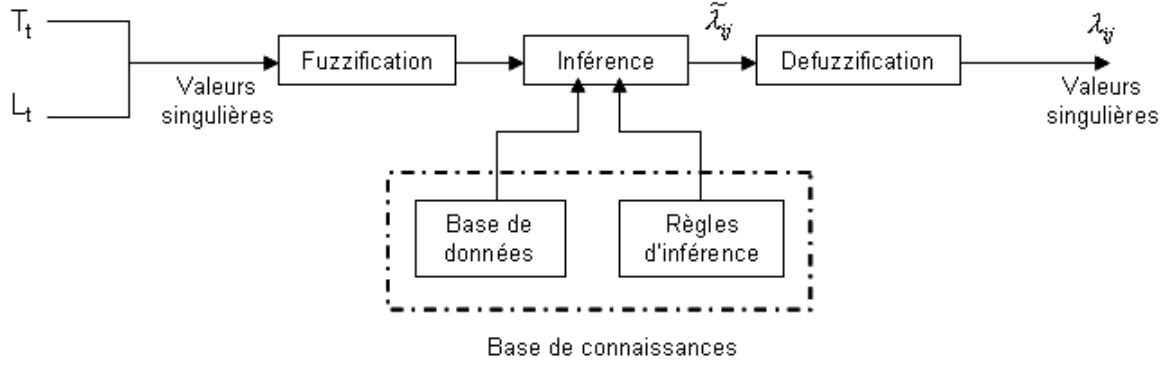


FIG. 2.13 – Modèle flou de Markov

Nous obtenons donc des valeurs singulières pour la matrice des taux de transition A .

$$A = \begin{pmatrix} 0 & \lambda_{1,2} & \dots & \lambda_{1,n} \\ \lambda_{2,1} & 0 & \dots & \lambda_{2,n} \\ \dots & \dots & \dots & \dots \\ \lambda_{n,1} & \lambda_{n,2} & \dots & 0 \end{pmatrix} \quad (2.129)$$

L'équation d'état du système est alors donnée par :

$$\frac{d}{dt}P(t) = AP(t) \quad (2.130)$$

Où $P(t) = [P_1(t), P_2(t), \dots, P_n(t)]^t$ et $P_i(t)$ représente la probabilité que le système soit dans l'état i à l'instant t .

Pour récapituler, nous présentons dans la figure (2.13) les différentes étapes qui permettent de construire le modèle flou de Markov.

2.3.2.8 Conclusion et perspectives

La modélisation des systèmes dynamiques par des chaînes de Markov présente bien des avantages dont notamment la possibilité d'effectuer des traitements plus précis et plus rapides que par simulation de Monte-Carlo en se ramenant à la résolution d'un système d'équations différentielles linéaires du premier ordre. Elle présente cependant deux inconvénients :

- L'emploi exclusif des taux de transition constants (et en particulier la loi exponentielle).
- L'explosion combinatoire des états (2^n états pour un système de n éléments à 2 états).

En outre, au delà d'une dizaine d'états, la construction d'un tel modèle devient très complexe et peut induire des erreurs. Afin de palier ces inconvénients, deux techniques sont couramment utilisées :

- La méthode des états fictifs.
- La génération automatique des matrices de Markov.

Dans le formalisme flou, l'utilisation des chaînes de Markov floues permet l'étude des systèmes Markoviens dont les états ne sont pas définis d'une manière précise. Dans la littérature, nous trouvons deux types de chaînes de Markov floues, le premier type est les chaînes de Markov à états non-flous et probabilités de transition floues, le deuxième est les chaînes de Markov à états flous et non-flous et probabilités de transitions singulières. D'un point de vue mathématique, nous pensons que plusieurs travaux doivent être menés dans le cadre des approches markoviennes floues afin d'établir notamment les propriétés de convergence des matrices de transition floues. D'un point de vue fiabiliste, ces approches markoviennes floues doivent être utilisées plus souvent dans l'évaluation de la sûreté de fonctionnement des systèmes à taux de défaillance imprécis et variant au cours du temps.

2.3.3 Réseaux de Petri

2.3.3.1 Introduction

Les systèmes automatisés complexes sont basés d'une part sur des mécanismes logiques complexes (parallélismes, synchronisation, partage de ressources, priorité, etc.) et d'autre part, sur des actions dont les durées sont le plus souvent connues de manière imparfaite. La conception de tels systèmes nécessite des modèles formels qui permettent de vérifier et évaluer leur propriétés avant d'en faire l'implantation.

Les réseaux de Petri ont été introduits par le mathématicien allemand Carl Adam Petri en 1962 [Petri, 1962]. Ils constituent un outil bien adapté à la modélisation des systèmes complexes grâce à leur représentation graphique et mathématique. La représentation graphique permet de visualiser d'une manière naturelle le parallélisme, la synchronisation, le partage de ressources, la priorité, etc. La représentation mathématique permet d'établir les équations d'état, à partir desquelles il est possible d'apprécier les propriétés du modèle et de les comparer avec le comportement du système modélisé [Khansa, 1997].

Dans cette section, nous présenterons les notions de base des réseaux de Petri. Puis, nous dressons un état de l'art des réseaux de Petri flous.

2.3.3.2 Définitions et concepts de base

Définition 1

Un réseau de Petri (RdP) non marqué est un graphe bipartite constitué de places, de transitions (correspondant aux sommets du graphe) et d'arcs qui relient les transitions aux places et les places aux transitions. Il est représenté par un quadruplet $R = \langle P, T, Pre, Post \rangle$ [Peterson, 1977] où :

- P est un ensemble fini de places,
- T est un ensemble fini de transitions,
- $Pre : P \times T \rightarrow \mathbb{N}$ (ensemble des entiers naturels) est l'application incidence *avant*, reliant des transitions à des places précédentes,
- $Post : P \times T \rightarrow \mathbb{N}$ est l'application incidence *arrière*, reliant des transitions à des places suivantes.

Définition 2

Soit R un réseau de Petri. Le marquage du réseau de Petri R est une fonction M de l'ensemble des places P dans $\mathbb{N}$. L'ensemble des marques d'un réseau de Petri R est donné par le vecteur de marquage $M(R)$. Un réseau marqué est le couple $\langle R, M_0 \rangle$ avec M_0 le marquage initial.

Définition 3

Soit R un réseau de Petri, M est un marquage et t une transition. La transition t est sensibilisée (validée) pour le marquage M , noté $M \geq \text{Pre}(\cdot, t)$ ou $M(t > \text{ ou } {}^tM \rightarrow)$, si et seulement si :

$$\forall p \in P, \quad M(p) \geq \text{Pre}(p, t) \quad (2.131)$$

Définition 4

Soit R un réseau de Petri, M est un marquage et t une transition sensibilisée pour M . Le tir (franchissement) de t donne le marquage M' tel que :

$$\forall p \in P, \quad M'(p) = M(p) - \text{Pre}(p, t) + \text{Post}(p, t) \quad (2.132)$$

Définition 5

Soit R un réseau de Petri, M est un marquage et t_1, t_2 deux transitions. Les transitions t_1 et t_2 sont en conflit structurel, si et seulement si, il existe une place $p \in P$ d'entrée telle que :

$$\text{Pre}(p, t_1) \cdot \text{Pre}(p, t_2) \neq 0 \quad (2.133)$$

où $(\cdot)$ dénote la multiplication d'entiers naturels.

Les transitions t_1 et t_2 sont en conflit effectif pour un marquage M , si elles sont en conflit structurel et sont sensibilisées par M .

2.3.3.3 Réseaux de Petri temporels

Les réseaux de Petri temporels ou t -temporels ont été introduits par Merlin [Merlin, 1974]. Ils se basent sur le principe de l'association d'une durée de sensibilisation bornée $\theta_s(t)$ à chaque transition t . Ceci permet de laisser disponibles les jetons dans la place d'entrée de la transition t jusqu'à son tir. D'autres transitions en conflit avec t peuvent ainsi consommer ces jetons lors de leur franchissement. Dans ces réseaux, un intervalle de temps $[a, b]$ est associé à chaque transition. Si une transition est restée sensibilisée pendant a unités de temps, alors elle peut être tirée. Aussi, si une transition est restée sensibilisée pendant b unités de temps, elle doit être tirée. L'intervalle $[a, b]$ associe donc une incertitude à la date de tir de la transition qui va être tirée entre a et b unités de temps relativement à sa date de sensibilisation [Valette and Kunzle, 1994; Boyer, 2001]. Toutes les valeurs de $\theta_s(t)$ telles que $a \leq \theta_s(t) \leq b$ sont des durées de sensibilisation possibles pour t . Un événement (franchissement de la transition t par exemple) est contraint d'arriver à au moins a unités de temps et au plus à b unités de temps après la sensibilisation de t .

Définition 6

Un réseau de Petri temporel est une paire $\langle R, I \rangle$ où :

- R est un réseau de Petri $\langle P, T, Pre, Post \rangle$ auquel est associé un marquage initial M_0 ,
- I est la fonction d'intervalle initial qui associe à chaque transition un intervalle fermé rationnel $I(t) = [a, b]$ qui décrit une durée de sensibilisation de la transition t .

Grâce à leur structure, les réseaux de Petri temporels ont l'avantage de représenter naturellement les chiens de garde. Dans la figure 2.14, le marquage de la place "Condition" signifie l'occurrence d'un événement. La transition t_2 est, de ce fait, tirée immédiatement. Si au bout de θ unités de temps l'événement ne s'est pas produit, alors la transition t_1 est tirée et une alarme est déclenchée.

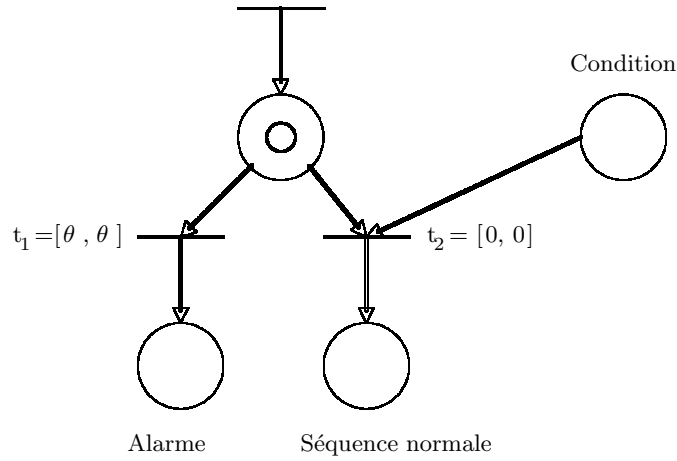


FIG. 2.14 – Représentation d'un chien de garde

2.3.3.4 Réseaux de Petri p-temporels

Les réseaux de Petri p-temporels ont été introduits par Khansa [Khansa, 1997]. Ils représentent un formalisme puissant et reconnu pour la modélisation de l'obligation de respect des temps de séjour (synchronisation sous obligation). Dans ce type de réseaux, un intervalle temporel est associé aux places. Un jeton dans une place p à laquelle est associé un intervalle $[a, b]$ doit être consommé au plus tôt a unités de temps et au plus tard b unités temps après le marquage de la place p . Si un jeton n'est pas consommé avant b unités de temps, alors il devient un jeton "mort".

Définition 7

Un réseau de Petri p-temporel est une paire $\langle R, I \rangle$ où :

- R est un réseau de Petri $\langle P, T, Pre, Post \rangle$ auquel est associé un marquage initial M_0 ,
- I est la fonction d'intervalle initial qui associe à chaque place un intervalle fermé rationnel $I(p) = [a, b]$ qui décrit une durée de marquage de la place p .

En général, ce type de réseau représente des durées séparant des occurrences d'événements. Ces durées peuvent, par exemple, être des durées opératoires séparant un événement début et un événement fin d'une même opération Op sur une pièce. L'exécution de Op est modélisée par un jeton dans une place. Dans ce cas, si le jeton correspondant à l'opération Op en cours devient un jeton "mort", alors cela permet de déduire que la pièce devant subir cette opération est restée trop longtemps dans cet état.

2.3.3.5 Réseaux de Petri stochastiques

Dans un réseau de Petri stochastique un temps aléatoire est associé au franchissement des transitions [Marson *et al.*, 1989]. A tout réseau de Petri stochastique, nous pouvons associer une chaîne de Markov correspondante au graphe de marquage du réseau. La méthode d'analyse employée consiste à construire le graphe de marquage accessible du réseau de Petri stochastique. Le comportement aléatoire du réseau de Petri stochastique est alors identique à celui de la chaîne de Markov ainsi déterminée. Dans la suite, nous allons détailler les modèles usuels d'évolution de ces réseaux.

Modèles d'évolution

Un réseau de Petri non temporisé représente un système déterministe et l'extension à un réseau stochastique permet d'ajouter un comportement dynamique aléatoire au système. L'aléatoire porte sur les durées d'exécution de tâches modélisées par les transitions, chaque tâche une fois achevée faisant changer d'état le système dynamique. La façon la plus naturelle de définir des durées aléatoires consiste à associer à chaque transition t une variable aléatoire de distribution $F_t(M)$ dépendante du marquage du réseau. Cette dépendance du marquage se fera à travers la moyenne $\lambda_t^{-1}(M)$ de la loi et non pas sur sa forme. Elle traduira une influence du marquage sur la vitesse d'exécution de la tâche et non sur sa distribution. Il suffira donc de définir un ensemble $\hat{F}$ de fonctions de répartition :

$$\hat{F} = \{\hat{F}_t, t \in T \mid \hat{F}_t : [0, \infty[\rightarrow [0, 1]\} \quad (2.134)$$

et un ensemble λ de taux de tir dépendant du marquage :

$$\lambda = \{\lambda_t, t \in T \mid \lambda_t : \mathbb{N}^m \rightarrow [0, \infty[\} \quad (2.135)$$

Où m est le nombre de transitions du réseau de Petri et θ la durée de sensibilisation de la transition.

Par conséquent la fonction de répartition associée à la transition t dans le marquage M s'exprimera par :

$$F_t(M, \theta) = \hat{F}_t(\lambda_t(M). \theta) \quad (2.136)$$

Où θ est la durée de sensibilisation d'une transition.

Il existe deux modèles pour le choix d'un état à partir d'un autre [Marson *et al.*, 1989] :

- Modèle de concurrence : un modèle de concurrence d'un réseau de Petri stochastique est basé sur la compétition entre tâches à durée aléatoire. Ce modèle considère que les états du système sont dénombrables. Quand le système se trouve dans un état, l'état suivant

sera choisi par l'observation de variables aléatoires. Chacune de ses variables est associée à une transition sensibilisée par le marquage M_i conduisant après son tir à un marquage M_j accessible. Soit W_t la variable aléatoire associée à la transition t . Si la plus petite réalisation de ses variables correspond à la variable W_s associée à la transition s , alors s sera tirée et l'état j atteint à partir de l'état i .

- Modèle de présélection : le modèle concurrent présente des lacunes, en particulier s'il existe deux activités de temps d'exécution identique. Selon ce modèle les deux transitions devraient être tirées en même temps. Le problème se résume au choix d'une transition parmi les transitions validées dans le marquage i considéré. Le modèle de présélection consiste dans un premier temps à choisir, pour un marquage i donné, une transition t en fonction des poids associés aux transitions validées par le marquage i . La transition retenue est alors franchie après une durée aléatoire dépendant de la loi associée. La durée de séjour dans le marquage M_i ne dépend plus des délais de tir associés aux autres transitions. La mise en place d'un tel modèle s'avère difficile car elle nécessite de connaître explicitement toutes les transitions validées par un marquage donné. Elle est cependant utilisée pour régler un conflit effectif entre des transitions immédiates, qui dans un modèle concurrent sont toujours franchies en premier.

Règles de mémorisation

Les deux modèles présentés exigent le calcul des lois des variables W_t conditionnellement à l'évolution du réseau de Petri jusqu'au temps d'atteinte du marquage M_i . Cette dépendance conditionnelle dépendra des règles de mémorisation du temps total de sensibilisation de chaque transition :

- "Age memory" : le comportement le plus général d'un réseau de Petri stochastique correspond à la règle d'age memory, qui consiste à mémoriser pour chaque transition la durée totale pendant laquelle elle a été sensibilisée, ceci jusqu'à son tir.
- "Enabling memory" : pour cette règle, la durée totale de sensibilisation mémorisée est remise à zéro à chaque désensibilisation de la transition.
- "Resampling" : pour la règle de resampling, il n'y a pas de mémorisation de la durée de sensibilisation. Il y a une remise à zéro à chaque changement d'état du réseau de Petri. Dans ce cas la loi de tir de chaque transition validée dans un certain marquage ne dépendra de l'évolution du réseau de Petri stochastique qu'à travers ce marquage. Les instants de changement d'état seront des points régénératifs du système.

En conclusion, nous pouvons dire que les réseaux de Petri stochastiques constituent un puissant outil de modélisation des systèmes dynamiques à états discrets. Ils sont particulièrement adaptés pour représenter les phénomènes de concurrence, de synchronisation et de communication. Actuellement plusieurs chercheurs comme Guy Juanole du LAAS participent à une réflexion sur le lien des réseaux de Petri stochastiques avec les réseaux de Petri temporels flous que nous allons décrire dans la suite de cette section.

2.3.3.6 Réseaux de Petri flous (RdPF)

Les différents types de logiques (classique, linéaire et floue) utilisées dans la description des systèmes, génèrent deux catégories de modèles RdPF [Racoceano, 2006] :

- La première catégorie est représentée par les réseaux qui modélisent des raisonnements flous exprimés par des règles floues décrivant le comportement générique d'un système à événements discrets. Les règles floues [Valette and Atabakhche, 1987; Valette and Kunzle, 1994], sont mises en correspondance avec les transitions. Le modèle global, correspondant à l'ensemble des règles, modélise un système expert flou. Les transitions qui modélisent l'ensemble des règles simultanément vraies seront franchies en parallèle et la conclusion sera le résultat obtenu par défuzzification [Hanna, 1994; Gomes and Steiger-Garçao, 1995]. Le mécanisme spécifique à un contrôleur flou et identifiable à un système expert flou présente alors une typologie simplifiée.
- L'autre catégorie de RdPF modélise un système physique dynamique. Nous obtenons des classes d'applications qui expriment l'incertitude des connaissances au niveau des états des ressources du système [Cardoso *et al.*, 1990; Valette *et al.*, 1989]. Ce type de RdPF modélise un système physique par application de la logique linéaire aux transformations de l'état des ressources, tout en conservant l'incertitude des connaissances concernant leur présence ou leur disponibilité. Les composantes du RdPF sont adaptées pour la représentation des connaissances incertaines par des variables floues associées aux places [Cao, 1993; Bugarin and Barro, 1994] et/ou aux transitions et/ou aux jetons [Chen *et al.*, 1990; Pedrycz and Gomide, 1994; Lipp and Gunter, 1994; Scarpelli and Gomide, 1994; Gomes and Steiger-Garçao, 1995] et/ou au marquage [Cardoso *et al.*, 1990; Cardoso, 1990; Valette *et al.*, 1989] et/ou aux arcs [Camargo, 1998; Lipp and Gunter, 1994].

Dans toutes ces approches, la présence du jeton est binaire, même si celui-ci se charge de valeurs floues au moment de son apparition dans la place.

RdPF pour la modélisation des raisonnements flous

Dans cette catégorie, le raisonnement est basé sur une représentation de la connaissance où des règles floues sont associées aux transitions et le modèle représente un système expert ou un système de contrôle. L'ensemble des transitions est en conformité avec un ensemble de règles floues R . Chaque transition est associée à une règle. Un ensemble D de propositions logiques définit la base de règles R [Racoceano, 2006].

Par exemple, dans la figure 2.15, nous présentons la modélisation de l'expression logique :

$$d_i \wedge d_j \rightarrow d_k \quad (2.137)$$

Cette expression est équivalente à une disjonction de deux expressions logiques élémentaires :

$$d_i \wedge d_j \rightarrow d_k \iff (d_i \rightarrow d_k) \vee (d_j \rightarrow d_k) \quad (2.138)$$

La transition qui modélise l'implication $d_i \rightarrow d_k$ peut être franchie dans l'intervalle :

$$\theta_i = [t_{0i}, t_{max-i}] \quad (2.139)$$

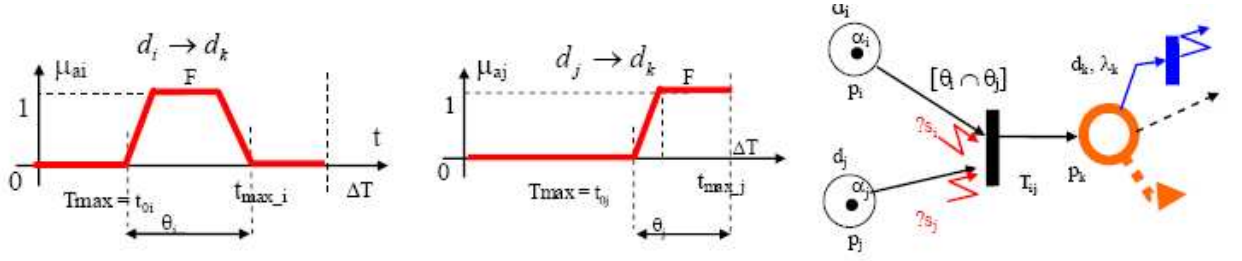


FIG. 2.15 – Modélisation de l'implication logique par une transition temporelle

Alors que la transition qui modélise l'implication $d_j \rightarrow d_k$ peut être franchie dans l'intervalle :

$$\theta_j = [t_{0j}, t_{max-j}] \quad (2.140)$$

Nous pouvons ainsi conclure que ces réseaux de Petri flous sont analogues à des systèmes experts flous obtenus en remplaçant simplement les transitions par des règles du type (IF ... THEN).

RdPF pour la modélisation des systèmes physiques dynamiques

De nombreux modèles flous basés sur les RdP ont été définis [Cardoso, 1999], le RdP temporel flou étant le modèle le plus représentatif. Les réseaux de Petri temporels flous associent aux transitions des réseaux de Petri des attributs temporels sous la forme d'intervalles de temps avec des distributions de possibilité. Ils permettent de prendre en compte une incertitude sur les données temporelles qui peut provenir des aléas auxquels le système est soumis, ou des besoins de flexibilité et d'autonomie dans son fonctionnement.

Le modèle développé dans le cadre de la surveillance des systèmes manufacturiers [Cardoso, 1990; Cardoso, 1999] est un modèle du système basé sur un réseau de Petri (à objets) temporel flou. Les places dénotent un état partiel (localisation d'un objet) du système et les transitions un changement possible d'état. Une séquence de tir de transitions représente un comportement possible. Un marquage précis représente le comportement normal du système tandis qu'un marquage flou exprime le fait que l'existence d'un objet (modélisé par un jeton) est connue mais sa localisation est imprécise ou floue (le jeton peut se trouver dans plus qu'une place) en raison, par exemple, d'une information imprécise ou incomplète. La notion de pseudo tir ou tir incertain est introduite pour une interprétation incertaine du franchissement de transition.

La connaissance de l'état du système est représentée par une distribution de possibilité π sur les places en fonction du temps. Cette distribution, dont un exemple est donné dans la figure 2.16, modélise la trajectoire d'évolution d'une pièce à travers des machines m_i représentées par les places p_i . A chaque instant t_1 , un degré de possibilité est associé à chaque place $\pi_{pt1}(p_i)$ représentant la possibilité pour la pièce p_i de se trouver à cet instant dans les machines m_i (cf. figure 2.17).

Cardoso [Cardoso, 1999] distingue deux méthodes de mise à jour du marquage (son évolution) en fonction de la prise en compte de l'information associée aux événements de tir de transition :

- Dans la première méthode, le temps est implicite et une fonction intitulée fonction d'autorisation $\eta \rightarrow \text{fausse, incertaine, vraie}$ est associée aux transitions. Cette fonction peut évaluer, par exemple, la condition de présence d'un signal et de son occurrence avant ou

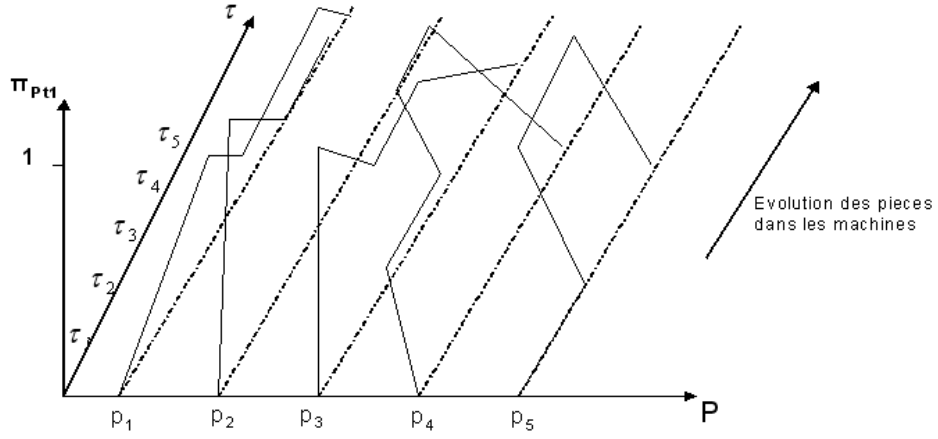


FIG. 2.16 – Distribution de possibilité sur les places p_i en fonction du temps : (P, p_0, T)

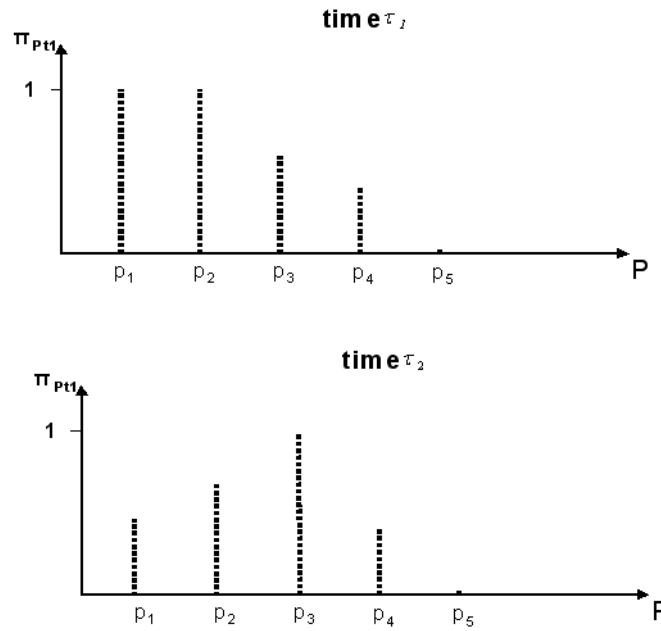


FIG. 2.17 – Distribution de possibilité sur les places p_i (π, P) aux instants $t = t_1$ et $t = t_2$

après une date donnée. La révision de la connaissance sur le réseau (évolution du marquage) est faite par un mécanisme de tir incertain [Cardoso, 1999].

- Dans la deuxième méthode, le temps est pris en compte explicitement et des dates possibles d'occurrence d'un événement $\pi_a(\tau)$ sont affectées aux transitions et délimitées par un ensemble flou A [Prade and Dubois, 1983]. Les notions qui en résultent comme l'intervalle flou, la longueur d'intervalle, le temps de transition au plus tôt et au plus tard sont définies de façon à mieux raisonner sur l'évolution temporelle du système à travers un mécanisme de propagation des contraintes temporelles floues. Le marquage flou est déduit de cette évaluation temporelle. En fonction de la date d'occurrence d'un événement, il est possible d'inférer sur le degré de possibilité que le système se trouve dans certains états partiels. Un degré de possibilité inférieur à 1 implique également une révision des dates possibles d'occurrence des événements subséquents, caractérisant un mécanisme de propagation de contrainte temporelle floue sur le RdPTF [Cardoso, 1999].

La prise en compte explicite du temps dans un contexte d'incertitude est également traitée par [Murata *et al.*, 1999] qui propose un modèle intitulé Fuzzy-Timing High-Level Petri Net (FTHLPN). En partant des concepts définis par [Prade and Dubois, 1983], les définitions de fuzzy timestamp (instant possible d'arrivée d'un jeton dans une place), fuzzy enabling time (liés aux instants possibles d'arrivée des jetons nécessaires pour le tir de transitions), fuzzy occurrence time (date floue d'occurrence d'un événement) et fuzzy delay (associé aux arcs) sont proposées. [Ribaric *et al.*, 1998] définissent dans une structure RdPTF similaire à celle de Cardoso, un jeton temporel flou qui porte l'information de distribution de possibilité sur sa date d'arrivée (π_a) et de fin de séjour (π_b) dans une place ainsi qu'un facteur de certitude. Ces informations sont mises à jour à chaque tir de transition en fonction de l'évaluation temporelle d'occurrence de l'événement.

2.3.3.7 Observations et limitations

Les réseaux de Petri sont un modèle formel qui permet la spécification, la validation et l'implémentation de systèmes dynamiques. Les réseaux de Petri sont utilisés dans divers domaines d'application. Dans ces types d'applications, deux aspects sont très importants : la dimension temporelle de l'évolution du système et l'incertitude de l'information dont chaque agent dispose.

Plusieurs extensions ont été définies, telles que les réseaux de Petri temporels qui permettent de représenter implicitement des contraintes temporelles, les réseaux de Petri stochastiques qui utilisent une fonction densité de probabilité sur les intervalles de temps et les réseaux de Petri flous qui manipulent les ensembles flous quand les informations traitées sont connues avec incertitude. Les réseaux de Petri temporels flous associent aux transitions des réseaux de Petri des attributs temporels sous la forme d'intervalles de temps avec des distributions de possibilités. Ils permettent de prendre en compte une incertitude sur les données temporelles qui peut provenir des aléas auxquels le système est soumis, ou des besoins de flexibilité et d'autonomie dans son fonctionnement.

Notons qu'il existe plusieurs problèmes à traiter par cette approche :

- Il faut déterminer et prouver les propriétés du graphe des classes d'états Flous (GCEF), en particulier quand il existe des boucles dans le système, et le cas de séquences infinies ;
- Il faut considérer le vrai parallélisme dans la génération du GCEF : si deux transitions

- sont franchissables avec dates de tir (floues) a et b , il faut calculer la possibilité $\Pi(a = b)$ (et nécessité $N(a = b)$) pour qu'elles puissent être franchies aux mêmes instants ;
- Etudier la propagation de sous-ensembles flous normalisés ainsi que de sous-ensembles flous ayant des formes non standard (non trapézoïdales).

2.4 Conclusion

Pour évaluer la sûreté de fonctionnement des systèmes complexes à défaillances rares comme les systèmes relatifs aux applications de sécurité, il est impératif de prendre en compte les connaissances imparfaites dont nous disposons concernant les données de fiabilité des composants de ces systèmes.

Pour ce faire, nous nous sommes focalisés sur les méthodes de sûreté de fonctionnement définies dans le cadre de la théorie de fiabilité PROBIST floue :

- Hypothèse probabiliste : le fonctionnement du système est complètement caractérisé par des mesures de probabilité.
- Hypothèse d'états binaires : le système ne peut avoir que deux états de fonctionnement ; l'état de défaillance et l'état de fonctionnement normal.

Nous avons d'abord introduit les notions de base concernant les théories de représentation et d'analyse des informations imparfaites :

- Théorie des probabilités.
- Théorie des sous-ensembles flous.
- Théorie des possibilités.
- Théorie des fonctions de croyance.

Puis, nous avons dressé un état de l'art de l'emploi de ces théories, notamment celle des sous-ensembles flous dans les principales méthodes de sûreté de fonctionnement en présence d'incertitudes :

- Arbres de défaillance.
- Approches markoviennes.
- Réseaux de Petri.

Dans la suite de nos travaux, nous avons choisi d'utiliser la méthode des arbres de défaillance flous pour l'évaluation de la sûreté de fonctionnement des SIS en se situant dans le cadre de la fiabilité PROBIST floue.

3

Evaluation des SIL par une approche probabiliste floue et études de sensibilité

Dans ce chapitre, nous proposons une approche probabiliste floue pour l'évaluation de la sûreté de fonctionnement des systèmes en présence de données de fiabilité imparfaites. Cette approche est utilisée pour la détermination des PFD_{avg} des SIS et l'évaluation des SIL. Ensuite, nous nous fixons comme objectif la réduction de l'incertitude entachant les SIL obtenus grâce à des études de sensibilité. Une application support est proposée afin de montrer l'intérêt de notre approche pour l'évaluation des SIL en présence d'informations imparfaites.

Ce chapitre est composé de quatre parties. La première partie propose l'approche probabiliste floue qui est basée sur l'utilisation des arbres de défaillance pour l'évaluation du PFD_{avg} des SIS. Nous parlerons indifféremment du SIL d'un SIS ou du SIL d'une SIF car nous supposons que chaque SIS ne réalise qu'une seule SIF. En outre, nous parlerons indifféremment du PFD_{avg} et de l'indisponibilité moyenne du SIS sur une période donnée. La deuxième partie permet de caractériser l'incertitude des SIL obtenus grâce à l'utilisation de mesures de possibilité et de nécessité. La troisième partie introduit d'abord les études de sensibilité basées sur les approches probabilistes, puis présente les facteurs d'importance flous que nous avons proposé. Ces facteurs d'importance vont apporter une aide à la réduction de l'incertitude des SIL en s'intéressant principalement aux composants présentant les facteurs les plus élevés. La dernière partie présente une application qui illustre l'utilisation de la méthodologie proposée et met en évidence l'intérêt des résultats obtenus avec cette approche.

3.1 Introduction

Les normes de sécurité IEC 61508 [IEC61508, 1998] et IEC 61511 [IEC61511, 2000] définissent les SIL auxquels doivent satisfaire les SIS (cf. section 1.3). Notre objectif est de déterminer les PFD_{avg} des SIS en présence de données de fiabilité imparfaites. Les méthodes quantitatives recommandées par ces normes de sécurité pour le calcul du PFD_{avg} des SIS à partir des paramètres de fiabilité de leurs composants sont :

- Les équations simplifiées ;
- Les arbres de défaillance ;
- Les approches markoviennes.

Comme l'ont souligné plusieurs chercheurs [Peguet and Boyer, 2006; Innal *et al.*, 2005; Innal *et al.*, 2006], il est préférable d'utiliser des méthodes de sûreté de fonctionnement classiques telles que les arbres de défaillance ou les approches markoviennes pour évaluer la disponibilité moyenne des SIS, plutôt que d'utiliser les équations simplifiées données par les normes IEC 61508 [IEC61508, 1998] et IEC 61511 [IEC61511, 2000]. En effet, ces équations simplifiées ne sont valables que sous plusieurs hypothèses et pour certains types d'architecture k parmi n . Par ailleurs, le rapport technique ISA-TR84.00.02-2002 [ISA-TR84.00.02-2002, 2002] a recommandé l'utilisation de la méthode des arbres de défaillance dans les procédés nécessitant des SIF de SIL 2 ou de SIL 3. Pour toutes ces raisons, nous avons donc choisi la méthode des arbres de défaillance pour l'évaluation des SIL.

L'analyse des arbres de défaillance conventionnelle est basée sur l'approche probabiliste. La probabilité d'occurrence de l'événement sommet (probabilité de défaillance du système complet) est calculée à partir des probabilités d'occurrence des événements élémentaires (probabilités de défaillance des composants du système) en utilisant les opérations arithmétiques classiques. Les probabilités de défaillance des composants sont calculées à partir de leurs taux de défaillance. Cependant, dans certains cas, il est difficile d'obtenir une quantité suffisante de données pour évaluer les taux de défaillance des composants du système d'une manière précise. Ce problème se pose, en particulier, pour les SIS faiblement sollicités. En effet, ceux-ci présentent des défaillances très rares ne permettant pas de valider statistiquement les calculs prévisionnels à partir des paramètres de fiabilité de leurs composants (cf. section 1.4). Il est donc intéressant de proposer d'autres approches pour analyser les arbres de défaillance en présence de données de fiabilité imprécises.

Dans ce chapitre, nous abordons le problème de la prise en compte des imprécisions relatives aux données de fiabilité des composants du SIS pour l'évaluation des SIL et la quantification des incertitudes des résultats obtenus. Nous proposons d'abord la modélisation des imprécisions des taux de défaillance des composants par des nombres flous triangulaires. L'approche peut être étendue aisément à tout type de forme de nombres flous (trapézoïdal, rectangulaire, etc.). Puis, nous utilisons des arbres de défaillance flous pour évaluer le PFD_{avg} et le SIL qui doit être satisfait par le SIS. Les arbres de défaillance flous sont traités par la méthode proposée par Kaufman et Gupta [Kaufman and Gupta, 1991] (cf. section 2.2.2.4). Les fonctions d'appartenance représentant les probabilités floues des événements de base sont découpées en α -coupes. Chaque α -coupe d'une probabilité floue d'un événement de base est utilisée pour calculer l' α -coupe correspondant à la probabilité de l'événement sommet qui représente la probabilité floue de défaillance du SIS lors de sa sollicitation. La fonction d'appartenance de la probabilité floue de l'événement som-

met est construite à partir des α -coupes calculées [Sallak *et al.*, 2006b]. L'introduction de deux nouveaux facteurs d'importance flous inspirés des facteurs d'importance classiques va permettre la mise en évidence des composants contribuant le plus à la défaillance du SIS et ceux contribuant le plus à l'incertitude entachant le SIL. Le calcul de ces facteurs va apporter une aide à la réduction de l'incertitude des SIL [Sallak *et al.*, 2006c]. Ensuite, nous présentons un exemple applicatif défini dans la littérature [ISA-TR84.00.02-2002, 2002] qui illustre l'approche proposée. Enfin, nous traçons le bilan de notre approche probabiliste floue et des facteurs d'importance flous introduits.

3.2 Evaluation des SIL par une approche probabiliste floue

3.2.1 Modélisation des taux de défaillance des composants des SIS

La première question que nous devons nous poser est : comment déterminer les fonctions d'appartenance des taux de défaillance des composants des SIS ? La première étape consiste à modéliser ces taux de défaillance en choisissant les fonctions d'appartenance adéquates.

L'imprécis et l'incertain peuvent être considérés comme deux points de vue sur une même réalité qu'est l'imperfection de l'information concernant les taux de défaillance des composants. Dans ce contexte, nous pouvons différencier clairement les concepts d'imprécis et d'incertain : l'imprécis concerne le contenu de l'information tandis que l'incertain est relatif à sa vérité, entendue au sens de sa conformité à une réalité [Bouchon-Meunier, 1995].

Nos travaux de thèse proposent seulement un traitement des taux de défaillance imprécis en utilisant des nombres flous sous la forme d'ensembles de valeurs avec des informations de caractère vague. Dans la langue française, il y a d'autres qualificatifs qui renvoient à l'imprécis, tels que "vague", "flou", "général" et "ambigu". L'ambigu est une forme d'imprécision liée au langage. Une information est ambiguë dans la mesure où elle renvoie à plusieurs contextes ou référentiels possibles. Ce type d'imprécision n'est pas celui qui sera considéré dans ce mémoire : nous supposons connu le référentiel associé à l'élément d'information. Le "général" est une forme d'imprécision liée au processus d'abstraction ; une information est générale si elle désigne un ensemble d'objets dont elle souligne une propriété commune. Nous nous intéressons au caractère vague ou flou d'une information qui réside dans l'absence de contours bien délimités de l'ensemble des valeurs affectées aux objets qu'elle concerne.

Considérons par exemple, un composant dont le taux de défaillance est obtenu à partir d'une base de données de fiabilité [IEEE, 1984; CCPS, 1991; CCPS, 2002]. Comme nous l'avons souligné précédemment (cf. section 1.4), ce taux de défaillance est donné sous forme d'une valeur moyenne m est d'un facteur d'erreur e . Nous pouvons dire que le taux de défaillance du composant est d'environ m défaillances par an. Dans ce cas, l'information sur le taux de défaillance du composant est vague ou floue. L'imperfection de cette information est considérée comme une imprécision qui sera, par exemple, modélisée par un nombre flou de valeur modale m et de support $2e$.

Nous proposons donc de modéliser les taux de défaillance imprécis par des nombres flous triangulaires (cf. figure 3.1). Le paramètre α utilisé dans les figures désigne le degré d'appartenance de chaque valeur (cf. figure 3.2). Cependant, l'approche proposée peut être appliquée à

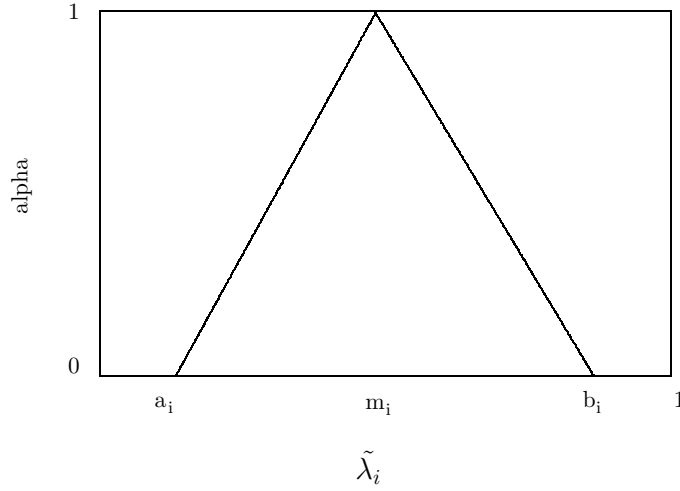


FIG. 3.1 – Modélisation du taux de défaillance imprécis λ_i par un nombre flou triangulaire

tout type de forme. Dans sa forme la plus générale, la fonction d'appartenance d'un nombre flou triangulaire λ_i est donnée par :

$$\begin{aligned} \mu(\lambda_i) &= 0, & \lambda_i < a_i \\ \mu(\lambda_i) &= \frac{\lambda_i - a_i}{m_i - a_i}, & a_i \leq \lambda_i \leq m_i \\ \mu(\lambda_i) &= \frac{b_i - \lambda_i}{b_i - m_i}, & m_i < \lambda_i < b_i \\ \mu(\lambda_i) &= 0, & \lambda_i > b_i \end{aligned}$$

Ces nombres flous triangulaires sont caractérisés par les 3 paramètres m_i , a_i et b_i . m_i représente la valeur modale de la fonction d'appartenance ($\mu(m_i) = 1$). Elle représente la valeur la plus probable du taux de défaillance λ_i . a_i est la limite à gauche de m_i , qu'on appelle aussi valeur basse. b_i est la limite à droite de m_i , qu'on appelle aussi valeur haute.

Par ailleurs, l'avantage d'utiliser une forme triangulaire pour modéliser les taux de défaillance réside dans le fait qu'elle peut être biaisée vers la droite ou vers la gauche par rapport à sa valeur la plus probable. Un autre avantage inhérent aux nombres flous triangulaires est la facilité d'utilisation grâce à la simplification des opérations arithmétiques floues [Kaufman and Gupta, 1991].

3.2.2 Evaluation du PFD_{avg} des SIS

Dans cette section, nous décrivons d'abord comment évaluer les probabilités de défaillance des composants à partir de leurs taux de défaillance imprécis. Dans nos travaux, nous supposons que les taux de défaillance des composants sont constants. Puisque les SIS sont des systèmes à défaillances rares, en utilisant l'approximation des événements rares [Collet, 1996], nous obtenons donc la probabilité de défaillance d'un composant à partir de son taux de défaillance λ à un instant t :

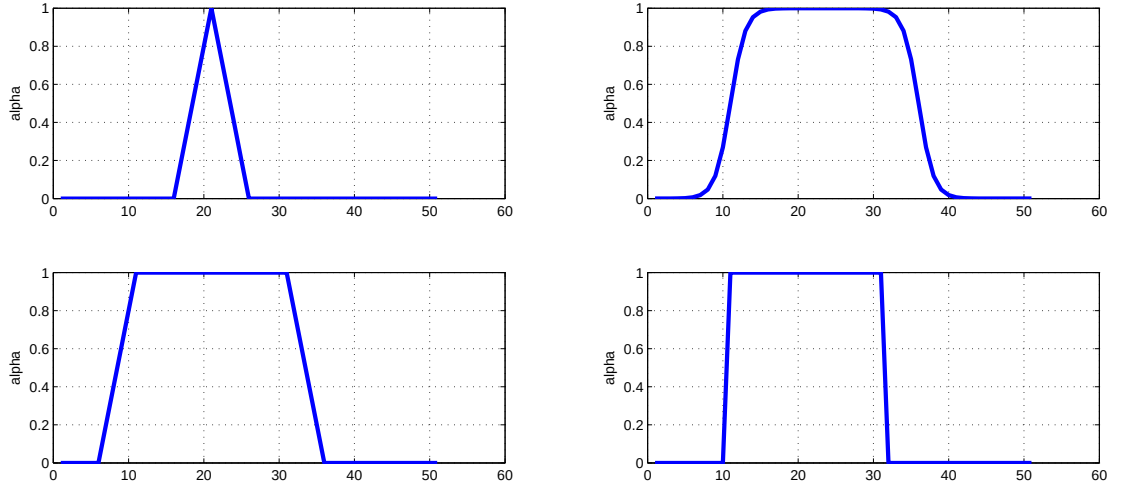


FIG. 3.2 – Différentes formes possibles des fonctions d'appartenance

$$P(t) = 1 - \exp(-\lambda.t) \simeq \lambda.t \quad (3.1)$$

Puis, nous construisons l'arbre de défaillance dont le sommet représente l'événement associé à la défaillance dangereuse du SIS. Pour évaluer le PFD_{avg} du SIS, nous commençons par évaluer la probabilité d'occurrence de l'événement sommet de l'arbre de défaillance à un instant t (qui représente l'indisponibilité instantanée du SIS) à partir des probabilités d'occurrences des événements élémentaires (probabilités de défaillance des composants du SIS). Ensuite, pour obtenir le PFD_{avg} , nous calculons la valeur moyenne des indisponibilités que nous avons calculé sur une période donnée.

Si nous considérons l'arbre de défaillance simple présenté dans la figure 3.4, en supposant que les événements X_i sont indépendants et ont de très faibles probabilités d'occurrences (approximation des événements rares [Collet, 1996]), alors la probabilité d'occurrence de l'événement sommet à un instant t est obtenue par :

$$\tilde{P}_T(y) = \sup_{\{y=p_{A1}+p_{A2}\}} \min\{\tilde{P}_{A1}(p_{A1}), \tilde{P}_{A2}(p_{A2})\} \quad (3.2)$$

Où :

$$\tilde{P}_{A1}(y) = \sup_{\{y=p_1p_2\}} \min\{\tilde{P}_{X1}(p_1), \tilde{P}_{X2}(p_2)\} \quad (3.3)$$

$$\tilde{P}_{A2}(y) = \sup_{\{y=p_3p_4\}} \min\{\tilde{P}_{X3}(p_3), \tilde{P}_{X4}(p_4)\}$$

$\tilde{P}_T$ représente la probabilité floue de défaillance dangereuse du système complet (probabilité d'occurrence de l'événement sommet) et $\tilde{P}_{X_i}$ représente la probabilité de défaillance du composant i (cf. figure 3.3).

Comme les opérations arithmétiques utilisées pour manipuler les probabilités floues requièrent beaucoup de ressources (cf. section 2.2.2.4), Kaufman et Gupta [Kaufman and Gupta, 1991] ont montré que ces efforts de calculs sont largement simplifiés par la décomposition des fonctions

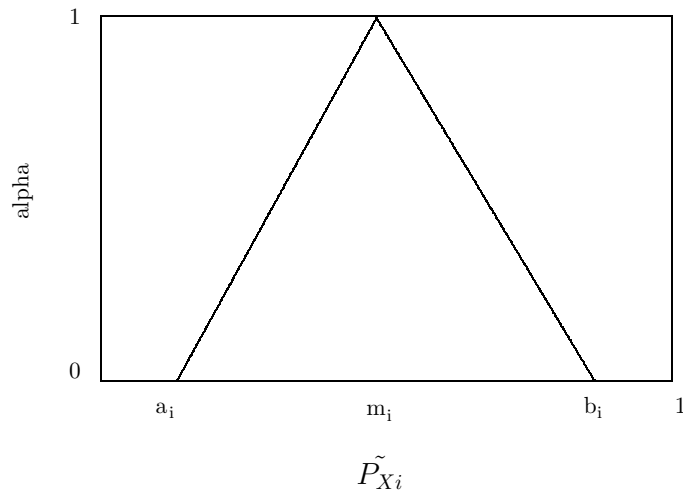


FIG. 3.3 – Probabilité floue de défaillance d'un composant i

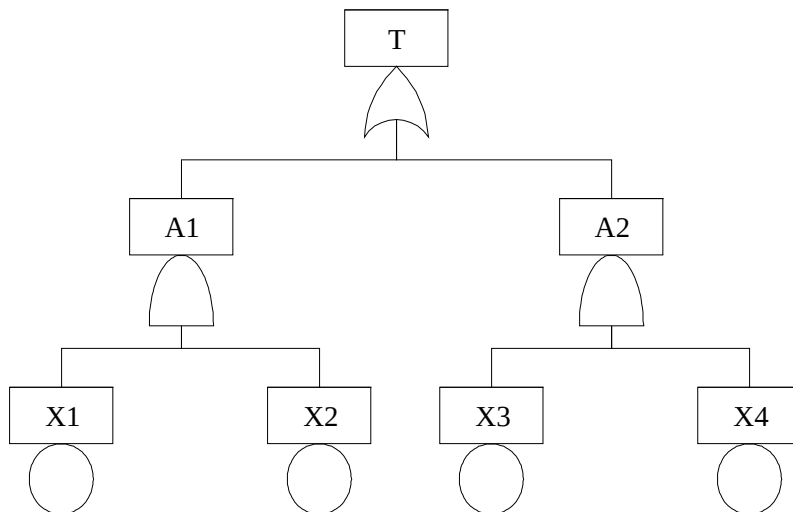
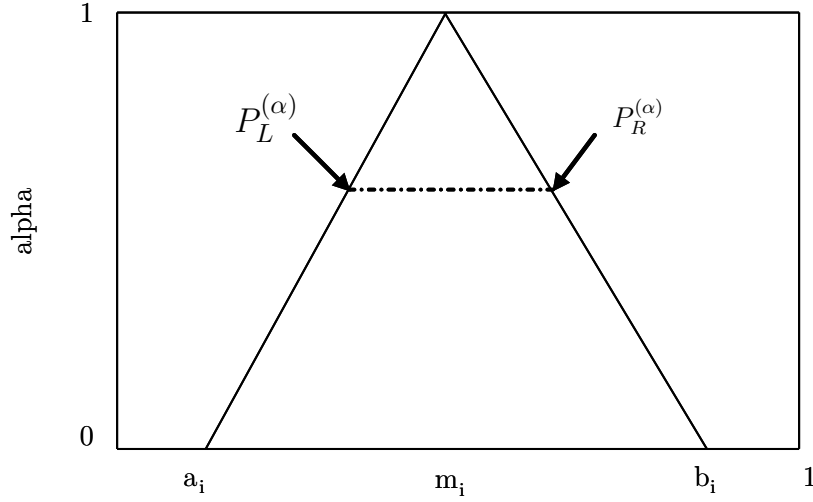


FIG. 3.4 – Arbre de défaillance


 FIG. 3.5 – α -coupes d'une probabilité floue

d'appartenance des nombres flous en α -coupes. En effet, si nous considérons un nombre flou $\widetilde{P}_A$ de fonction d'appartenance $\mu_{\widetilde{P}_A}(x)$ (cf. figure 3.5), nous pouvons obtenir plusieurs intervalles emboîtés en utilisant la méthode des α -coupes. $P_{A_L}^{(\alpha)}$ et $P_{A_R}^{(\alpha)}$ représenteront respectivement les limites gauche et droite de la fonction d'appartenance $\mu_{\widetilde{P}_A}(x)$ à chaque α -coupe.

Les opérations arithmétiques appliquées à deux probabilités floues $\widetilde{P}_A$ et $\widetilde{P}_B$ donnent les expressions suivantes (cf. équations 2.20, 2.21, et 2.22) :

$$\widetilde{P}_C = \widetilde{P}_A + \widetilde{P}_B \rightarrow [P_{C_L}^{(\alpha)}, P_{C_R}^{(\alpha)}] = [P_{A_L}^{(\alpha)} + P_{B_L}^{(\alpha)}, P_{A_R}^{(\alpha)} + P_{B_R}^{(\alpha)}]$$

$$\widetilde{P}_C = \widetilde{P}_A - \widetilde{P}_B \rightarrow [P_{C_L}^{(\alpha)}, P_{C_R}^{(\alpha)}] = [P_{A_L}^{(\alpha)} - P_{B_L}^{(\alpha)}, P_{A_R}^{(\alpha)} - P_{B_R}^{(\alpha)}]$$

$$\widetilde{P}_C = \widetilde{P}_A \cdot \widetilde{P}_B \rightarrow [P_{C_L}^{(\alpha)}, P_{C_R}^{(\alpha)}] = [\min(P_{A_L}^{(\alpha)} \cdot P_{B_L}^{(\alpha)}, P_{A_R}^{(\alpha)} \cdot P_{B_L}^{(\alpha)}, P_{A_L}^{(\alpha)} \cdot P_{B_R}^{(\alpha)}, P_{A_R}^{(\alpha)} \cdot P_{B_R}^{(\alpha)}), \max(P_{A_L}^{(\alpha)} \cdot P_{B_L}^{(\alpha)}, P_{A_R}^{(\alpha)} \cdot P_{B_L}^{(\alpha)}, P_{A_L}^{(\alpha)} \cdot P_{B_R}^{(\alpha)}, P_{A_R}^{(\alpha)} \cdot P_{B_R}^{(\alpha)})]$$

Finalement, nous calculons la probabilité floue d'occurrence de l'événement sommet $\widetilde{P}_T(y)$ à plusieurs instants t_i sur une période T_d . Le $P\widetilde{F}\widetilde{D}_{avg}$ est la probabilité floue moyenne des probabilités floues d'occurrence de l'événement sommet sur la période T_d .

3.3 Caractérisation de l'incertitude des résultats obtenus

Il est clair que le SIL serait déterminé directement à partir du tableau 1.3 défini dans la norme IEC 61508 [IEC61508, 1998] si nous avions obtenu une valeur singulière du PFD_{avg} . Or, nous sommes en présence d'une fonction d'appartenance du PFD_{avg} qui conduit à une incertitude dans la détermination du SIL. Nous sommes donc naturellement amené à caractériser l'incertitude concernant le SIL obtenu. Pour caractériser cette incertitude et dire dans quelle mesure un SIL i est possible et dans quelle mesure nous sommes certain de l'avoir, nous avons choisi d'introduire des mesures de possibilité et de nécessité [Prade and Dubois, 1983; Dubois and Prade, 1988] définies dans le cadre de la théorie des possibilités (cf. section 2.2.3). L'introduction de ces

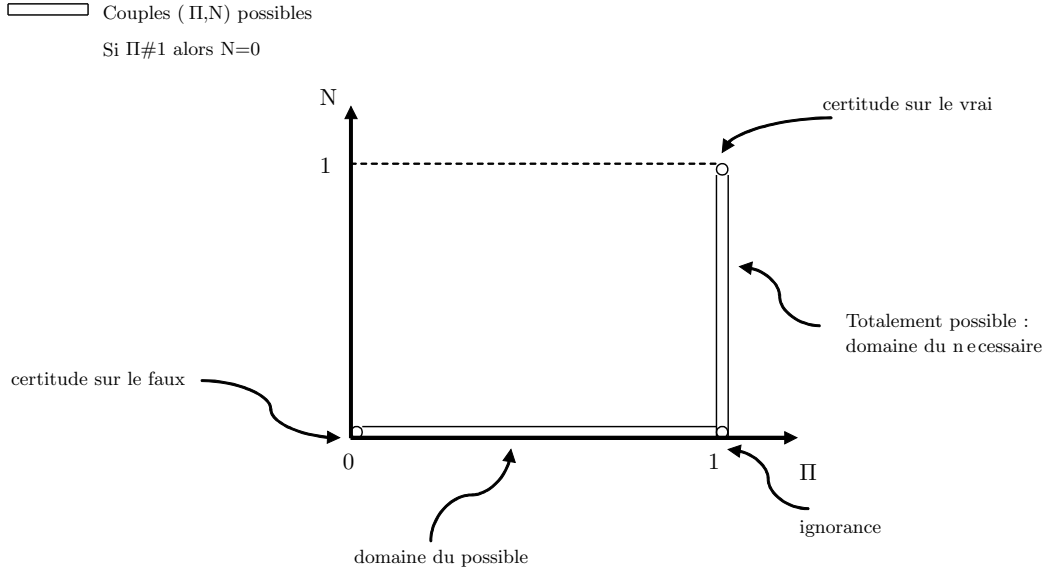


FIG. 3.6 – Traduction de l'incertitude par les mesures de possibilité et de nécessité

mesures a pour unique objectif de caractériser l'incertitude des résultats obtenus. D'après la figure 3.6, on remarque qu'un événement qui n'est pas possible ($\Pi \neq 1$) ne peut pas être certain ($N = 0$) et qu'une fois ($\Pi = 1$) on peut avoir des degrés de nécessité différents de 0, jusqu'à l'événement certain ($\Pi = N = 1$).

3.4 Etudes de sensibilité

Pour réduire l'incertitude qui entache le SIL du SIS, il est légitime de s'intéresser en priorité aux composants qui contribuent le plus à cette incertitude. Pour ce faire, nous avons introduit des facteurs d'importance adaptés à notre approche probabiliste floue pour identifier ces composants. Nous commençons d'abord par un rappel des facteurs d'importance probabilistes et ceux intégrant les connaissances imparfaites puis nous introduisons ensuite les facteurs flous que nous avons proposés [Sallak *et al.*, 2006c].

3.4.1 Facteurs d'importance probabilistes

Il est souvent utile d'associer à chaque composant d'un système un ensemble de mesures permettant de rendre compte de l'importance qu'il occupe vis-à-vis des mesures de fiabilité ou de disponibilité. Ces mesures sont appelées facteurs d'importance. La notion de facteurs d'importance a été introduite par Birnbaum [Birnbaum, 1969] et développée par de nombreux auteurs [Barlow and Proschan, 1973; Chatterjee, 1974; Fussell, 1975]. Lambert [Lambert, 1975] en a dressée une liste assez complète et a développé le premier programme informatique de calcul de tels facteurs.

Au niveau de la conception des systèmes, il est important de pouvoir répondre à des questions telles que :

- Quel composant faut-il améliorer en priorité pour augmenter la fiabilité du système ?

- Quels sont les paramètres des composants qui ont le plus d'influence sur la fiabilité du système ?

Au niveau de l'exploitation, il s'avère intéressant de répondre à des questions telles que :

- Quelles sont les coupes minimales les plus importantes ?
- Sachant que le système est en panne, quel composant faut-il réparer en priorité ?

L'objectif du calcul des facteurs d'importance est de répondre à ces questions. Dans cette section, nous étudions les quatre principaux facteurs d'importance probabilistes utilisés et leurs différentes interprétations :

- Le facteur d'importance marginal qui permet de mesurer la variation de l'indisponibilité du système en fonction de l'indisponibilité de ses composants [Birnbaum, 1969].
- Le facteur d'importance critique qui représente la probabilité qu'un composant ait provoqué la défaillance du système sachant que le système est défaillant [Lambert, 1975].
- Le facteur d'importance critique des coupes minimales qui représente la probabilité qu'une coupe minimale ait provoqué la défaillance du système sachant que le système est défaillant [Fussell, 1975].
- Le facteur d'importance de diagnostic qui représente la probabilité pour qu'un composant soit en panne à l'instant t sachant que le système est en panne à cet instant [Vesely, 1981].

3.4.1.1 Hypothèses et notations

Nous supposons que chaque composant ne possède que deux états : un état de fonctionnement, noté 1, et un état de défaillance noté 0.

Nous désignons par $x_i(t)$ l'état du composant i à l'instant t et par $x(t)$ l'état du système à cet instant.

Notons φ l'application de $\{0, 1\}^n$ dans $\{0, 1\}$ définie par $\varphi(x(t)) = 1$ si le système fonctionne à l'instant t et $\varphi(x(t)) = 0$ si le système est défaillant à cet instant. La fonction φ s'appelle la fonction de structure du système.

La relation entre l'état du système $x(t)$ et l'état de ses composants est donnée par :

$$x(t) = \varphi(x_1(t), x_2(t), \dots, x_n(t)) \quad (3.4)$$

$E(x(t))$ désignera l'espérance mathématique de $x(t)$.

Pour $x(t) = \varphi(x_1(t), \dots, x_n(t)) \in \mathbb{R}^n$ et $1 \leq i \leq n$, nous posons :

$$(1_i, x(t)) = (x_1(t), \dots, x_{i-1}(t), 1, x_{i+1}(t), \dots, x_n(t)) \quad (3.5)$$

$$(0_i, x(t)) = (x_1(t), \dots, x_{i-1}(t), 0, x_{i+1}(t), \dots, x_n(t)) \quad (3.6)$$

Ainsi, $\varphi(1_i, x(t))$ désigne l'état du système lorsque le composant i fonctionne ($x_i(t) = 1$), et $\varphi(0_i, x(t))$ désigne l'état du système lorsque i est défaillant ($x_i(t) = 0$).

Nous définissons aussi une relation d'ordre sur $x(t)$ par :

$$(1_i, x(t)) \geq (0_i, x(t)) \quad (3.7)$$

Un système est monotone, si nous avons pour tout $x(t) \in \{0, 1\}^n$:

$$\varphi(1_i, x(t)) \geq \varphi(0_i, x(t)) \quad (3.8)$$

Un système est cohérent, lorsqu'il est monotone et vérifie en plus les deux égalités :

$$\varphi(1) = 1 \quad \varphi(0) = 0 \quad (3.9)$$

Par conséquent, nous avons pour $x(t) \in \{0, 1\}^n$:

$$\varphi(1_i, x(t)) - \varphi(0_i, x(t)) \in \{0, 1\}, \quad (3.10)$$

$$\varphi(1_i, x(t)) - \varphi(0_i, x(t)) = 1 \Leftrightarrow \varphi(1_i, x(t)) = 1, \varphi(0_i, x(t)) = 0 \quad (3.11)$$

Enfin, nous désignerons par $q_i(t)$ l'indisponibilité du composant i à l'instant t et par $Q(t)$ l'indisponibilité du système à cet instant. Nous avons donc :

$$Q(t) = 1 - E(\varphi(x(t))) \quad (3.12)$$

Puisque les variables aléatoires $x_1(t), \dots, x_n(t)$ sont indépendantes et que la fonction φ est affine par rapport à chaque variable, nous obtenons finalement :

$$Q(t) = 1 - \varphi(E(x_1(t)), E(x_2(t)), \dots, E(x_n(t))) = 1 - \varphi(q_1(t), q_2(t), \dots, q_n(t)) \quad (3.13)$$

3.4.1.2 Facteur d'importance marginal

Ce facteur d'importance a été introduit par Birnbaum [Birnbaum, 1969]. Il est aussi appelé *facteur d'importance de Birnbaum*. Il existe deux interprétations possibles de ce facteur que nous allons présenter.

Première interprétation

Le facteur d'importance marginal d'un composant i est défini par :

$$B_i(t) = \frac{\partial Q}{\partial q_i}(t) \quad (3.14)$$

Il permet donc de mesurer la variation de l'indisponibilité Q du système en fonction de l'indisponibilité q_i du composant i .

Deuxième interprétation

Pour les systèmes cohérents, la fonction de structure φ s'écrit (formule de Shannon) :

$$\varphi(x_1, \dots, x_n) = x_i \cdot \varphi(1_i, x(t)) + (1 - x_i) \cdot \varphi(0_i, x(t)) \quad (3.15)$$

d'où :

$$\frac{\partial \varphi}{\partial x_i}(x_1, \dots, x_n) = \varphi(1_i, x(t)) - \varphi(0_i, x(t)) \quad (3.16)$$

D'après la définition (3.14), nous obtenons :

$$B_i(t) = \frac{\partial \varphi(q_1(t), \dots, q_n(t))}{\partial q_i}(t) = \varphi(1_i, q(t)) - \varphi(0_i, q(t)) \quad (3.17)$$

et puisque l'indisponibilité du système est définie par :

$$\varphi(q_i, q(t)) = E(\varphi(x_i, x(t))) \quad (3.18)$$

Nous obtenons :

$$B_i(t) = E(\varphi(1_i, x(t))) - E(\varphi(0_i, x(t))) = E(\varphi(1_i, x(t)) - \varphi(0_i, x(t))) \quad (3.19)$$

Si i est un composant critique du système, alors :

$$\varphi(1_i, x(t)) - \varphi(0_i, x(t)) = 1 \quad (3.20)$$

D'où :

$$B_i(t) = P(\varphi(1_i, x(t)) - \varphi(0_i, x(t)) = 1) \quad (3.21)$$

Nous obtenons finalement :

$$B_i(t) = P(\varphi(1_i, x(t)) = 1, \varphi(0_i, x(t)) = 0) \quad (3.22)$$

Le facteur d'importance marginal représente donc la probabilité que le système fonctionne à l'instant t lorsque le composant i fonctionne à cet instant et que le système soit défaillant à l'instant t lorsque le composant i est défaillant à cet instant.

A partir de l'équation (3.19), une autre expression de B_i est souvent utilisée. Elle est donnée par la relation :

$$B_i(t) = Q(1_i, q(t)) - Q(0_i, q(t)) \quad (3.23)$$

3.4.1.3 Facteur d'importance critique

Ce facteur d'importance a été introduit par Lambert [Lambert, 1975]. Il a également deux interprétations :

Première interprétation

Le facteur d'importance critique ou *facteur d'importance de Lambert* d'un composant i est défini par :

$$L_i(t) = \frac{q_i(t)}{Q(t)} \frac{\partial Q}{\partial q_i}(t) \quad (3.24)$$

Il représente donc le produit du facteur d'importance de Birnbaum et le rapport de l'indisponibilité du composant i et de l'indisponibilité du système.

Deuxième interprétation

Pour les systèmes cohérents, nous pouvons écrire que :

$$L_i(t) = \frac{q_i(t)}{Q(t)} \frac{\partial \varphi(q_1(t), \dots, q_n(t))}{\partial q_i}(t) \quad (3.25)$$

D'où :

$$L_i(t) = \frac{q_i(t)}{Q(t)} (\varphi(1_i, q(t)) - \varphi(0_i, q(t))) \quad (3.26)$$

En utilisant le résultat (3.25) et l'indépendance des variables aléatoires $x_i(t)$, nous obtenons :

$$L_i(t) = \frac{1}{P(\varphi(x(t)) = 0)} P(x_i(t) = 0) P(\varphi(1_i, x(t)) = 1, \varphi(0_i, x(t)) = 0) \quad (3.27)$$

D'où :

$$L_i(t) = \frac{1}{P(\varphi(x(t)) = 0)} P(x_i(t) = 0, \varphi(1_i, x(t)) = 1, \varphi(0_i, x(t)) = 0) \quad (3.28)$$

D'après le théorème de Bayes :

$$L_i(t) = P(x_i(t) = 0, \varphi(1_i, x(t)) = 1, \varphi(0_i, x(t)) = 0 / \varphi(x(t)) = 0) \quad (3.29)$$

Le facteur d'importance critique est donc la probabilité que le composant i soit en panne à l'instant t et que le système ait le même état que le composant i , sachant que le système est en panne à l'instant t .

Cela peut s'interpréter comme étant la probabilité que le composant i ait provoqué la défaillance du système sachant que le système est défaillant.

Il faut noter que le composant i n'est pas forcément le seul en panne, mais dans ce cas, il est le dernier à être tombé en panne (c'est-à-dire que sa défaillance a provoqué la défaillance du système).

3.4.1.4 Facteur d'importance critique des coupes minimales

Le facteur d'importance critique d'une coupe minimale peut aussi être calculé. Cette notion a été introduite par Fussell et Vesely [Fussell, 1975] ; il est aussi appelé *facteur d'importance critique d'une coupe minimale de type Fussell-Vesely*. Il représente la probabilité qu'une coupe minimale ait provoqué la défaillance du système sachant que le système est défaillant. Il permet d'indiquer ainsi le poids respectif de chaque coupe minimale dans sa contribution à la défaillance du système. Nous obtenons simplement :

$$L_i = \frac{Q_i(t)}{Q(t)} \quad (3.30)$$

$Q_i(t)$ étant l'indisponibilité liée à la coupe minimale i et $Q(t)$ l'indisponibilité du système complet.

3.4.1.5 Facteur d'importance de diagnostic

Cette notion a été introduite par Vesely [Vesely, 1981] et a été utilisée par Fussell [Fussell, 1975]. Elle a également deux interprétations :

Première interprétation

Le facteur d'importance de diagnostic ou *facteur d'importance de Vesely-Fussell* d'un composant i est défini par :

$$VF_i(t) = \frac{q_i(t)}{Q(t)} \varphi(1_i, q(t)) \quad (3.31)$$

Deuxième interprétation

Pour les systèmes cohérents, nous pouvons écrire que :

$$VF_i(t) = \frac{q_i(t)}{Q(t)} P(\varphi(1_i, q(t)) = 0) \quad (3.32)$$

D'où :

$$VF_i(t) = \frac{P(x_i(t) = 0, \varphi(1_i, q(t)) = 0)}{P(\varphi(x(t)) = 0)} \quad (3.33)$$

D'après le théorème de Bayes :

$$VF_i(t) = P(x_i(t) = 0 / \varphi(x(t)) = 0) \quad (3.34)$$

Le facteur d'importance de diagnostic est donc la probabilité pour que le composant i soit en panne à l'instant t sachant que le système est en panne à cet instant.

3.4.1.6 Observations et limitations

- Ces facteurs d'importance peuvent donner des résultats différents pour les composants d'un même système. Cela est dû aux égalités qui les définissent.
- Pour identifier les composants qui peuvent améliorer significativement la fiabilité d'un système, le facteur d'importance de Birnbaum est le plus approprié.
- Pour identifier les composants qui ont causé la défaillance du système, nous pouvons utiliser soit le facteur d'importance critique, soit le facteur de Vesely-Fussell. Nous pourrions ainsi établir une liste des composants qui doivent être réparés en priorité.
- Ces quatre facteurs d'importance dépendent du temps : à différents instants nous pouvons avoir pour un même composant des facteurs d'importance différents.
- Les systèmes réparables et non réparables peuvent être considérés lors du calcul des facteurs d'importance. Cependant, il est clair que l'interprétation de ces facteurs est beaucoup plus difficile dans le cas des systèmes réparables.

3.4.2 Facteurs d'importance intégrant les connaissances imparfaites

Pour compléter l'étude de l'intégration des connaissances imparfaites dans les études de fiabilité, il est essentiel d'étudier les facteurs d'importance qui ont été définis dans ce cadre. Dans la littérature, il existe principalement deux types :

- Les facteurs d'importance sur l'incertitude ;
- Les facteurs d'importance flous.

3.4.2.1 Facteurs d'importance sur les incertitudes

Ces facteurs sont utilisés pour l'étude de l'influence des variances des probabilités des événements de base sur la variance de la probabilité de l'événement sommet.

Facteur d'importance de Pan-Tai

Analogue à celui de Birnbaum. Ce facteur a été introduit par Pan et Tai [Pan and Tai, 1988]. Il est défini par :

$$F_{PN}^i = \frac{\partial var(Q)}{\partial var(q_i)} \quad (3.35)$$

Il peut s'exprimer aussi par :

$$F_{PN}^i = E\left\{\left(\frac{\partial Q}{\partial q_i}\right)^2\right\} \quad (3.36)$$

Facteur d'importance de Bier

Adapté de celui de Lambert. Ce facteur a été introduit par Bier [Bier, 1983]. Il est défini par :

$$F_{BR}^i = \frac{\text{var}(q_i)}{\text{var}(Q)} \frac{\partial \text{var}(Q)}{\partial \text{var}(q_i)} \quad (3.37)$$

Observations et limitations

- Le calcul des facteurs d'importance introduits par Pan et Tai [Pan and Tai, 1988] et Bier [Bier, 1983] devient une tâche très difficile dans le cas de systèmes complexes contenant un nombre important de composants de base.
- Ces facteurs d'importance peuvent aussi être utilisés dans le cadre des sous-ensembles flous. En effet, nous pouvons utiliser les définitions de la variance et des moments flous donnés par Matarazzo [Matarazzo and Munda, 2001] qui définit à un instant donné t , la variance d'une indisponibilité q (un sous-ensemble flou en général) d'un composant par :

$$E(q) = \frac{\int_{-\infty}^{+\infty} x \pi_q(x) dx}{\int_{-\infty}^{+\infty} \pi_q(x) dx} \quad (3.38)$$

$$\text{var}(q) = \frac{\int_{-\infty}^{+\infty} (x - E(q))^2 \pi_q(x) dx}{\int_{-\infty}^{+\infty} \pi_q(x) dx} \quad (3.39)$$

où $\pi_q(x)$ est la distribution de possibilité de q à l'instant t .

3.4.2.2 Facteurs d'importance flous

Les facteurs d'importance flous ont été principalement introduits dans les études des arbres de défaillance. Ils indiquent la contribution des probabilités de défaillance des composants de base (dans les arbres de défaillance cela se traduit par les probabilités d'occurrence des événements de base) à la probabilité de défaillance du système complet (probabilité d'occurrence de l'événement sommet). Birnbaum [Birnbaum, 1969] a considéré que la contribution totale de la défaillance d'un composant à celle du système complet était égale à la différence entre les indisponibilités du système complet quand le composant i fonctionne, puis quand il est défaillant :

$$B_i(t) = Q(1_i, q(t)) - Q(0_i, q(t)) \quad (3.40)$$

Dans l'approche floue, les indisponibilités $Q(1_i, q(t))$ et $Q(0_i, q(t))$ ne sont plus des valeurs singulières mais des nombres flous. Il faut donc apporter une nouvelle définition de ces facteurs d'importance.

Facteur de mesure flou

A partir d'une analyse floue des arbres de défaillance, Suresh *et al.* [P. V. Suresh and Raj, 1996] ont défini un *facteur de mesure flou* :

$$Sa = ED[Q(1_i, q) - Q(0_i, q)] \quad (3.41)$$

Où ED désigne la distance euclidienne entre les deux nombres flous $Q(1_i, q)$ et $Q(0_i, q)$.
En effet, si nous posons :

$$Q(1_i, q) = [a, b] \quad Q(0_i, q) = [c, d] \quad (3.42)$$

Alors, nous obtenons :

$$Sa = \sum_{\alpha_i=0,1,\dots,n} ((a-c)^2 + (b-d)^2)^{0.5} \quad (3.43)$$

Où les α_i représentent les α -coupes des distributions des nombres flous.

Si nous nous basons sur les opérations arithmétiques définies dans [Bouchon-Meunier, 1995], alors nous obtenons la relation suivante :

$$Sa = \sup_{\{(x,y)/z=x-y\}} \min(\pi_Q(x), \pi_Q(y)) \quad (3.44)$$

où $\pi_Q(x)$ est la distribution de possibilité de q .

Ainsi, si l'indisponibilité d'un composant $q_i(t)$ est représentée par un nombre flou triangulaire, alors :

$$q_i(t) = f(a_{i1}(t), a_{i2}(t), a_{i3}(t)) \quad (3.45)$$

Où :

$a_{i2}(t)$ est la valeur modale de $q_i(t)$ avec $\pi_{q_i(t)}(a_{i2}(t)) = 1$, $a_{i1}(t)$ est la largeur de son support à gauche de $a_{i2}(t)$ et $a_{i3}(t)$ celle de son support à droite de $a_{i2}(t)$. f est la fonction qui associe à chaque indisponibilité $q_i(t)$ d'un composant, un ensemble de paramètres $a_i(t)$ qui caractérisent sa distribution de possibilité.

Facteur d'incertitude flou

Nous considérons la fonction de structure du système φ définie par :

$$\varphi : q(t) \mapsto Q(t) \quad (3.46)$$

où :

$$q(t) = (q_1(t), q_2(t), \dots, q_n(t)) \quad (3.47)$$

La relation entre l'indisponibilité du système $Q(t)$ et les indisponibilités de ses composants sera donnée par :

$$Q(t) = \varphi(q_1(t), q_2(t), \dots, q_n(t)) \quad (3.48)$$

Ainsi, le facteur d'incertitude flou Sb_i d'un composant i sera défini par :

$$Sb_i(t) = \frac{\partial Q(t)}{\partial a_i(t)} \quad (3.49)$$

or :

$$Q(t) = \varphi(q(t)) = \varphi(q_1(t), q_2(t), \dots, q_n(t)) \quad (3.50)$$

d'où :

$$Sb_i(t) = \frac{\partial[\varphi(q(t))]}{\partial a_i(t)} = \frac{\partial[\varphi(q_1(t), q_2(t), \dots, q_n(t))]}{\partial a_i(t)} \quad (3.51)$$

or :

$$\forall i = 1, \dots, n \quad q_i(t) = f(a_i(t)) \quad (3.52)$$

D'où :

$$Sb_i(t) = \frac{\partial[\varphi(f(a_1(t), a_2(t), \dots, a_n(t)))]}{\partial a_i(t)} \quad (3.53)$$

D'après la formule de dérivation du composé de deux fonctions u et v :

$$d(uov) = du(v) * dv \quad (3.54)$$

Nous obtenons finalement la formule générale qui permet de calculer le facteur d'incertitude flou :

$$Sb_i(t) = \frac{\partial \varphi(f(a_1(t), a_2(t), \dots, a_n(t)))}{\partial a_i(t)} \frac{\partial f(a_1(t), a_2(t), \dots, a_n(t))}{\partial a_i(t)} \quad (3.55)$$

3.4.3 Facteurs d'importance proposés

Nous avons choisi de définir nos propres facteurs d'importances qui seront adaptés aux fonctions d'appartenances des probabilités floues que nous obtenons [Sallak *et al.*, 2006a]. En outre, nous avons choisi pour la définition de ces facteurs les opérateurs arithmétiques flous qui nous simplifieront les calculs.

3.4.3.1 Facteur d'importance flou (FPIM)

Le premier facteur que nous proposons permet d'identifier les composants critiques du système du point de vue de la fiabilité et de la disponibilité [Sallak *et al.*, 2006c; Sallak *et al.*, 2006b]. Ainsi, le facteur d'importance flou $FPIM_i$ d'un composant i est donné par la relation suivante :

$$FPIM_i = defuz(\widetilde{FPIM}_i) \quad (3.56)$$

où $defuz$ est la méthode de défuzzification du centre de gravité de la distribution de probabilité utilisée pour obtenir une valeur singulière à partir de la distribution de probabilité floue de $\widetilde{FPIM}_i$ qui est donnée par :

$$\widetilde{FPIM}_i = \widetilde{P} - \widetilde{P}^i \quad (3.57)$$

où $\widetilde{P}$ est la distribution de la probabilité floue de défaillance du SIS quand le composant i fonctionne et $\widetilde{P}^i$ la distribution de la probabilité floue de défaillance du SIS quand le composant i est défaillant.

3.4.3.2 Facteur d'incertitude flou (FPUM)

Nous proposons aussi, un autre facteur d'incertitude floue $FPUM_i$ d'un composant i qui est donné par la relation suivante [Sallak *et al.*, 2006c; Sallak *et al.*, 2006b] :

$$FPUM_i = defuz(\widetilde{FPUM}_i) \quad (3.58)$$

La distribution de possibilité de $\widetilde{FPUM}_i$ est donnée par :

$$\widetilde{FPUM}_i = \widetilde{P} - \widetilde{P}_{P_i=cts} \quad (3.59)$$

où $\widetilde{P}$ est la distribution de la probabilité floue de défaillance du SIS et $\widetilde{P}_{P_i=cts}$ la distribution de la probabilité floue de défaillance du SIS quand nous annulons l'incertitude qui entache la

probabilité de défaillance du composant i (c'est-à-dire que nous considérons une valeur précise de la probabilité de défaillance de ce composant).

Ce facteur d'incertitude flou permet d'identifier les composants dont l'incertitude de la probabilité de défaillance contribue significativement à l'incertitude entachant la probabilité de défaillance du système.

3.4.4 Conclusion partielle

L'étude des facteurs d'importance constitue une partie essentielle des études de sûreté de fonctionnement des systèmes. Selon les problèmes rencontrés, nous pouvons utiliser différents facteurs d'importance.

En présence de connaissances suffisantes pour établir des valeurs de taux de défaillance précises, il est préférable d'utiliser les facteurs d'importance probabilistes. Ainsi, pour identifier les composants qui permettent d'augmenter la fiabilité des systèmes, le facteur de Birnbaum semble le plus adapté. Pour dresser la liste des composants qui ont causé la défaillance du système et qui doivent donc être réparés en priorité, nous nous devons d'utiliser le facteur d'importance de Lambert ou de Vesely-Fussel. L'emploi d'autres facteurs définis dans la littérature (Facteur de réduction du risque, etc.) peut s'avérer nécessaire selon l'objectif des études réalisées.

En présence de connaissances imprécises, l'utilisation des facteurs flous est préférable. Les facteurs d'importance introduits par Pan et Tai [Pan and Tai, 1988] et Bier [Bier, 1983] sont convenables lorsque nous disposons de paramètres tels que la variance ou la moyenne mais nécessitent souvent des calculs très longs en présence de systèmes complexes. Les facteurs d'importance que nous avons proposé semblent les plus simples pour traiter des problèmes prenant en compte les imprécisions des paramètres de fiabilité des composants.

3.5 Cas d'application

3.5.1 Description du système

Nous considérons un système constitué d'un réservoir sous pression contenant un liquide inflammable volatil (cf. figure 3.7). Ce réservoir peut rejeter des gazs dans l'atmosphère. Nous supposons que le risque acceptable est défini sous forme d'un taux moyen de rejet de gaz de fréquence inférieure à 10^{-4} par an. Une analyse des phénomènes dangereux liés à ce système a montré que les systèmes de protection disponibles (alarmes et niveaux de protection) sont insuffisants pour assurer ce risque acceptable (le non dépassement du seuil imposé pour le rejet des gazs) et qu'une nouvelle fonction instrumentée de sécurité (SIF) de niveau SIL2 doit être implémentée dans un SIS pour réduire le taux de rejet du réservoir sous pression. Notre objectif est de vérifier si le SIS proposé au concepteur (cf. figure 3.8) est capable de satisfaire à l'exigence SIL 2 requise pour réaliser la fonction instrumentée de sécurité de réduction du rejet des gazs. Le SIS utilisé a été défini dans le document ISA-TR84.00.02-2002 [ISA-TR84.00.02-2002, 2002] qui explique les différentes techniques qui permettent de calculer les niveaux de SIL.

La figure 3.9 représente l'arbre de défaillance du SIS [ISA-TR84.00.02-2002, 2002] lors de sa sollicitation pour réaliser la fonction SIF demandée (réduction du rejet des gazs). Notre premier objectif est de calculer à partir des probabilités floues d'occurrence des événements de base de

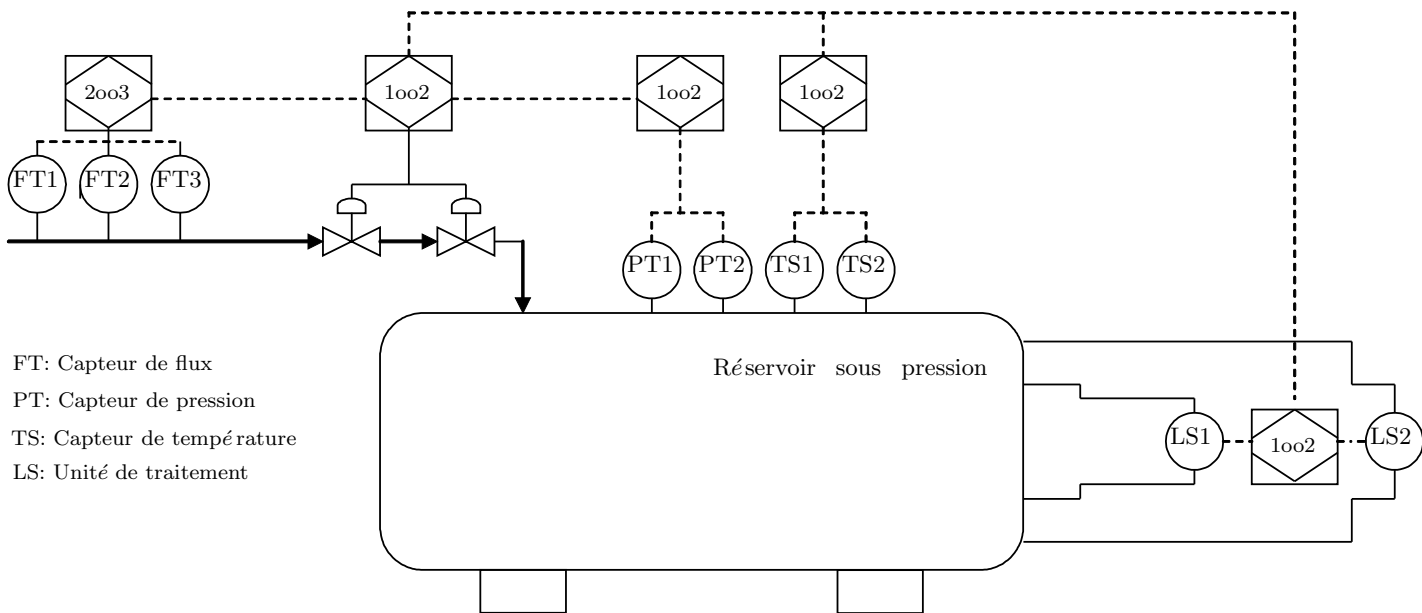


FIG. 3.7 – Réservoir sous pression

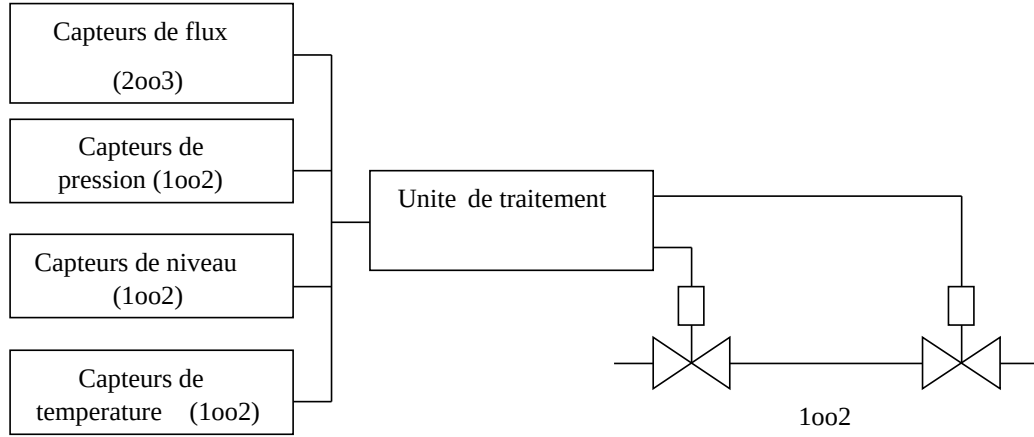


FIG. 3.8 – Configuration du SIS

l'arbre de défaillance (chaque événement de base représente la défaillance d'un composant du SIS), la probabilité floue d'occurrence de l'événement redouté qui représente la défaillance du SIS lors de sa sollicitation.

3.5.2 Hypothèses

- Nous nous plaçons dans le cas où le SIS est faiblement sollicité (le SIS est sollicité moins d'une fois par an) d'où le besoin d'évaluer le PFD et non pas le PFH (probabilité de défaillance par heure). Dans ce cas, le PFD instantané est assimilé à une indisponibilité instantanée.
- Nous nous intéressons à l'évaluation du PFD_{avg} du SIS. C'est pourquoi nous utilisons les taux de défaillance λ des composants qui désignent les taux de défaillance dangereuses non détectés λ^{DU} et les taux de défaillance détectés λ^{DD} mais non réparés :

$$\lambda = \lambda^{DU} + \lambda^{DD} \quad (3.60)$$

Ces défaillances dangereuses font passer le système de l'état normal à l'état de défaillance dangereux.

- Les composants sont non réparables.
- Il n'existe pas de dépendance entre les défaillances des composants (le facteur de causes communes de défaillances β est nul).

3.5.3 Approche probabiliste floue

Toutes les probabilités floues de défaillance des composants sont du type triangulaire et caractérisées par les 3 paramètres m_i , a_i et b_i , tel que m_i est la valeur modale avec $\mu_{PFD}(m_i) = 1$, a_i est la limite à gauche de m_i et b_i est la limite à droite de m_i (cf. figure 3.3). Dans le tableau 3.1, nous donnons les valeurs des 3 paramètres m_i , a_i et b_i pour chaque composant du SIS. Les valeurs données ainsi que les facteurs d'erreurs sont en conformité avec les valeurs usuelles données dans les bases de données de fiabilité [OREDA, 2002; Exida, 2005; Hauge *et al.*, 2006; Goble and Cheddie, 2006].

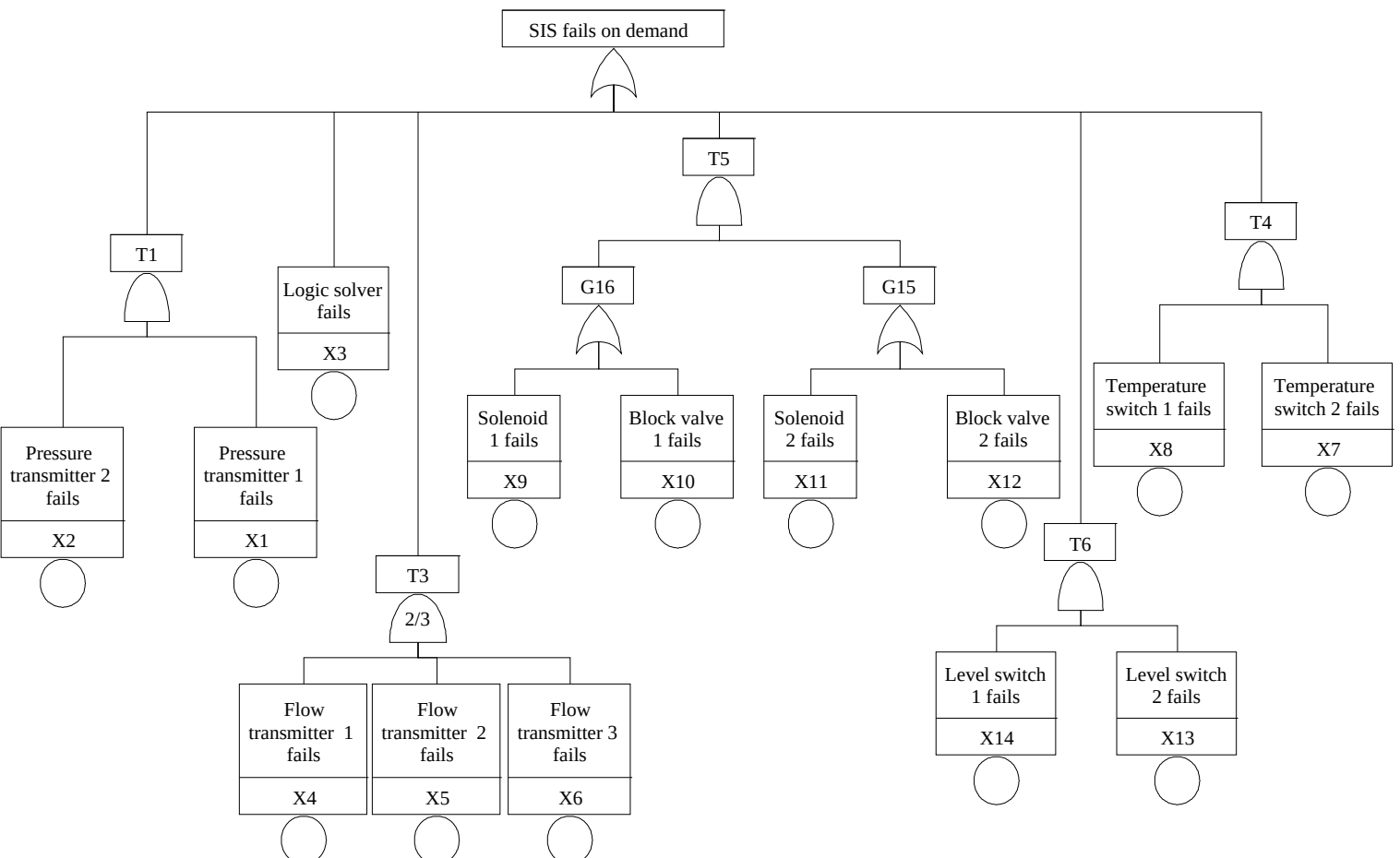


FIG. 3.9 – Arbre de défaillance du système complet

TAB. 3.1 – Paramètres flous des probabilités de défaillance des composants du SIS

Composants du SIS	$a_i(\times 10^{-2})$	$m_i(\times 10^{-2})$	$b_i(\times 10^{-2})$
X1, X2 : Capteurs de pression	2.13	3.2	4.8
X3 : Unité de traitement	0.5	0.6	0.72
X4, X5, X6 : Capteurs de flux	1.31	1.7	2.21
X9, X11 : Solénoïdes des vannes	1.65	2.8	4.76
X7, X8 : Capteurs de température	3.64	4	4.4
X10, X12 : Bloc-vannes	1.65	2.8	4.76
X13, X14 : Capteurs de niveau	3.07	3.99	5.19

TAB. 3.2 – Coupes minimales de l'arbre de défaillance

Coupes minimales
$C_1 = \{X3\}$
$C_2 = \{X1, X2\}$
$C_3 = \{X10, X11\}$
$C_4 = \{X10, X12\}$
$C_5 = \{X13, X14\}$
$C_6 = \{X4, X5\}$
$C_7 = \{X4, X6\}$
$C_8 = \{X5, X6\}$
$C_9 = \{X7, X8\}$
$C_{10} = \{X9, X11\}$
$C_{11} = \{X9, X12\}$

En utilisant la méthode des α -coupes, les opérations arithmétiques définies précédemment et les coupes minimales de l'arbre de défaillance (cf. tableau 3.2), nous déterminons la probabilité floue d'occurrence de l'événement sommet (probabilité floue de défaillance du SIS complet) à partir des probabilités floues de défaillance des composants.

La Figure 3.10 donne la distribution de la probabilité floue d'occurrence de l'événement sommet. Cette probabilité de défaillance varie de 7.4×10^{-3} jusqu'à 2.24×10^{-2} , ce qui donne un SIL 1 ($PFD_{avg} \in [10^{-2}, 10^{-1}]$) ou un SIL2 ($PFD_{avg} \in [10^{-3}, 10^{-2}]$) selon la tableau 1.3 qui définit le niveau de SIL en fonction de la valeur du PFD_{avg} . Nous remarquons qu'il existe une incertitude concernant le niveau de SIL du SIS (1 ou 2), c'est pourquoi nous allons utiliser les facteurs d'importance probabilistes flous pour tenter de réduire cette incertitude.

3.5.4 Caractérisation de l'incertitude

Pour donner une mesure de l'incertitude du SIL obtenu, par analogie avec les mesures caractérisant l'incertitude utilisées dans la théorie de possibilités, nous utilisons des mesures de possibilités et de nécessité [Gouriveau, 2003; Sallak *et al.*, 2006a]. Après avoir calculé le PFD_{avg} du SIS, la mesure de possibilité d'obtention du SIL i est donnée par :

$$\Pi(\widetilde{SIL}_i, \widetilde{PFD}_{avg}) = \sup_{u \in [0,1]} \{ \min(\widetilde{SIL}_i(u), \widetilde{PFD}_{avg}(u)) \}$$

Où $\widetilde{SIL}_i(\cdot)$ est la distribution de possibilité du SIL i et $\widetilde{PFD}_{avg}(\cdot)$ est la fonction d'appartenance du PFD_{avg} .

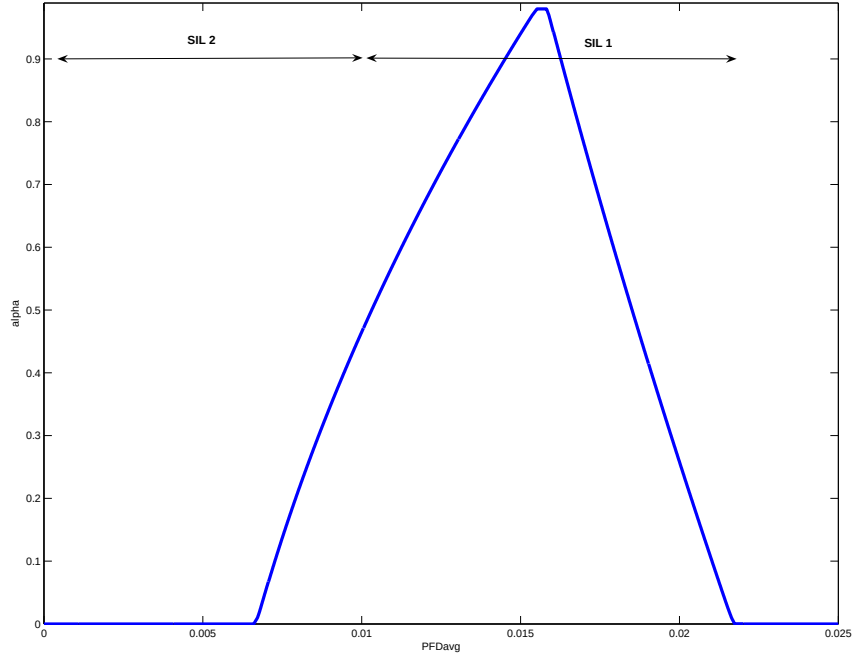


FIG. 3.10 – PFD_{avg} du SIS

SIL	Π	N
1	1	0.45
2	0.55	0
3	0	0
4	0	0

TAB. 3.3 – Mesures de possibilité (Π) et de nécessité (N) des SIL obtenus

Les résultats du calcul des mesures de possibilités montrent qu'il est tout a fait possible d'obtenir un SIL1 ($\Pi = 1$) et qu'il est impossible d'obtenir un SIL3 ou un SIL4 ($\Pi = 0$) (cf. tableau 3.3). Le SIL2 est également possible ($\Pi = 0.55$). Puisqu'il est possible d'avoir soit un SIL1 soit un SIL2, nous allons utiliser des mesures de nécessité.

La mesure de nécessité est donnée par la relation :

$$N(\widetilde{SIL}_i, \widetilde{PFD}_{avg}) = \inf_{u \in [0,1]} \{ \max(\widetilde{SIL}_i(u), 1 - \widetilde{PFD}_{avg}(u)) \}$$

Selon les résultats présentés dans le tableau 3.3, il n'est pas tout a fait certain d'obtenir un SIL1 ($N = 0.35$), et il n'est absolument pas certain du tout d'obtenir un SIL2 ($N = 0$).

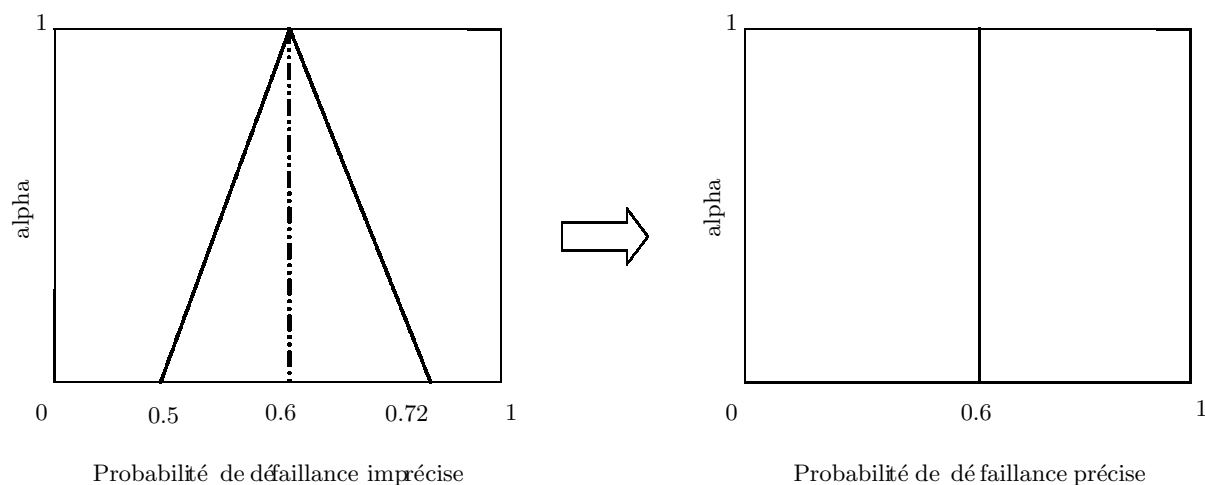


FIG. 3.11 – Probabilités de défaillance précise et imprécise de l'unité de traitement

Composants du SIS	FPIM	Rang
Capteur de pression	0.035	7
Unité de traitement	0.992	1
Capteur de flux	0.038	6
Capteur de température	0.042	3
Solénoïde des vannes	0.042	3
Bloc-vannes	0.042	3
Capteur de niveau	0.082	2

TAB. 3.4 – Facteur d'importance flou (FPIM)

3.5.5 Etudes de sensibilité et réduction de l'incertitude du SIL

3.5.5.1 Facteur d'importance flou (FPIM)

Les résultats du calcul des facteurs d'importance flous (FPIM) des composants du SIS sont donnés dans le tableau 3.4. Nous notons que le composant le plus important est l'unité de traitement avec un facteur de 0.992. Cela signifie que l'unité de traitement est le composant le plus critique pour la fiabilité et l'indisponibilité du SIS.

Pour réduire l'incertitude qui entache le niveau de SIL obtenu, nous proposons de supprimer l'incertitude qui entache la probabilité de défaillance de l'unité de traitement (nous supposons que cette probabilité de défaillance est précise (cf. figure 3.11)).

La figure 3.12 représente la probabilité floue de défaillance du SIS avant (courbe en trait plein) et après la suppression de l'incertitude de la probabilité de défaillance de l'unité de traitement (courbe tiretée). Nous remarquons que l'incertitude qui entache le niveau de SIL n'a pratiquement pas été réduite. Le SIL du SIS peut toujours être 1 ou 2 (cf. tableau 3.4). Nous devons donc utiliser un autre facteur d'importance pour réduire cette incertitude.

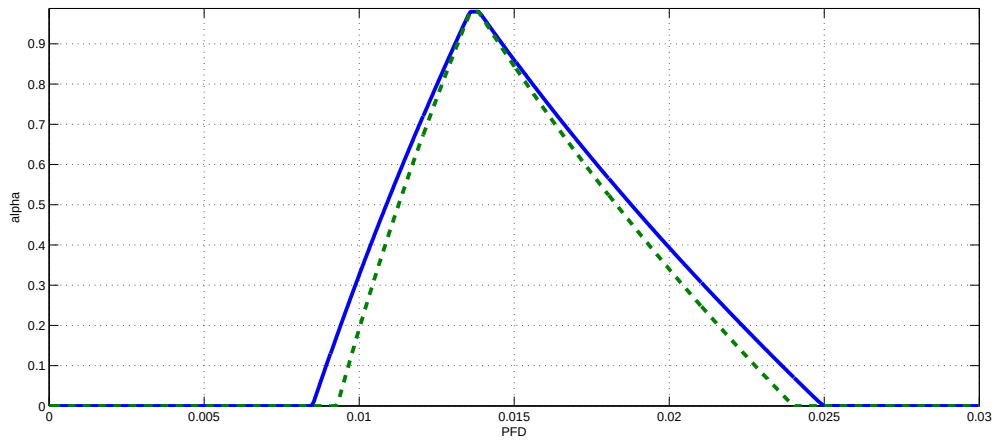


FIG. 3.12 – PFD_{avg} du SIS selon les modifications du SIS après le calcul du FPIM

SIL	Π	N
1	1	0.25
2	0.75	0
3	0	0
4	0	0

TAB. 3.5 – Mesures de possibilité (Π) et de nécessité (N) après modifications selon le facteur (FPIM)

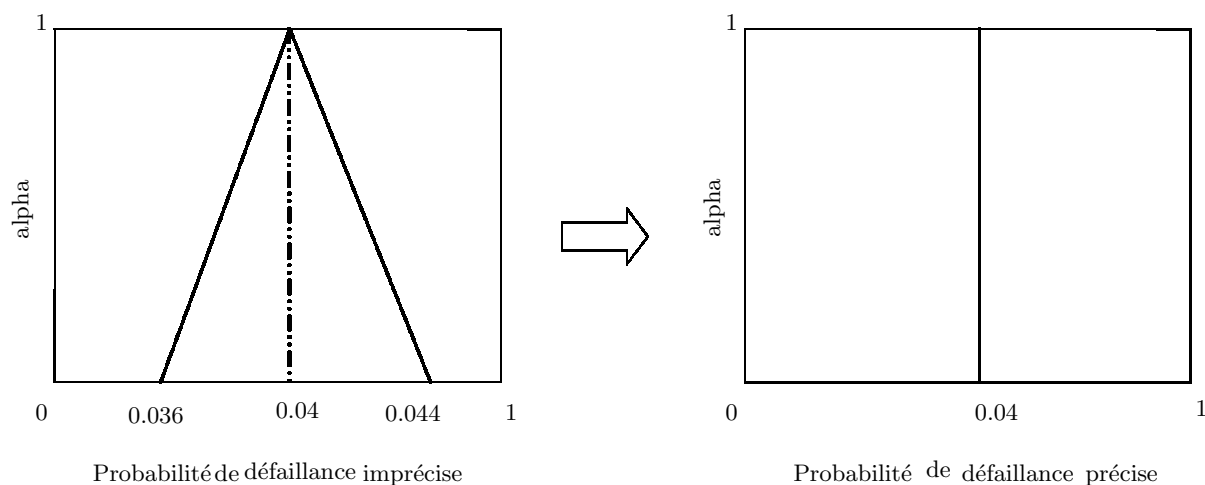


FIG. 3.13 – Probabilités de défaillance précise et imprécise du capteur de température

Composants du SIS	FPUM	Rang
Capteur de pression	0.038	5
Unité de traitement	0.039	2
Capteur de flux	0.006	7
Capteur de température	0.045	1
Solénoïde des vannes	0.039	2
Bloc-vannes	0.039	2
Capteur de niveau	0.039	2

TAB. 3.6 – Facteur d'incertitude flou (FPUM)

3.5.5.2 Facteur d'incertitude flou (FPUM)

Le tableau 3.6 donne les résultats du calcul des facteurs d'incertitude flous (FPUM) des composants du SIS. Nous notons que le classement des composants du SIS n'est plus le même qu'à la section 3.5.5.1. Le composant le plus important est le capteur de température avec un facteur de 0.045.

Pour réduire l'incertitude qui entache le niveau de SIL, nous proposons de supprimer l'incertitude qui entache la probabilité de défaillance du capteur de température (cf. figure 3.13).

La figure 3.14 représente la distribution de la probabilité floue de défaillance du SIS après la suppression de l'incertitude sur la probabilité de défaillance du capteur de température (courbe en pointillés). Nous remarquons que l'incertitude qui entache le niveau de SIL a été pratiquement éliminée. Nous sommes sûr que le SIL du SIS est 1 (cf. tableau 3.7). Dans notre cas, nous concluons que pour réduire l'incertitude du niveau de SIL, le facteur d'incertitude flou est le plus efficace.

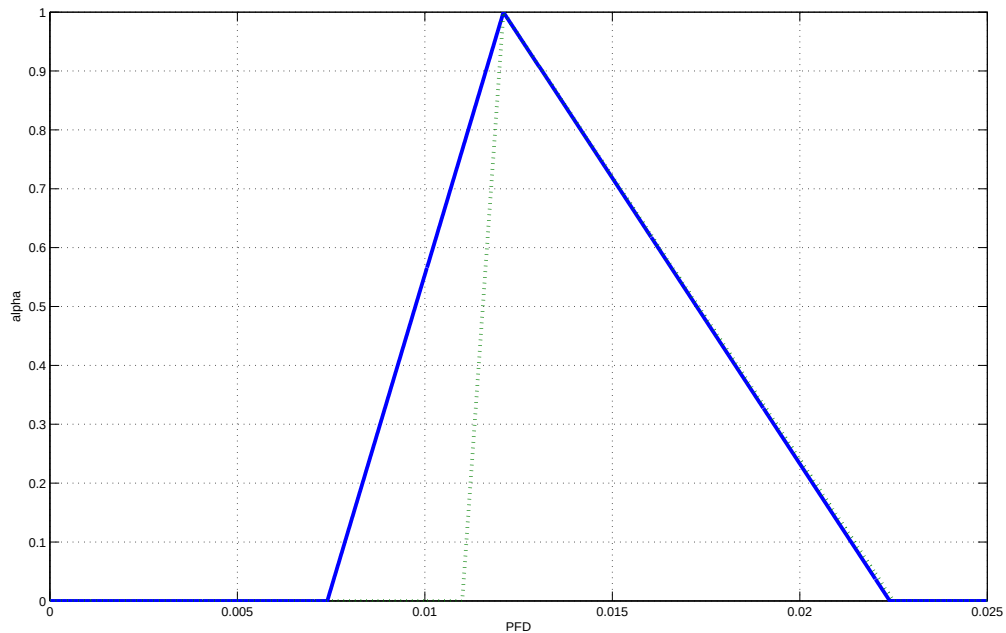


FIG. 3.14 – PFD_{avg} du SIS selon les modifications du SIS après le calcul du FPUM

SIL	Π	N
1	1	1
2	0	0
3	0	0
4	0	0

TAB. 3.7 – Mesures de possibilité (Π) et de nécessité (N) après modifications selon le facteur $FPUM$

3.6 Conclusions et perspectives

Dans ce chapitre, nous avons proposé l'utilisation des arbres de défaillance flous pour évaluer le PFD_{avg} des SIS. En utilisant la méthode de Soman et Misra [Soman and Misra, 1993] basée sur les α -coupes, nous avons obtenu une fonction d'appartenance du PFD_{avg} flou du SIS. Cette fonction a mis en évidence l'existence d'incertitude concernant le SIL du SIS.

Par ailleurs, l'utilisation des mesures de possibilité et de nécessité nous a permis de caractériser l'incertitude du SIL des SIS. La quantification de cette incertitude a souligné le fait que plusieurs SIL étaient possibles pour un même SIS quand nous prenons en compte l'imprécision des taux de défaillance des composants du SIS.

En outre, l'introduction de nouveaux facteurs d'importance a permis d'identifier les composants dont l'incertitude de la probabilité de défaillance contribue significativement à l'incertitude entachant la probabilité de défaillance du SIS complet, et donc son SIL. Nous avons apporté ainsi une aide à la décision aux fiabilistes, pour réduire d'une manière efficace les incertitudes dans la détermination des SIL.

L'application de l'approche proposée ainsi que les facteurs d'importance flous à un exemple applicatif a montré que les résultats obtenus sont très satisfaisants. En effet, nous avons caractérisé l'incertitude des résultats obtenus et nous avons réussi à réduire efficacement l'incertitude du SIL obtenu grâce aux études de sensibilité. Il faut souligner enfin que les résultats obtenus pour les facteurs flous à partir des facteurs d'importance de Birnbaum peuvent être étendus aux autres facteurs d'importance (facteurs d'importance de Lambert, facteurs d'importance de Vesely-Fussell, etc.).

4

Conception optimale des SIS

Nous présentons dans ce chapitre une nouvelle approche de conception optimale des SIS pour l'obtention d'un SIL exigé. Nous ramenons le problème de conception optimale à un problème de minimisation du coût global du SIS sous contrainte du SIL exigé qui est exprimé en fonction de la disponibilité moyenne A_{avg} du SIS. L'approche optimale est basée sur l'utilisation des réseaux de fiabilité et des algorithmes génétiques (AG).

Nous allons d'abord modéliser le SIS par un réseau de fiabilité. Nous utiliserons ensuite un AG pour l'optimisation de la structure du réseau de fiabilité et l'obtention de configurations optimales du SIS. Une application illustrera la méthodologie proposée et donnera les configurations optimales obtenues pour les différents SIL exigés sous forme de schémas de connexion, ce qui permettra de faciliter l'implémentation des SIS. En outre, nous proposerons un autre modèle d'aide à la conception optimale des SIS avec des composants à taux de défaillance imprécis modélisés par des nombres flous. Pour cette approche, nous donnerons également les configurations optimales obtenues pour les différents SIL exigés sous forme de schémas de connexion.

4.1 Introduction

Dans le cadre de la conception des SIS, les fiabilistes assignent aux SIS la réalisation des objectifs de sûreté de fonctionnement (PFD_{avg} ou disponibilité moyenne $A_{avg} = 1 - PFD_{avg}$) dès la phase d'expression du besoin en réduction de risque. Cet assignement a pour but, d'une part, d'aider le concepteur à rationaliser ses choix de composants et, d'autre part, de garantir à l'exploitant les objectifs de sûreté de fonctionnement exigés. Dans la figure 4.1, nous présentons un organigramme simplifié illustrant les étapes nécessaires pour la conception des SIS [Gruhn and Cheddie, 2006].

Si nous considérons le SIS à réaliser se décomposant successivement en sous-systèmes, ensembles et composants, nous sommes amenés à procéder, dès la phase de faisabilité, à la distribution des objectifs au niveau sous-systèmes, puis au niveau ensembles jusqu'au niveau composants. Cette distribution des objectifs intervient dans une logique descendante visant à allouer aux composants les exigences de sûreté de fonctionnement exprimées au niveau système [Haire *et al.*, 1985] et, dans une logique ascendante visant à consolider au niveau système les prévisions de données de sûreté de fonctionnement exprimées à des niveaux inférieurs.

La distribution des objectifs de sûreté de fonctionnement paraît donc indispensable pour :

- Pouvoir exprimer les exigences de sûreté de fonctionnement au niveau des composants.
- Vérifier que les principes retenus (architecture, concepts technologiques, etc.) sont compatibles avec les exigences de sûreté de fonctionnement émises par les fiabilistes.
- Procéder à la comparaison des concepts envisageables dans l'optique d'une optimisation du coût. Il est certain qu'une stratégie de conception, dont le souci est purement technique, permet d'atteindre les objectifs de sûreté de fonctionnement, mais elle le fait au détriment du coût de conception. Par contre, si la stratégie de conception cherche uniquement à réduire le coût de conception, le résultat est un nombre important de défaillances dangereuses et le non respect des objectifs.

Il est donc primordial de trouver une stratégie d'allocation des objectifs de sûreté de fonctionnement au niveau des composants du SIS qui permette d'établir le meilleur compromis entre le coût de conception du SIS et les objectifs de sûreté de fonctionnement qui sont exprimés sous forme du SIL requis.

Dans nos travaux, le problème de conception des SIS est ramené à un problème de minimisation du coût global du SIS sous contrainte du SIL exigé. Le SIL est exprimé en fonction de la disponibilité moyenne A_{avg} du SIS. La conception optimale du SIS se fait donc par le choix du nombre et des types de composants dans chaque sous-système du SIS qui garantissent un coût global minimal et le SIL exigé. Ce problème d'allocation fait partie des problèmes d'allocation de fiabilité et de redondance. Cependant, comme l'a souligné Elegbede [Elegbede, 2000], les problèmes d'allocation de fiabilité et de redondance ont traité principalement les systèmes à structure série parallèle et parallèle série. Cette restriction est due au fait que pour ce type de système, nous disposons d'une expression explicite de la fiabilité totale qui facilite le processus d'optimisation [Sallak *et al.*, 2006d], ce qui n'est pas le cas pour des structures quelconques. L'une des originalités de nos travaux est de proposer une nouvelle approche qui est appliquée à tout type de structure représentant un système cohérent et pouvant être modélisé par un réseau de fiabilité. Cette approche sera étendue pour prendre en compte les informations imprécises

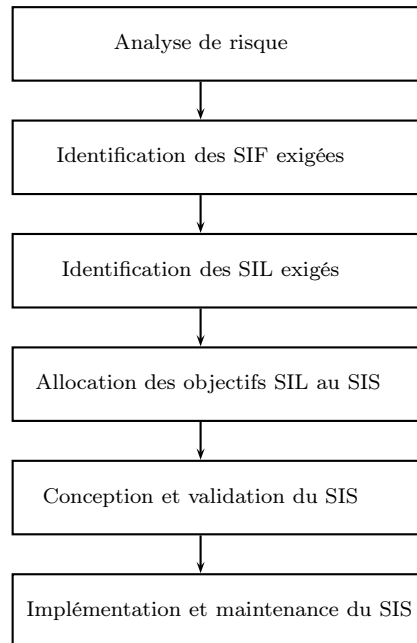


FIG. 4.1 – Organigramme simplifié de conception des SIS

sur les données de fiabilité des composants utilisés. Nous donnerons également les configurations optimales obtenues par l'approche floue.

Dans ce chapitre, nous traitons le problème de conception optimale des SIS pour le respect du SIL exigé. Nous commençons par donner un état de l'art des méthodes d'allocation de fiabilité et de redondance. Ensuite, deux méthodes de modélisation des SIS sont introduites : les BdF et les réseaux de fiabilité. Pour procéder à l'optimisation de la structure des SIS, nous allons d'abord les modéliser par des réseaux de fiabilité [Sallak *et al.*, 2007b]. Nous utiliserons ensuite un AG pour l'optimisation de la structure du réseau de fiabilité et l'obtention de configurations optimales des SIS. Pour prendre en compte les informations imprécises sur les données de fiabilité des composants utilisés, nous adapterons la méthodologie proposée aux probabilités floues de défaillances des composants. Une application illustrera la méthodologie proposée et donnera les configurations optimales obtenues pour les SIL exigés sous forme de schémas de connexion.

4.2 Etat de l'art de l'allocation de fiabilité et de redondance

Allain-Morin et Cloarec [Allain-Morin and Cloarec, 1994] divisent les méthodes d'allocation de fiabilité et de redondance en deux grandes catégories :

- Les méthodes basées sur les facteurs de pondération.
- Les méthodes basées sur l'optimisation.

4.2.1 Méthodes pondérées

Les méthodes pondérées sont des méthodes qui affectent des objectifs de fiabilité aux composants, sous-systèmes ou parties d'un système afin que la fiabilité totale du système puisse satisfaire à l'objectif exigé au départ (approche descendante). La même définition peut être utilisée pour l'affectation d'un objectif de disponibilité ou de maintenabilité. En anglais, ce type de méthode est appelé "reliability apportionment".

En 1957, la méthode AGREE a été introduite pour tenir compte de l'importance et de la complexité des sous-systèmes [AGREE, 1957; of Defence, 1998] pour l'allocation de fiabilité des équipements électroniques. La méthode Karmioli a été introduite en 1965 pour tenir compte de la complexité, l'état de l'art technologique, le profil opérationnel et la criticité du système par rapport aux objectifs [of Defence, 1998]. La méthode ARINC permet de prendre en compte le retour d'expérience concernant les composants [Alven, 1964]. En l'absence d'informations concernant l'importance des composants d'un système constitué de n sous-systèmes en série, la méthode d'égale allocation paraît la plus appropriée [of Defence, 1998]. Dans la suite, nous donnons une description sommaire de ces méthodes.

4.2.1.1 Méthode AGREE

La méthode AGREE a été développée en 1957 pour l'allocation de fiabilité des équipements électroniques [AGREE, 1957; of Defence, 1998]. L'affectation de fiabilité reflète l'importance et la complexité des sous-systèmes. La complexité d'un sous-système est représentée par le nombre de composants qu'il contient. L'importance d'un sous-système représente la probabilité de défaillance du système quand le sous-système est défaillant. L'inconvénient principal de cette méthode est qu'elle ne peut être appliquée qu'aux systèmes séries.

L'allocation AGREE de fiabilité d'un sous-système j est exprimée par la relation :

$$R(t_j) = 1 - \frac{1 - R^*(T)^{n_j/N}}{E_j} \quad (4.1)$$

Où :

$R^*(T)$ est l'objectif de fiabilité du système.

n_j est le nombre de composants dans le sous-système j .

E_j est le facteur d'importance de Birnbaum du sous-système j .

t_j est le nombre d'heures de fonctionnement du sous-système j en T heures de fonctionnement du système.

N est le nombre de composants du système.

4.2.1.2 Méthode Karmioli

La méthode Karmioli prend en compte l'influence des facteurs suivants [Allain-Morin and Cloarec, 1994] :

- La complexité (Cx) ;
- L'état de l'art technologique (A) ;
- Le profil opérationnel (O) ;

- La criticité (Cr).

Le produit des quatre facteurs représente le facteur effet n de chaque sous-système :

$$n = Cx.A.O.Cr \quad (4.2)$$

La fiabilité R_k affectée au sous-système k est alors donnée par la relation :

$$R_k = R^A.R^B \quad (4.3)$$

$$A = (n_1 + n_2 + \dots + n_m)/(2.N.n_k) \quad (4.4)$$

$$B = F_k/(2(F_1 + F_2 + \dots + F_m)) \quad (4.5)$$

$$N = \sum_{i=1}^m (n_1 + n_2 + \dots + n_m)/n_i \quad (4.6)$$

Où :

R est la fiabilité du système ;

n_k est le facteur effet du sous-système k ;

F_k est le nombre de composants dans le sous-système k ;

m est le nombre de sous-systèmes.

4.2.1.3 Méthode ARINC

La méthode ARINC est dédiée aux systèmes composés de sous-systèmes séries [Alven, 1964; of Defence, 1998]. La méthode suppose que :

- Les sous-systèmes ont des taux de défaillance constants.
- La défaillance de chaque sous-système entraîne la défaillance du système.
- La durée de mission de chaque sous-système est égale à la durée de mission du système.

L'allocation ARINC est réalisée en plusieurs étapes :

1. L'objectif est de choisir les λ_i^* tel que :

$$\sum_{i=1}^n \lambda_i^* \leq \lambda^* \quad (4.7)$$

Où :

λ_i^* est le taux de défaillance affecté au sous-système i .

λ^* est le taux de défaillance objectif du système.

n est le nombre de sous-systèmes.

2. Déterminer le taux de défaillance λ_i du sous-système i à partir du retour d'expérience des composants du sous-système i .

Méthodes	Avantages	Inconvénients
AGREE	- Prise en compte de l'importance et la complexité	- Applicable uniquement pour les systèmes séries
Egale allocation	- Simplicité	- Applicable uniquement pour les systèmes séries
ARINC	- Simplicité et objectivité	- Applicable uniquement pour les systèmes séries - Avoir un retour d'expérience des composants
Karmiol (facteur produit)	- Bien détaillée	- Subjectivité

TAB. 4.1 – Principales méthodes d'allocation pondérées

3. Affecter des facteurs pondérés w_i à chaque sous-système i selon les taux de défaillance déterminés dans (2) :

$$w_i = \frac{\lambda_i}{\sum_{i=1}^n \lambda_i} \quad (4.8)$$

4. Affecter les taux de défaillance requis pour chaque sous-système i :

$$\lambda_i^* = w_i \cdot \lambda^* \quad (4.9)$$

4.2.1.4 Méthode d'égle allocation

En l'absence d'informations concernant un système constitué de n sous-systèmes en série, la méthode d'égle allocation paraît la plus convenable [of Defence, 1998]. En effet, cette méthode affecte la même fiabilité pour chaque sous-système. L'inconvénient principal de cette méthode est qu'elle ne tient pas compte de la difficulté de réalisation des solutions proposées au niveau des sous-systèmes.

L'allocation de fiabilité d'un sous-système j est exprimée par la relation :

$$R_j^* = (R^*)^{1/N} \quad j = 1, \dots, N \quad (4.10)$$

Où :

R^* est l'objectif fiabilité du système.

R_j^* est l'objectif fiabilité du sous-système j .

N est le nombre de sous-systèmes.

Dans le tableau 4.1, nous récapitulons les avantages et les inconvénients des principales méthodes d'allocation pondérées que nous avons présenté.

4.2.2 Méthodes d'optimisation

L'optimisation de la fiabilité est un domaine qui a suscité beaucoup d'intérêt de la part des chercheurs et des industriels à partir des années 60. En effet, la prise en compte de la fiabilité des composants lors de la conception des systèmes permet d'avoir une idée précise du coût et de la fiabilité totale du système. Pour augmenter la fiabilité d'un système, plusieurs méthodes

peuvent être utilisées :

- Affectation de la fiabilité aux composants.
- Ajout de composants en redondance.
- Combinaison des deux techniques précédentes.
- Affectation de composants interchangeables.

Chaque technique décrite précédemment nécessite des ressources. La diversité des structures des systèmes, des contraintes sur les ressources et des options pour améliorer la fiabilité ont abouti à la construction et l'analyse de plusieurs modèles d'optimisation.

Kuo *et al.* [Kuo *et al.*, 2001] classent les méthodes d'optimisation de fiabilité selon :

- La structure des systèmes : séries parallèles, systèmes possédant une structure en réseaux, systèmes à redondance k parmi n , systèmes à structure indéfinie, etc.
- Le type de problème.
- Les techniques d'optimisation : heuristiques, méta heuristiques (algorithmes génétiques, recherche tabou, etc.), algorithmes exacts (programmation dynamique, méthodes du gradient, etc.) ou autres méthodes (décomposition, apportement flou, etc.).

L'objectif de l'allocation de redondance est de déterminer une configuration optimale du système qui maximise sa fiabilité sous certaines contraintes (coût, poids, volume, etc.). Les variables de décision représentent le nombre de composants à placer en parallèle dans chaque sous-système du système. En général, nous supposons que les composants sont identiques dans chaque sous-système. Les sous-systèmes étant montés en série.

Le problème d'allocation de redondance est exprimé sous la forme (**P1**) :

$$\text{maximiser} \quad R_s = f(x_1, \dots, x_n) \quad (4.11)$$

$$\text{sous les contraintes :} \quad g_i(x_1, \dots, x_n) \leq b_i, \quad i = 1, \dots, m; \quad (4.12)$$

$$l_j \leq x_j \leq u_j, \quad j = 1, \dots, n; \quad (4.13)$$

Où :

R_s est la fiabilité du système.

x_j est le nombre de composants en parallèle dans le sous-système j .

n est le nombre de sous-systèmes du système.

m est le nombre de ressources.

$g_i(x_1, \dots, x_n)$ est la consommation de la ressource i par le système.

l_j et u_j sont les bornes inférieure et supérieure de x_j .

Le problème d'allocation de redondance peut aussi s'exprimer sous la forme (**P2**) :

$$\text{maximiser} \quad R_s = f(x_1, \dots, x_n) \quad (4.14)$$

$$\text{sous les contraintes :} \quad \sum_{j=1}^n g_{i,j}(x_1, \dots, x_n) \leq b_i, \quad i = 1, \dots, m; \quad (4.15)$$

$$l_j \leq x_j \leq u_j, \quad j = 1, \dots, n; \quad (4.16)$$

$g_{i,j}(x_1, \dots, x_n)$ est la consommation de la ressource i par le sous-système j . Les fonctions contraintes $g_{i,j}$ dans (P2) doivent être séparables pour chaque sous-système du système.

Pour résoudre les problèmes (P1) et (P2), plusieurs approches ont été introduites. Kuo *et al.* [Kuo *et al.*, 2001] ont recensé la plupart de ces approches. Les problèmes d'allocation de redondance se sont surtout focalisés sur les structures séries parallèles. Chern a démontré que l'allocation de redondance des systèmes séries parallèles avec un seul objectif est un problème NP-difficile [Chern, 1992]. Elegbede [Elegbede, 2000] a montré que dans le cas des systèmes séries parallèles, si la fonction coût est convexe, la solution optimale est obtenue quand tous les composants dans chaque sous-système ont la même fiabilité.

Le problème d'allocation de redondance est un problème difficile à résoudre [Chern, 1992]. C'est pourquoi plusieurs heuristiques ont été développées pour trouver des solutions optimales à ce type de problème. Jusqu'en 1977, les heuristiques proposées pour résoudre le problème d'allocation de redondance étaient basées sur le même principe. A chaque itération du processus d'optimisation, les solutions proposées sont obtenues à partir des solutions précédentes en ajoutant un composant en parallèle dans le sous-système le moins fiable. En 1977, Nakagawa et Nakashima [Nakagawa and Nakashima, 1977] ont proposé une heuristique importante utilisant le même principe. Le choix de placement du composant est basé sur une étude de facteurs d'importance des sous-systèmes. Le processus d'optimisation s'arrête quand les contraintes sont violées. Kuo *et al.* ont ensuite étendu cette méthode aux systèmes complexes [Kuo *et al.*, 1978]. Li [Li, 1996] a aussi présenté en 1996 une heuristique importante basée sur la recherche des régions de l'ensemble des solutions possibles. Les résultats obtenus étaient très satisfaisants. En 1981, Nakagawa et Miyazaki [Nakagawa and Miyazaki, 1981] ont comparé les résultats obtenus (temps de calcul et précision des résultats) avec les méthodes heuristiques [Sharma and Venkateswaran, 1971; Nakagawa and Nakashima, 1977; Gopal *et al.*, 1978; Kuo *et al.*, 1978] d'allocation de redondance sous contraintes non linéaires. En 2003, Hsieh [Hsieh, 2003] a proposé le premier modèle qui permet de transformer le problème d'allocation de redondance en un problème de programmation linéaire standard. Cette méthode est basée sur une approximation linéaire de la fiabilité du système. Ramirez *et al.* [Ramirez-Marquez *et al.*, 2004] ont présenté une heuristique max-min. Le principe de cette méthode est de transformer le problème d'allocation de redondance en un problème de maximisation de la fiabilité minimale des sous-systèmes. Dans le cas de l'allocation de fiabilité et de redondance des systèmes parallèles séries, Elegbede *et al.* [Elegbede *et al.*, 2003] ont proposé l'algorithme ECAY pour les composants à fiabilités identiques. Dans le cas où les fiabilités peuvent n'être choisies que dans un ensemble fini de valeurs, Yalaoui *et al.* [Yalaoui *et al.*, 2005] ont développé la méthode YCC. Les méta heuristiques ont aussi été très utilisées ces dernières années et ont donné des résultats satisfaisants. Elles sont basées sur un raisonnement artificiel plutôt que des théorèmes mathématiques. Elles incluent principalement les algorithmes génétiques et le recuit simulé [Holland, 1975; Goldberg, 1994; Painton and Campbell, 1995; Michalewicz and Fogel, 2000; Deb *et al.*, 2006].

Dans la section suivante, nous allons nous focaliser sur l'optimisation par les algorithmes génétiques que nous avons utilisée. Nous allons d'abord justifier le choix des AG comme méthode d'optimisation dans nos travaux, puis nous allons introduire les principes de l'approche de conception optimale que nous avons proposée.

4.3 Algorithmes génétiques (AG)

4.3.1 Introduction

La résolution du problème de minimisation du coût global du SIS sous contrainte du SIL exprimé en fonction de la disponibilité du SIS n'est pas aisée. Pour résoudre ce type de problème, il existe diverses méthodes qui se divisent principalement en deux catégories : les méthodes déterministes et les méthodes stochastiques. Les techniques stochastiques tournent principalement autour des algorithmes stochastiques d'évolution de populations (AG, recuit simulé, etc.), qui sont des méthodes d'optimisation globale. Elles sont robustes, parallélisables et permettent de déterminer l'optimum global d'une fonctionnelle. Leur inconvénient majeur réside dans le nombre important d'évaluations nécessaires pour obtenir l'optimum recherché. Les méthodes déterministes de type gradient présentent en revanche l'avantage de converger rapidement vers un optimum. Cependant, elles ne sont pas aussi robustes que les techniques stochastiques et, n'assurent pas que l'optimum déterminé est un optimum global. En outre, elles dépendent beaucoup du point de départ de recherche de l'extremum [Rao, 1996; Li and Aggarwal, 2000; Shia and Shenb, 2007].

Développés par Holland [Holland, 1975] à l'université du Michigan, les AG sont des méthodes d'optimisation de fonctions. Ces algorithmes s'inspirent de l'évolution génétique des espèces. Schématiquement, ils copient de façon extrêmement simplifiée certains comportements des populations naturelles. Ainsi, ces techniques reposent toutes sur l'évolution d'une population de solutions qui sous l'action de règles précises optimisent un comportement donné, exprimé sous forme d'une fonction, dite fonction sélective (fitness) [Holland, 1975; Goldberg, 1994; Lutton, 1999; Michalewicz and Fogel, 2000; Deb *et al.*, 2006].

4.3.2 Nomenclature des AG

Les AG étant basés sur des phénomènes biologiques, il convient de rappeler au préalable les termes utilisés dans les AG et leur signification biologique (cf. tableau 4.2).

4.3.3 Avantages des AG

L'utilisation des AG présente plusieurs avantages :

- L'utilisation unique de l'évaluation de la fonction objectif sans se soucier de sa nature. En effet, nous n'avons besoin d'aucune propriété particulière sur la fonction à optimiser (continuité, dérivabilité, convexité, etc.), ce qui lui donne plus de souplesse et un large domaine d'application.
- Génération d'une forme de parallélisme en travaillant sur plusieurs points en même temps (population de taille N) au lieu d'un seul itéré dans les algorithmes classiques.
- L'utilisation des règles de transition probabilistes (probabilités de croisement et de mutation), contrairement aux algorithmes déterministes où la transition entre deux itérations successives est imposée par la structure et la nature de l'algorithme. Cette utilisation permet, dans certaines situations, d'éviter des optimums locaux et de se diriger vers un optimum global.

Terme	AG	Signification biologique
Gène	trait, caractéristique	Une unité d'information génétique transmise par un individu à sa descendance.
locus	Position dans une chaîne	L'emplacement d'un gène dans son chromosome.
Allèle	Différentes versions d'un même gène	Une des différentes formes que peut prendre un gène, les allèles occupent le même locus.
Chromosome	Chaîne	Une structure contenant les gènes.
Génotype	structure	L'ensemble des allèles d'un individu portés par l'ADN d'une cellule vivante.
Phénotype	Ensemble de paramètres ou une structure décodé.	Aspect physique et physiologique observable de l'individu obtenu à partir de son génotype.
Epistasie	Non linéarité	Terme utilisé pour définir les relations entre deux gènes distincts, lorsque la présence d'un gène empêche la présence d'un autre gène non allèle.

TAB. 4.2 – Résumé de la terminologie utilisée en AG

4.3.4 Inconvénients des AG

L'utilisation des AG présente aussi quelques inconvénients :

- Si l'AG nous garantit de trouver une bonne solution approchante, il ne nous garantit pas de trouver la valeur optimale.
- Suivant le problème, la résolution par l'AG peut être coûteuse en temps.

4.3.5 Concepts de base

Un AG est un algorithme itératif de recherche d'optimum, il manipule une population formée de points candidats appelés chromosomes. La taille constante de la population entraîne un phénomène de compétition entre les chromosomes. Chaque chromosome représente le codage d'une solution potentielle au problème à résoudre. Il est constitué d'un ensemble d'éléments appelés gènes, pouvant prendre plusieurs valeurs appartenant à un alphabet non forcément numérique [Ludovic, 1994]. A chaque itération, appelée génération, une nouvelle population est créée avec le même nombre de chromosomes. Cette génération consiste en des chromosomes mieux adaptés à leur environnement tel qu'il est représenté par la fonction sélective. Au fur et à mesure des générations, les chromosomes vont tendre vers un optimum de la fonction sélective. La création d'une nouvelle population à partir de la précédente se fait par application des opérateurs génétiques que sont : la sélection, le croisement et la mutation [Goldberg, 1994]. Ces opérateurs sont stochastiques. La sélection des meilleurs chromosomes est la première opération dans un AG. Au cours de cette opération l'algorithme sélectionne les éléments pertinents qui optimisent mieux la fonction. Le croisement permet de générer deux chromosomes nouveaux "enfants" à partir de deux chromosomes sélectionnés "parents", tandis que la mutation réalise l'inversion d'un ou plusieurs gènes d'un chromosome. La figure 4.2 illustre la séquence des opérations qui interviennent dans un AG de base [Ouarzizi, 1997].

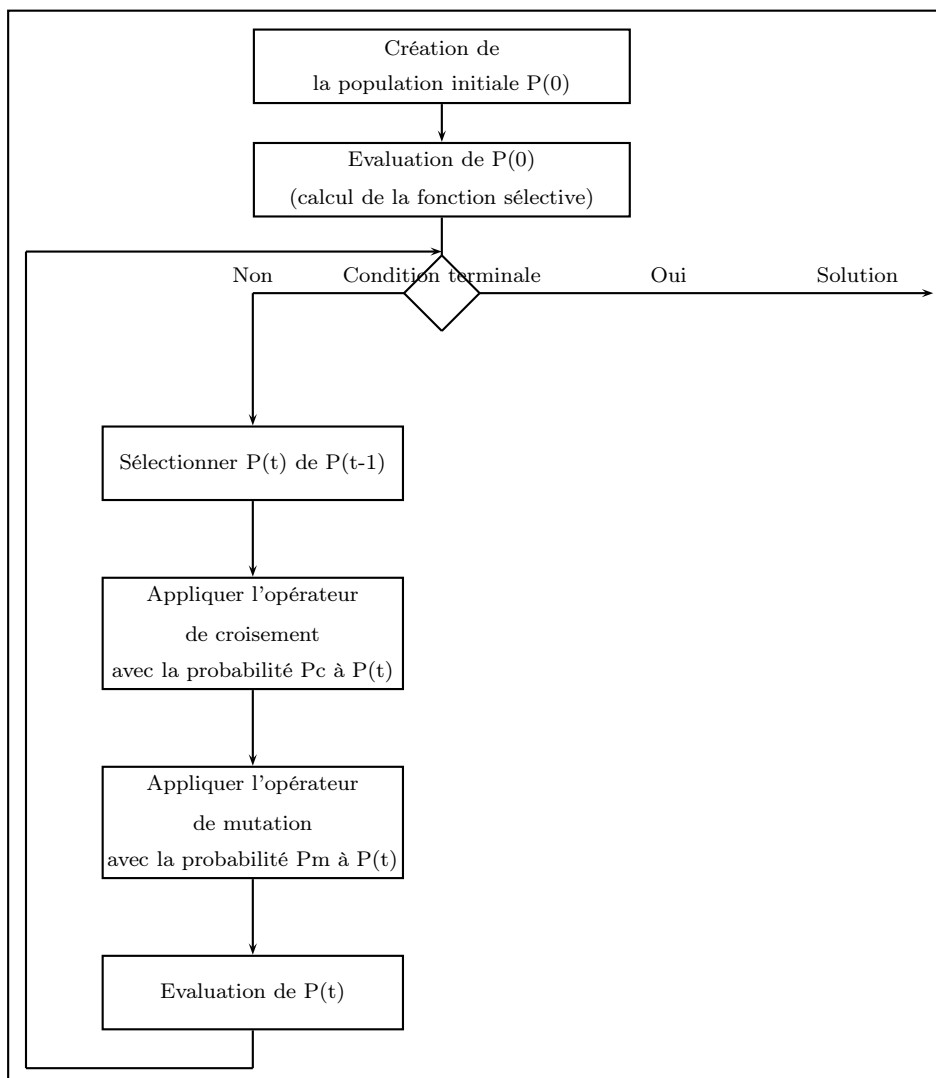


FIG. 4.2 – Algorithme génétique de base

4.3.6 Opérateurs des AG

Les opérateurs jouent un rôle prépondérant dans la possible réussite d'un AG. Nous en dénombrons trois principaux : la sélection, le croisement et la mutation. Si le principe de chacun de ces opérateurs est facilement compréhensible, il est toutefois difficile d'expliquer l'importance isolée de chacun de ces opérateurs dans la réussite de l'AG. Cela tient pour partie au fait que chacun de ces opérateurs agit selon divers critères qui lui sont propres (valeur sélective des individus, probabilité d'activation de l'opérateur, etc.). Nous détaillons par la suite les éléments de l'AG que nous avons utilisé.

4.3.6.1 Codage des solutions

Dans un AG, nous ne travaillons pas directement avec les données du problème mais avec une représentation de celles-ci appelées codage. La forme codée d'une solution est une chaîne qu'on appellera chromosome. Ce chromosome est à son tour constitué d'éléments qu'on appellera gènes. Dans une population, nous parlerons indifféremment de chromosomes et d'individus. Le choix du codage dépend de la spécificité du problème traité. Il conditionne fortement l'efficacité de l'AG. Dans la littérature, nous trouvons trois types de codage :

- Numérique si l'alphabet est constitué de chiffres.
- Symbolique si l'alphabet est un ensemble de lettres alphabétiques ou de symboles.
- Alphanumérique si nous utilisons un alphabet combinant les lettres et les chiffres.

4.3.6.2 Population initiale

Une fois le codage choisi, une population initiale formée de solutions admissibles (chromosomes) du problème doit être déterminée. Plusieurs mécanismes de génération de la population initiale sont utilisés dans la littérature. La population initiale peut être générée aléatoirement, par duplication et évolution ou en s'appuyant sur une heuristique.

4.3.6.3 Taille des populations

Il n'y a pas de standardisation quant au choix de la taille des populations. Des tailles de population faibles augmenteront la vitesse de convergence de l'algorithme, mais aussi le risque de convergence prématurée vers des solutions non optimales. Des tailles de population trop grandes risquent au contraire de ralentir fortement la progression de l'algorithme. Cette importance de la taille est essentiellement due à la notion de parallélisme implicite qui implique que le nombre d'itérations de l'AG soit au moins proportionnelle au cube du nombre d'individus. La taille fréquemment utilisée est 30 [Ouarzizi, 1997; Michalewicz and Fogel, 2000].

4.3.6.4 Sélection

La sélection a pour objectif d'identifier les individus qui doivent se reproduire. Cet opérateur ne crée pas de nouveaux individus mais identifie les individus sur la base de leur fonction d'adaptation. Les individus les mieux adaptés sont sélectionnés alors que les moins bien adaptés sont écartés. Ceci permet de donner aux individus dont la valeur de la fonction sélective est plus grande une probabilité plus élevée de contribuer à la génération suivante. Vladimir [Vladimir, 2003] a souligné le fait que lorsqu'un AG est utilisé pour maximiser une fonction objectif, dans certains cas, c'est le processus de sélection qui assure la convergence vers un optimum global.

Il existe plusieurs types de sélection [Goldberg, 1994]. Une des méthodes les plus connues est la roue de loterie biaisée (roulette wheel) de Goldberg [Goldberg, 1994].

4.3.6.5 Croisement

Le croisement a pour but d'enrichir la diversité de la population en manipulant la structure des chromosomes. Classiquement, les croisements sont envisagés avec deux parents et génèrent deux enfants. Dans la littérature, plusieurs techniques de croisement sont utilisées dont les principaux sont le croisement barycentrique et le croisement à un ou plusieurs points [Goldberg, 1994; Michalewicz and Fogel, 2000]. Le croisement à deux points est l'un des croisements les plus utilisés. Il consiste à couper le chromosome en deux points choisis aléatoirement et recombinaison les morceaux en croisant les chromosomes. Une probabilité de croisement P_C signifie que, quand deux parents sont candidats à la reproduction, nous tirons un réel x aléatoirement selon une loi uniforme sur l'intervalle $[0,1]$, si x est inférieur à P_C , nous croisons alors les parents. Les valeurs généralement admises sont comprises entre 0,5 et 0,9 [Goldberg, 1994].

4.3.6.6 Mutation

L'opérateur de mutation permet d'introduire un facteur aléatoire dans les solutions générées, et d'élargir ainsi l'espace des solutions explorées (Goldberg, 1994; Michalewicz 1996; Koza, 1992) pour éviter à l'AG de s'enliser dans des optima locaux. Pour les codages en nombre réels, la mutation consiste à modifier légèrement quelques gènes des chromosomes. En général, on choisit une faible probabilité de mutation. Cette probabilité de mutation représente la fréquence à laquelle les gènes d'un chromosome sont mutés. Par exemple, une mutation très utilisée est de tirer aléatoirement un seul gène dans le chromosome et à le remplacer par une valeur aléatoire avec une probabilité de mutation P_m . Cet opérateur est donc d'une grande importance. La probabilité de mutation P_m est généralement faible puisqu'un taux élevé risque de conduire à une solution sous-optimale. La mutation est traditionnellement considérée comme un opérateur marginal bien qu'elle confère aux AG la propriété d'ergodicité (*i.e.* tous les points de l'espace de recherche peuvent être atteints).

4.3.6.7 Remplacement

Cet opérateur est le plus simple, son travail consiste à réintroduire les descendants obtenus par application successive des opérateurs de sélection, de croisement et de mutation (la population P') dans la population de leurs parents (la population P).

Ce faisant, ils vont remplacer une certaine proportion de ceux-ci, proportion pouvant bien sûr être choisie. Le rapport entre le nombre d'individus nouveaux allant être introduits dans la population P et le nombre d'individus de cette population est connu sous le nom de generation gap.

Nous trouvons essentiellement deux méthodes de remplacement différentes :

- Le remplacement stationnaire : dans ce cas, les enfants remplacent automatiquement les parents sans tenir compte de leurs performances respectives, et le nombre d'individus de la population ne varie pas tout au long du cycle d'évolution simulé, ce qui implique donc d'initialiser la population initiale avec un nombre suffisant d'individus.
- Le remplacement élitiste : dans ce cas, nous gardons au moins l'individu possédant les meilleures performances d'une génération à la suivante. En général, nous pouvons partir

du principe qu'un nouvel individu (enfant) ne prend place au sein de la population que s'il remplit le critère d'être plus performant que le moins performant des individus de la population précédente. Donc les enfants d'une génération ne remplaceront pas nécessairement leurs parents comme dans le remplacement stationnaire et par conséquent la taille de la population n'est pas figée au cours du temps. Ce type de stratégie améliore les performances des AG dans certains cas. Mais présente aussi l'inconvénient d'augmenter le taux de convergence prématurément.

4.3.7 Choix des paramètres des AG

Comme pour toute heuristique d'optimisation, l'efficacité d'un algorithme génétique dépend du choix de ses paramètres (probabilités liées aux opérateurs d'évolution, taille des populations, etc.) qui gouvernent l'exploration des solutions et des conditions initiales. Il n'y a pas de règle générale pour le choix de ces paramètres. Pour qu'un AG ait des bonnes performances, il doit être exécuté plusieurs fois avec différentes tailles de population, probabilités de croisement et de mutation afin de trouver l'ensemble des paramètres qui conviennent le plus à l'utilisateur.

4.4 Modélisation des SIS

La conception optimale d'un SIS nécessite sa modélisation. Plusieurs approches sont utilisées selon le système à étudier et les critères d'analyse de défaillance fixés. Dans cette section, nous allons présenter les deux approches que nous avons utilisées pour modéliser les SIS : les blocs diagrammes de fiabilité (BdF) et les réseaux de fiabilité.

4.4.1 Blocs diagrammes de Fiabilité (BdF)

La méthode des blocs diagrammes de fiabilité est une des premières méthodes à avoir été utilisée pour analyser les systèmes et permettre des calculs de fiabilité [Kleinerman and Weiss, 1954; Blanton, 1957; Villemeur, 1998]. Elle est aussi appelée la Méthode du Diagramme de Succès (MDS). C'est une représentation de la logique de fonctionnement des systèmes car elle est souvent proche de leur schéma fonctionnel. Cette méthode est basée sur l'utilisation de blocs pour représenter les composants, les sous-systèmes ou les fonctions. La modélisation consiste à rechercher les liens existant entre ces blocs. En outre, dans cette modélisation, les systèmes doivent vérifier les deux hypothèses suivantes :

- Hypothèse d'états binaires ;
- Indépendance des états de fonctionnement ou de défaillance des composants.

Les BdF sont ainsi utilisés dans de nombreux domaines industriels pour les systèmes non réparables [Dhillon and Yang, 1997; Levitin, 2007], mais ils peuvent également, sous certaines conditions, être utilisés pour les calculs de fiabilité de systèmes réparables [Guo and Yang, 2006].

4.4.1.1 Diagramme série

La panne de l'un des éléments E_i du système entraîne la panne du système (cf. figure 4.3). Si nous désignons par R_S la fiabilité du système et R_i la fiabilité du composant E_i , alors la fiabilité du système est donnée par :

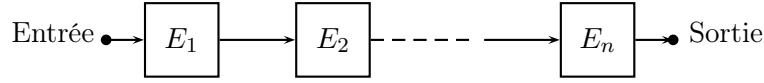


FIG. 4.3 – BdF d'un système série

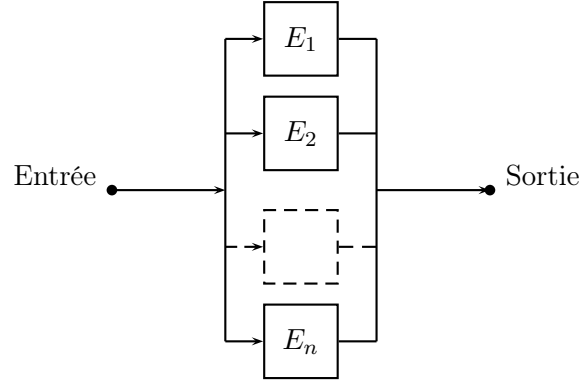


FIG. 4.4 – BdF d'un système parallèle

$$R_S = \prod_{i=1}^n R_i \quad (4.17)$$

où n est le nombre de composants du système.

4.4.1.2 Diagramme parallèle

La panne de tous les éléments E_i du système entraîne la panne du système (cf. figure 4.4). Si un seul des éléments fonctionne alors il conduit au fonctionnement du système. Dans ce cas, la fiabilité du système est donnée par :

$$R_S = 1 - \prod_{i=1}^n (1 - R_i) \quad (4.18)$$

4.4.1.3 Diagramme en redondance k/n

Les systèmes à n composants qui fonctionnent si et seulement si au moins k de leurs composants fonctionnent, sont appelés des systèmes k parmi n (k/n) (cf. figure 4.5).

On ne dispose pas d'une expression générale de la fiabilité d'un système k/n . Néanmoins, dans le cas où tous les composants du système ont la même fiabilité R , la fiabilité totale du système est donnée par :

$$R_S = \sum_{i=k}^n C_n^i R^i (1 - R)^{n-i} \quad (4.19)$$

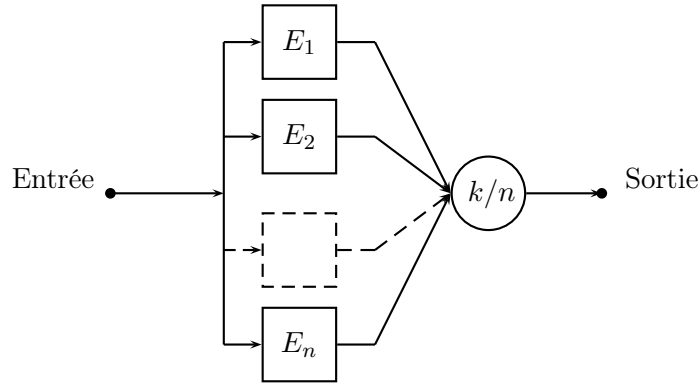


FIG. 4.5 – BdF d'un système k/n

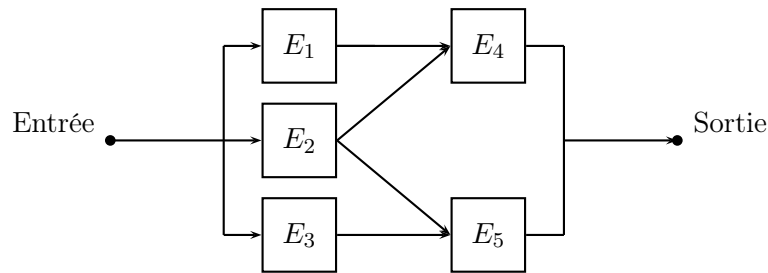


FIG. 4.6 – Schéma de connexion d'un système complexe

avec :
$$C_n^i = \frac{n!}{i!(n-i)!}$$

4.4.1.4 Diagramme complexe

Les diagrammes complexes sont des diagrammes qui ne peuvent pas être réduits à des combinaisons series et/ou parallèles [Villemeur, 1987].

Pour traiter les diagrammes complexes (cf. figure 4.6), nous pouvons utiliser :

- les liens minimaux et les coupes minimales.
- la fonction de structure.
- le théorème des probabilités totales.
- la table de vérité.

4.4.1.5 Liens minimaux et coupes minimales

Dans un schéma de connexion, un lien est un ensemble de blocs dont le fonctionnement assure le succès de la mission du système. Un lien minimal est une des plus petites combinaisons de blocs qui lorsqu'ils sont en fonction permettent d'assurer la fonction requise pour le système (il ne contient pas d'autres liens).

Si L_i est un lien minimal du système alors la fiabilité totale du système est donnée par :

$$R = P\left[\sum_{i=1}^l L_i\right] \quad (4.20)$$

où l est le nombre de liens minimaux du système.

Les coupes minimales représentent les plus petites combinaisons de défaillances des blocs qui compromettent la fonction requise pour le système.

Si C_i est une coupe minimale du système alors la fiabilité totale du système est donnée par :

$$R = 1 - P\left[\sum_{i=1}^c C_i\right] \quad (4.21)$$

où c est le nombre de coupes minimales du système.

4.4.2 Réseaux de fiabilité

Dans cette section, nous verrons comment décrire la structure d'un système à l'aide d'un réseau de fiabilité. Nous allons d'abord introduire le vocabulaire de base de la théorie des graphes nécessaire pour comprendre les réseaux de fiabilité. Ensuite, nous allons présenter les réseaux de fiabilité tels qu'ils ont été défini par Kaufmann *et al.* [Kaufmann *et al.*, 1975]. Enfin, nous donnerons le réseau de fiabilité du SIS correspondant au schéma de connexion qui a été présenté dans la section précédente.

4.4.2.1 Rappel sur la théorie des graphes

De manière générale, un graphe permet de représenter la structure et les connexions d'un ensemble complexe en exprimant les relations entre ses éléments : réseaux de communication, réseaux routiers, circuits électriques, etc. [Berge, 1958; Cogis and Robert, 2003]. Les graphes constituent donc une méthode qui permet de modéliser une grande variété de problèmes en se ramenant à l'étude de sommets et d'arcs. Dans cette partie, nous nous contentons de donner les éléments nécessaires pour introduire aisément la théorie des réseaux de fiabilité.

Considérons un ensemble fini S et le produit $S \times S$. Soit U un sous ensemble de $S \times S$. le couple :

$$G = (S, U)$$

est appelé un multigraphe (ou encore graphe r -appliqué) où r est le nombre maximal des arcs ayant même extrémité initiale et même extrémité terminale. Les éléments de S sont appelés les sommets du graphe ; les éléments de U , qui sont des couples de sommets, sont appelés les arcs du graphe.

La figure 4.7 donne une représentation d'un graphe. Les sommets sont représentés par des points, et les arcs par des lignes continues joignant deux sommets, et portant une flèche.

Dans ce graphe, l'ensemble des sommets est :

$$S = \{A, B, C, D\} \quad (4.22)$$

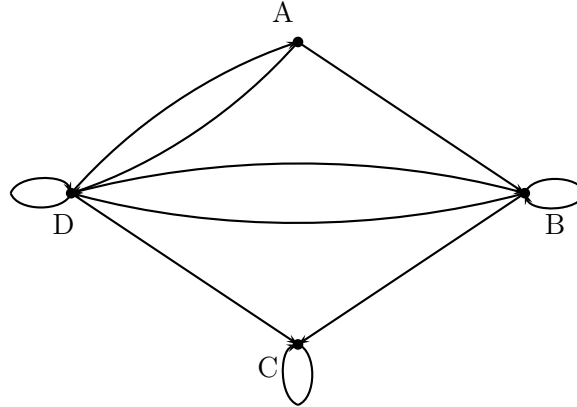


FIG. 4.7 – Dessin d'un graphe

$$\begin{pmatrix} 0 & 1 & 0 & 1 \\ 0 & 1 & 1 & 1 \\ 0 & 0 & 1 & 0 \\ 1 & 1 & 1 & 1 \end{pmatrix}$$

FIG. 4.8 – Matrice booléenne

et l'ensemble des arcs est :

$$U = \{(AB), (AD), (BB), (BC), (BD), (CC), (DA), (DB), (DC), (DD)\}. \quad (4.23)$$

Étant donné un arc tel que (A, B) , A est appelé l'extrémité initiale de l'arc et B son extrémité terminale. Un arc tel que (B, B) est appelé une boucle. Le sommet B est dit suivant de A si (A, B) est un arc ; par exemple, dans le graphe de la figure 4.7, C n'est pas un suivant de A .

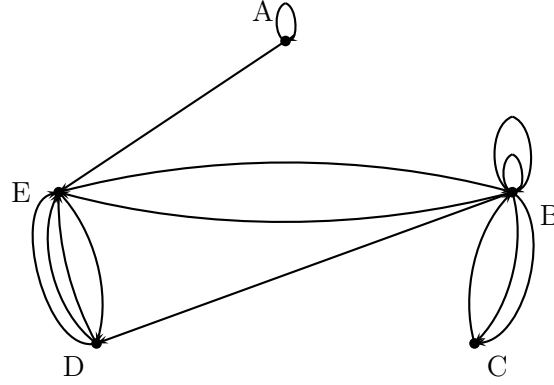
Un graphe peut aussi être décrit par sa matrice booléenne (appelée aussi matrice de connexion), c'est-à-dire par une matrice carrée dont les lignes et les colonnes correspondent aux sommets du graphe. Les éléments valent 1 ou 0 suivant que le couple de sommets correspondant appartient ou non à U comme le montre la figure 4.8.

Pour mettre en mémoire un graphe sur un ordinateur, nous utilisons le plus souvent le dictionnaire des suivants ou le dictionnaire des précédents. Comme son nom l'indique, le dictionnaire des suivants donne, pour chaque sommet du graphe, la liste des suivants. Cela revient à décrire la matrice booléenne ligne par ligne (cf. figure 4.9). Le dictionnaire des précédents donne au contraire, pour chaque sommet, la liste des sommets qui sont l'extrémité initiale d'un arc ayant pour extrémité terminale le sommet considéré. Il correspond aux colonnes de la matrice booléenne.

Nous allons maintenant définir sur les graphes r -appliqués quelques concepts que nous aurons à utiliser dans la suite de ce document.

A	B, D
B	B, C, D
C	C
D	A, B, C, D

FIG. 4.9 – Dictionnaire des suivants

FIG. 4.10 – Graphe r -appliqué

Étant donné un graphe $G = (S, U)$, est appelé graphe partiel de G un graphe $G_p = (S, U')$ tel que $U' \subset U$. Nous obtenons donc un graphe partiel de G en supprimant certains arcs. Ainsi, le graphe de la figure 4.11 est un graphe partiel du graphe représenté sur la figure 4.10.

Un chemin est une suite d'arcs :

$$\mu = (u_1, u_2, \dots, u_l) \quad \text{avec} \quad u_i \in U, \quad i = 1, \dots, l. \quad (4.24)$$

tel que l'extrémité terminale de chaque arc u_i coïncide avec l'extrémité initiale de l'arc suivant u_{i+1} . Un chemin tel que l'extrémité terminale de l'arc u_l coïncide avec l'extrémité initiale de l'arc u_1 est appelé un circuit. On appelle longueur d'un chemin le nombre d'arcs qu'il comporte. Un chemin est dit élémentaire s'il ne passe pas deux fois par le même sommet.

Considérons un sous-ensemble de sommets $S_1 \subset S$. On appelle coupe du graphe G relative au sous-ensemble S_1 l'ensemble $\omega(S_1)$ des arcs dont l'extrémité initiale n'appartient pas à S_1 , et dont l'extrémité finale appartient à S_1 . Par exemple, si l'on pose $S_1 = \{A, E\}$ dans le graphe de la figure 4.11, la coupe correspondante est :

$$C1(\{A, E\}) = \{(BE), (DE)_1, (DE)_2, (DE)_3\} \quad (4.25)$$

La théorie des graphes est un très vaste domaine en constante évolution, tant du point de vue des recherches fondamentales que de celui des applications. Pour plus de détails, il est intéressant de consulter [Berge, 1958; Kaufmann, 1968; Kaufmann, 1972; Kaufmann *et al.*, 1975;

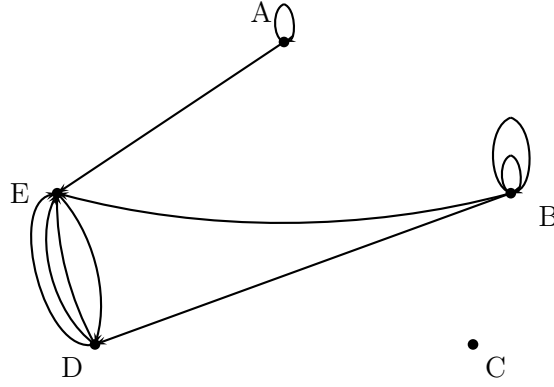


FIG. 4.11 – Graphe partiel

Cogis and Robert, 2003]. Dans la suite, nous allons introduire les réseaux de fiabilité qui sont basés sur l'utilisation de la théorie des graphes.

4.4.2.2 Concepts de base des réseaux de fiabilité

Un réseau de fiabilité R défini sur un ensemble $e = \{e_1, e_2, \dots, e_r\}$ de composants est constitué par :

- Un graphe r -appliqué $G = (S, U)$ sans boucles, dans lequel deux sommets $O \in S$ et $Z \in S$ sont distingués et appelés respectivement "origine" et "extrémité".
- Une application $\Delta : U \rightarrow e$ telle que :

$$\Omega(u_j) = (S_i, S_k), \quad \Omega(u_{j'}) = (S_i, S_k) \Rightarrow \Delta(u_j) \neq \Delta(u_{j'}) \quad (4.26)$$

où Ω est l'application qui fait correspondre à chaque arc le couple de ses extrémités.

L'application Δ fait correspondre à chaque arc du graphe un composant. Plusieurs arcs peuvent correspondre à un même composant (il se peut qu'un arc ne corresponde à aucun composant). La figure 4.12 donne un exemple de réseau de fiabilité où :

$$e = \{e_1, e_2, e_3, e_4\} \quad (4.27)$$

$$S = \{O, Z, A, B, C\} \quad (4.28)$$

$$U = \{(OA)_{e_2}, (OA)_{e_3}, (OB), (AB), (AZ), (BC), (BZ), (CB), (ZB)\}. \quad (4.29)$$

L'application Δ est indiquée par les e_i affectés aux arcs.

4.4.2.3 Lien d'un réseau de fiabilité

A tout sous-ensemble de composants $e_i \subset e$, nous pouvons faire correspondre le graphe partiel $G_p(e_i)$ du graphe G , obtenu en ne conservant que les arcs de G auxquels correspond un composant appartenant à e_i :

$$G_p(e_i) = (S, U_p(e_i)) \quad (4.30)$$

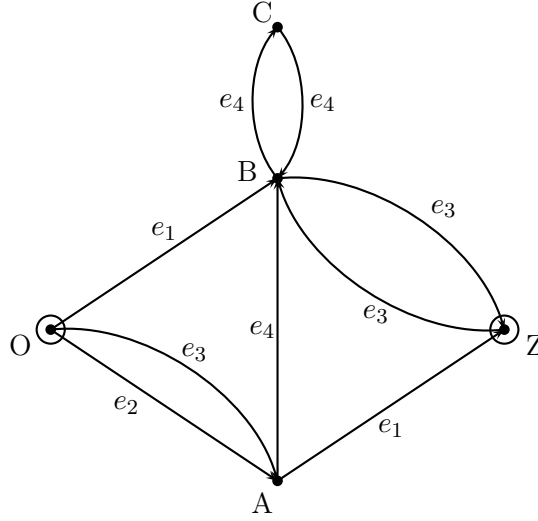


FIG. 4.12 – Réseau de fiabilité

avec :

$$U_p(e_i) = \{u \in U \mid \Delta(u) \in e_i\}. \quad (4.31)$$

Nous appellerons "lien" d'un réseau R un sous-ensemble de composants $a \subset e$ tel qu'il existe dans le graphe $G_p(a)$ un chemin de O à Z .

Par exemple, $\{e_1, e_2, e_3\}$ est un lien du réseau de la figure 4.12, comme le montre la figure 4.13, qui représente le graphe partiel correspondant. A tout lien correspondent un ou plusieurs chemins du graphe. Inversement, à un chemin $\mu = (u_1, u_2, \dots, u_l)$ de O à Z correspond le lien a formé des images des arcs $u_1, u_2, \dots, u_l$ par l'application Δ . Par exemple, au chemin $\mu = ((OA)_{e3}, AB, BZ)$ correspond le lien $a = \{e_3, e_4\}$.

4.4.2.4 Coupe d'un réseau de fiabilité

Nous appellerons "coupe" d'un réseau R un sous-ensemble de composants $b \subset e$ tel que le sous-ensemble d'arcs $U_p(b)$ contienne une coupe du graphe G relative à un sous-ensemble de sommets incluant Z et n'incluant pas O .

Par exemple, $\{e_1, e_2, e_3\}$ est une coupe du réseau (en même temps qu'un lien) ; le sous-ensemble complémentaire $\{e_4\}$ n'est ni une coupe, ni un lien. A toute coupe b du réseau correspondent donc une ou plusieurs coupes du graphe incluses dans $U_p(b)$. Dans l'exemple ci-dessus, l'ensemble d'arcs $\{AZ, BZ\}$ est une coupe du graphe qui n'inclut pas les arcs OB , $(OA)_{e2}$ et $(OA)_{e3}$.

4.4.2.5 Lien minimal et coupe minimale d'un réseau de fiabilité

Un lien a est minimal si aucun sous-ensemble $a' \subset a$ n'est un lien du réseau. Une coupe b est minimale si aucun sous-ensemble $b' \subset b$ n'est une coupe du réseau.

Dans la figure 4.12, $\{e_1, e_2, e_3\}$ n'est pas un lien minimal ; $\{e_1, e_3\}$ en est un, c'est également une coupe minimale. $\{e_1, e_4\}$ est un autre exemple de coupe minimale.

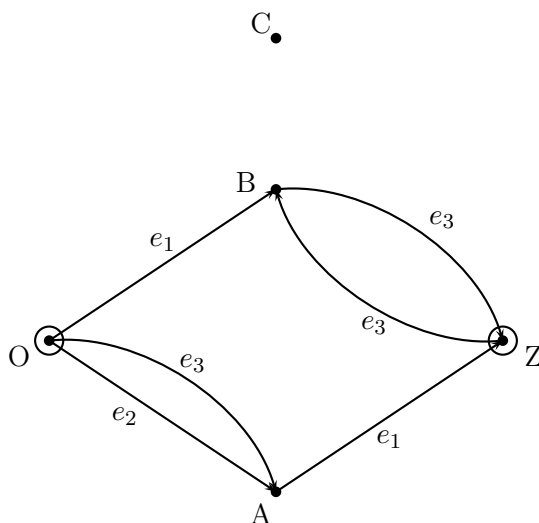


FIG. 4.13 – Réseau de fiabilité partiel

4.4.2.6 Remarque

Un réseau de fiabilité n'est pas un schéma de connexion physique des éléments mais seulement un modèle formelle du comportement dysfonctionnel; si nous connaissons toutes les coupes ou tous les liens alors nous avons deux réseaux de fiabilité équivalents.

4.5 Conception optimale

4.5.1 Introduction

Dans cette partie, nous proposons une méthodologie générale pour résoudre le problème de conception optimale des SIS. Ce problème a été ramené à un problème de minimisation du coût global du SIS sous contrainte du SIL exigé. Ensuite, nous appliquerons l'approche proposée au réservoir sous pression défini dans le document technique [ISA-TR84.00.02-2002, 2002]. Enfin, nous présenterons les différentes architectures optimales obtenues et nous discuterons de leur implémentation.

4.5.2 Pourquoi les AG ?

La première raison est que la mise en oeuvre des AG ne nécessite aucune hypothèse ou information sur le système optimisé (pas de calcul de gradient par exemple), ce qui correspond à notre problématique où nous devons optimiser une fonction qui n'est pas connue à priori et qui n'est pas continue.

La deuxième raison est que les AG permettent un équilibre entre exploitation et exploration [Beasley *et al.*, 1993]. Le mot équilibre est justifié par le fait que les deux procédures sont antagonistes. L'exploitation d'une direction de recherche consiste essentiellement à encourager

l'apparition de ses représentants dans la population tandis que l'exploration plaide en faveur de nouvelles directions de recherche. L'AG apporte une solution à ce dilemme en allouant un nombre exponentiel croissant à la meilleure direction observée.

La troisième raison est que les AG ont montré de très bonnes performances dans la résolution de problèmes d'allocation de fiabilité et de redondance sur lesquels peu d'informations sont disponibles ou pour lesquels il faut considérer de multiples critères d'optimisation [Kuo *et al.*, 2001].

4.5.3 Pourquoi les réseaux de fiabilité ?

La première raison est que les réseaux de fiabilité permettent de modéliser aisément tout type de structure quelque soit la complexité des systèmes étudiés.

La deuxième raison est que suite à l'utilisation des réseaux de fiabilité, nous pouvons utiliser les matrices booléennes de connexion dans le processus d'optimisation. En effet, il est plus aisé de procéder à une optimisation sur une matrice qui est une représentation de la structure du système.

4.5.4 Méthodologie générale

Dans cette partie, nous présentons les différentes étapes de notre méthodologie de conception optimale des SIS.

1. Définition de l'objectif SIL.

Lors de l'étude d'un process, quand il y a un danger éventuel, il faut procéder à une analyse de risques et dangers. Les risques existants y sont décrits et les mesures existantes et supplémentaires prises pour réduire ces risques y sont définies. Le risque subsistant doit toujours rester à un niveau tolérable. Après l'affectation des tâches de sécurité aux niveaux de protection, nous nous intéressons aux spécifications relatives aux exigences de sécurité pour les SIS. Ces exigences étant exprimées sous la forme d'un SIL. L'objectif est de concevoir le SIS pour qu'il puisse atteindre le SIL exigé.

2. Proposition de différents types de composants pour la conception du SIS.

Pour concevoir le SIS, les fiabilistes disposent de différents types de composants disponibles sur le marché. Pour chaque composant, nous disposons de son taux de défaillance et de son coût. Un SIS est généralement constitué de trois parties :

- Partie Capteurs ;
- Partie Unités de traitement ;
- Partie Actionneurs.

Chaque partie pouvant contenir plusieurs composants de différents types placés souvent en parallèle, notre objectif est de choisir les composants de chaque partie et leur connexion qui permettent d'atteindre le SIL exigé avec un coût global minimal.

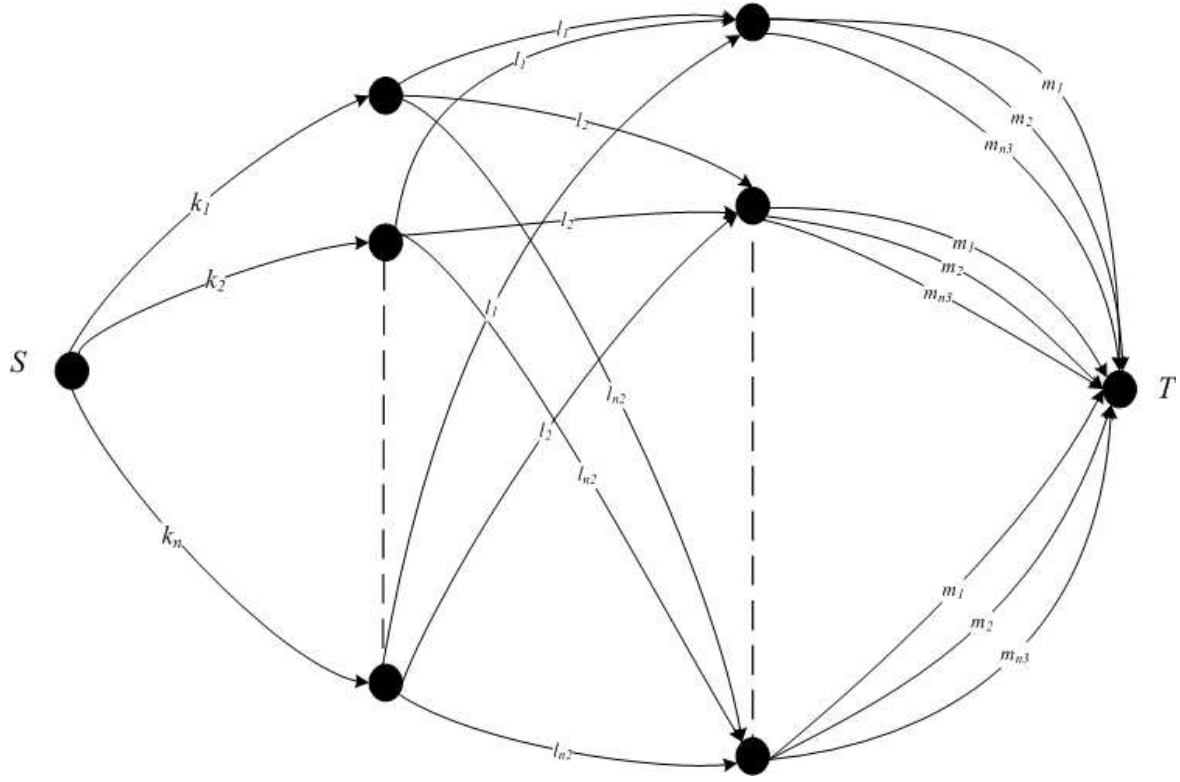


FIG. 4.14 – Réseau de fiabilité général d'un SIS

3. Modélisation de la structure générale du SIS par un réseau de fiabilité.

Nous modélisons le SIS par un réseau de fiabilité défini sur l'ensemble e de ses composants :

$$e = \{k_1, k_2, \dots, k_{n1}, l_1, l_2, \dots, l_{n2}, m_1, m_2, \dots, m_{n3}\} \quad (4.32)$$

Le composant k_i de la partie Capteurs représente le capteur de type i . Conformément à la structure générale des SIS, le capteur k_i ne peut être connecté qu'aux composants l_i de la partie adjacente Unités de traitement. De même, les composants l_i ne peuvent être connectés qu'aux composants m_i de la partie adjacente Actionneurs. Les deux sommets initial et final du réseau sont S et T (cf. figure ??). Nous faisons correspondre à chaque arc du réseau un composant. Les extrémités de chaque arc représentent les points de connexions. Plusieurs arcs peuvent correspondre à un même composant du SIS (cf. remarque 4.4.2.6). Les liens du réseau de fiabilité représentent les chemins de succès de S jusqu'à T . Par exemple $\{k_1, l_1, m_1\}$ et $\{k_2, l_1, m_1\}$ sont des liens du réseaux. Ils sont aussi des liens minimaux.

4. Codage des solutions et génération de la population initiale (configuration initiale du SIS).

Nous cherchons les configurations optimales du réseau de fiabilité du SIS (*i.e.*, les connexions entre les différents composants du SIS) qui permettent d'atteindre les objectifs souhaités. Les chromosomes x (les solutions potentielles) sont des chaînes constituées de plusieurs gènes x_{ij} . Chaque chromosome représente une configuration particulière du réseau de fiabilité du SIS. Chaque gène x_{ij} est un nombre binaire (cf. figure 4.15). Il prend la valeur 1

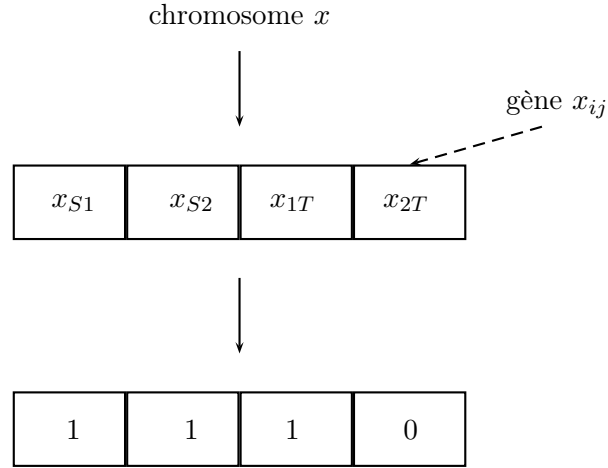


FIG. 4.15 – Codage du réseau de fiabilité

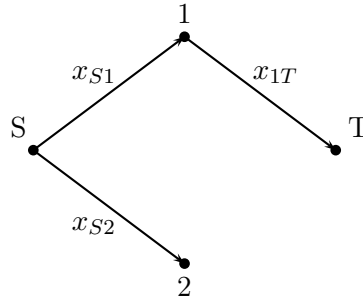


FIG. 4.16 – Exemple d'un réseau de fiabilité

s'il existe un arc entre le nœud i et le nœud j , et 0 sinon. Le nombre de gènes (la taille du chromosome) est déterminé par le nombre de connexions possibles entre les différents composants du SIS. Chaque composant ne pouvant être connecté qu'à des composants du sous-système adjacent.

La figure 4.16 montre un réseau de fiabilité composé de quatre nœuds et trois arcs. S est le nœud de départ et T est le nœud d'arrivée. Par exemple le gène x_{2T} est égale à 0 car il n'existe pas de connexion entre le nœud 2 et le nœud T . Nous rappelons que le nœud S ne peut pas être connecté au nœud T car chaque nœud ne peut être connecté qu'à des nœuds adjacents. Le nombre de connexions possibles dans ce réseau étant 4, nous déduisons que la taille du chromosome x est 4 (cf. figure 4.15). Le chromosome x est défini par :

$$x = [x_{S1} \quad x_{S2} \quad x_{1T} \quad x_{2T}] = [1 \quad 1 \quad 1 \quad 0] \quad (4.33)$$

La figure 4.17 donne des exemples de réseaux de fiabilité ainsi que les schémas de connexion correspondants.

Le choix de la taille de la population initiale de chromosomes conditionne fortement la

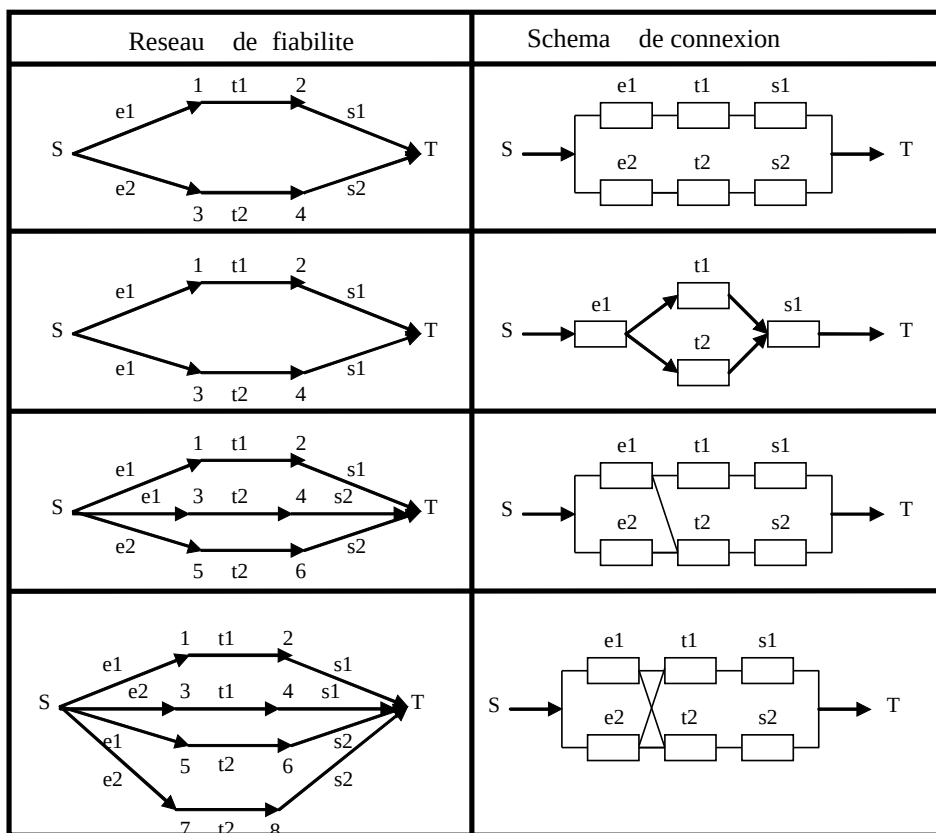


FIG. 4.17 – Correspondance entre réseaux de fiabilité et schémas de connexion

rapidité et l'efficacité de l'AG. Dans notre cas où nous ne connaissons pas à priori le type de configuration que nous pouvons obtenir, nous avons choisi que la population initiale soit répartie uniformément sur tout le domaine de recherche en faisant des tirages uniformes. La taille de la population utilisée est fixée à 20. le nombre de génération de la population est fixé à 300.

5. Génération des liens minimaux du réseau de fiabilité du SIS.

Les liens l_i ou chemin de succès du SIS sont les ensembles de composants dont le fonctionnement assure le succès de la mission du SIS. Dans le réseau de fiabilité, les liens représentent les chemins de succès qui mènent du nœud de départ S vers le nœud d'arrivée T . Les liens sont minimaux quand ils ne contiennent pas d'autres liens du réseau. Ce qui est intrinsèquement vérifié dans notre cas. En effet, avec les hypothèses précédentes, tous les liens sont minimaux.

6. Calcul de $A_{avg}(x) = 1 - PFD_{avg}(x)$ et du coût $C(x)$ du SIS.

La disponibilité moyenne A_{avg} du SIS est calculée à partir des liens l_i minimaux du réseau de fiabilité. La disponibilité instantanée est définie par l'équation :

$$A(t) = \sum_{i=1}^n P_{l_i(t)} \quad (4.34)$$

Où :

$P_{l_i(t)}$ est la disponibilité instantanée du lien minimal i .
 n est le nombre de liens minimaux du réseau de fiabilité.

En procédant à une disjonction des différents termes qui figurent dans les liens minimaux, nous obtenons :

$$A(t) = P_{l_1(t)} + \bar{P}_{l_1(t)}P_{l_2(t)} + \dots + \bar{P}_{l_1(t)}\bar{P}_{l_2(t)}\dots\bar{P}_{l_{n-1}(t)}P_{l_n(t)} \quad (4.35)$$

Le calcul de disponibilité moyenne A_{avg} du SIS est donné par :

$$A_{avg} = \frac{1}{T} \int_0^T A(t) dt \quad (4.36)$$

Où T est la période d'étude.

Le coût total du SIS $C(x)$ est la somme des coûts des composants impliqués dans l'architecture optimale du SIS.

7. Calcul de la fonction objectif (fitness).

Pour assurer les contraintes imposées par le respect du SIL, nous devons assurer que la disponibilité moyenne A_{avg} du SIS reste au dessus d'une valeur minimale A_{min} et en dessous d'une valeur maximale A_{max} . Pour cela, nous utilisons un terme de correction de la

disponibilité de type erreur quadratique moyenne. En outre, nous ciblons en premier lieu le respect du SIL exigé. Dans l'équation 4.37, le paramètre δ permet de minimiser le coût uniquement si le SIL exigé est atteint et c_{max} représente le coût maximal des composants du SIS.

La fonction objectif est donc la somme du coût total des composants et une erreur quadratique moyenne de la disponibilité moyenne du SIS. La fonction objectif est donnée par :

$$f(x) = C(x) + \delta \left(\frac{(A_{avg}(x) - A_{min})^2}{2} \right) \quad (4.37)$$

$$\delta = 1 \quad \text{si} \quad A_{avg} \leq A_{min}, \quad \text{et} \quad 0 \quad \text{sinon.} \quad (4.38)$$

8. Choix des paramètres de l'AG

Comme pour toute heuristique d'optimisation, l'efficacité d'un algorithme génétique dépend du choix de ses paramètres (probabilités liées aux opérateurs d'évolution, taille des populations, etc.) qui gouvernent l'exploration des solutions, et des conditions initiales. Il n'y a pas de règle générale pour le choix de ces paramètres. Pour qu'un AG ait des bonnes performances, Kimbrough [Kim,] a suggéré de l'exécuter plusieurs fois avec différentes tailles de population, probabilités de croisement et de mutation afin de trouver l'ensemble des paramètres qui conviennent le plus à l'utilisateur. C'est cette méthode que nous avons choisie.

9. Sélection.

La sélection permet d'identifier statistiquement les meilleurs chromosomes de la population et d'éliminer les moins bons. Nous avons choisi la méthode de sélection du tournoi binaire stochastique qui est aujourd'hui la technique de sélection la plus populaire en raison de sa simplicité et de son efficacité. A chaque fois qu'il faut sélectionner un chromosome, nous tirons aléatoirement deux chromosomes de la population, sans tenir compte de la valeur de leur fonction d'adaptation, et nous choisissons le meilleur chromosome parmi les deux. L'opération est évidemment répétée autant de fois que nous avons de chromosomes à sélectionner pour la reproduction.

10. Croisement.

Nous avons choisi le croisement à deux points. Il consiste à couper le chromosome en deux points choisis aléatoirement et recombinaison les morceaux en croisant les chromosomes. Une probabilité de croisement P_c signifie que, quand deux parents sont candidats à la reproduction, nous tirons un réel x aléatoirement selon une loi uniforme sur l'intervalle $[0, 1]$. Si x est inférieur à P_c , nous croisons alors les parents. La probabilité de croisement est fixée à 0.60.

11. Mutation.

La mutation choisie est une mutation très utilisée dont le principe est un tirage aléatoire d'un seul gène dans le chromosome et son remplacement par une valeur aléatoire avec une probabilité de mutation $P_m = 0.05$.

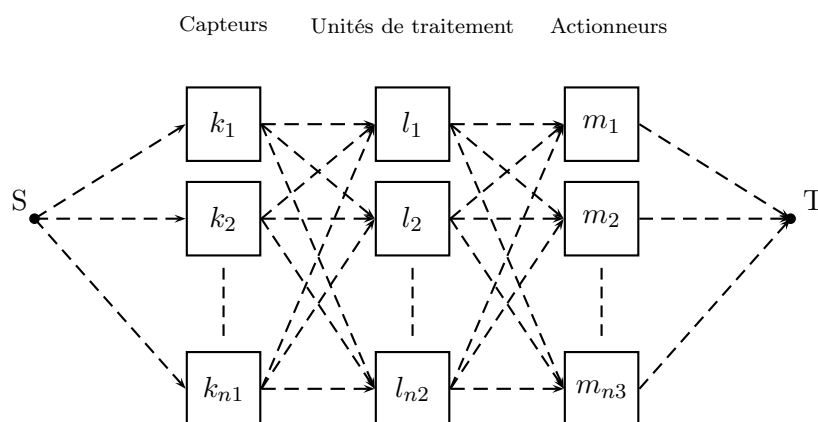


FIG. 4.18 – Schéma de connexion général d'un SIS

12. Obtention d'une solution optimale.

En définissant un critère d'arrêt de l'AG, nous obtenons une solution optimale qui représente une configuration optimale du réseau de fiabilité du SIS. Pour obtenir une représentation simple de la structure du SIS obtenu, nous transformons le réseau de fiabilité du SIS en un schéma de connexion.

13. Modélisation de la structure optimale du SIS par un schéma de connexion.

Le schéma de connexion du SIS est constitué de trois sous-systèmes qui représentent les trois parties du SIS (cf. figure 4.18). Chaque bloc d'un sous-système représente un composant. Les arcs représentent les liens qui existent entre les différents composants. En outre, le SIS doit vérifier les trois hypothèses suivantes :

- Hypothèse d'état binaires.
- Indépendance des états de fonctionnement des composants.
- Les composants sont non réparables.

L'avantage du schéma de connexion est de donner une représentation graphique du SIS facile à interpréter et qui permet une réalisation aisée de l'architecture optimale du SIS obtenue après l'optimisation. L'organigramme de la figure 4.19 résume les principales étapes de conception optimale d'un SIS pour lequel un SIL est exigé.

4.5.5 Cas d'application

Dans cette partie, nous proposons d'appliquer la méthodologie proposée à la conception optimale d'un SIS dont un SIL i est exigé. Nous reprenons l'application présentée dans le chapitre 3 concernant le réservoir sous pression défini dans le document technique ISA-TR84.00.02-2002

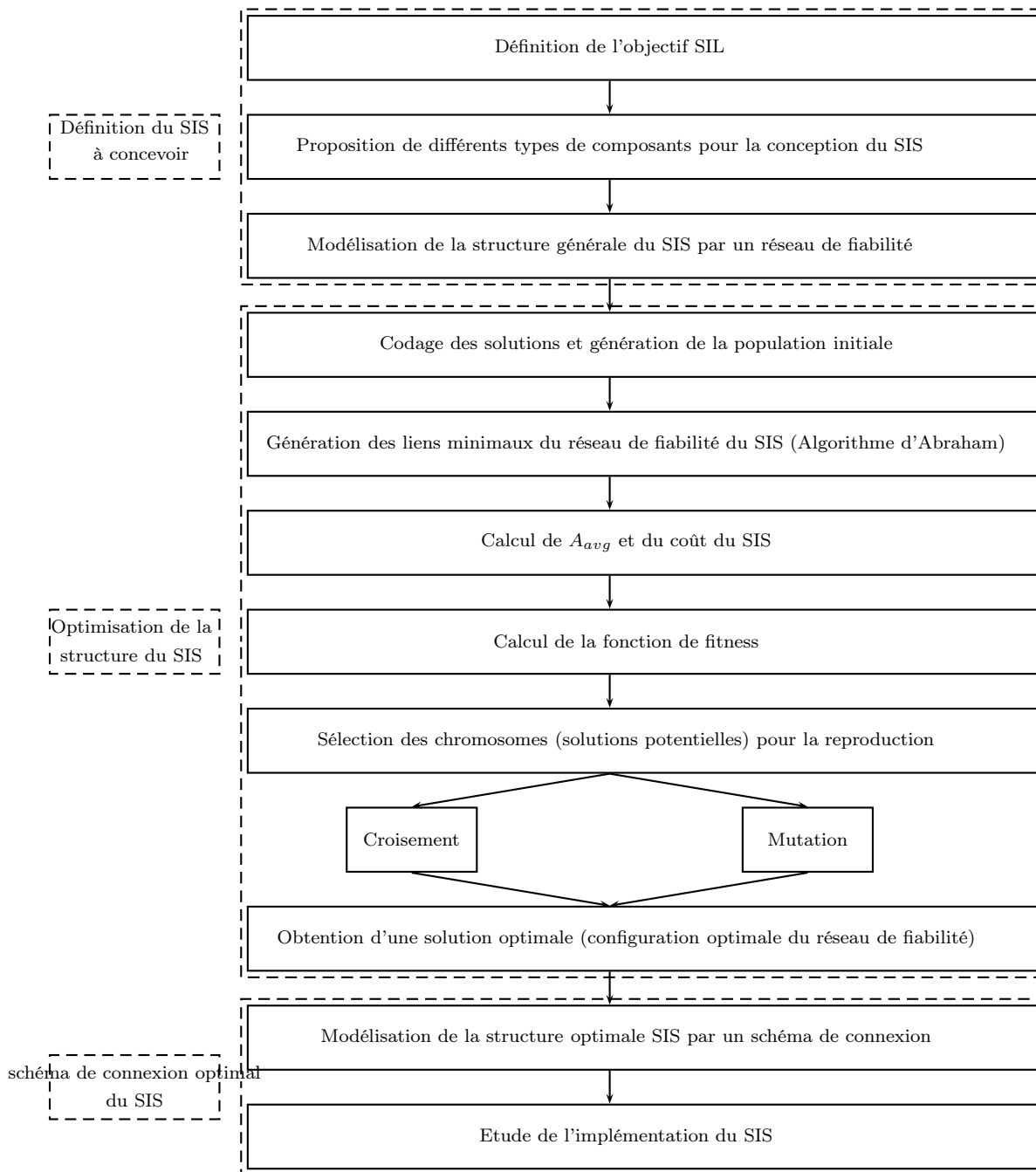


FIG. 4.19 – Organigramme de conception optimale d'un SIS

Composants du SIS	Sous systèmes					
	Capteurs		Unités de traitements		Actionneurs	
	c_1 (unités)	p_1	c_2 (unités)	p_2	c_3 (unités)	p_3
Type 1	2100	0.039	1400	0.09	2500	0.1
Type 2	1500	0.07	2100	0.05	3500	0.06
Type 3	2000	0.03	1200	0.07	4100	0.04

TAB. 4.3 – Coûts et probabilités de défaillance des composants du SIS

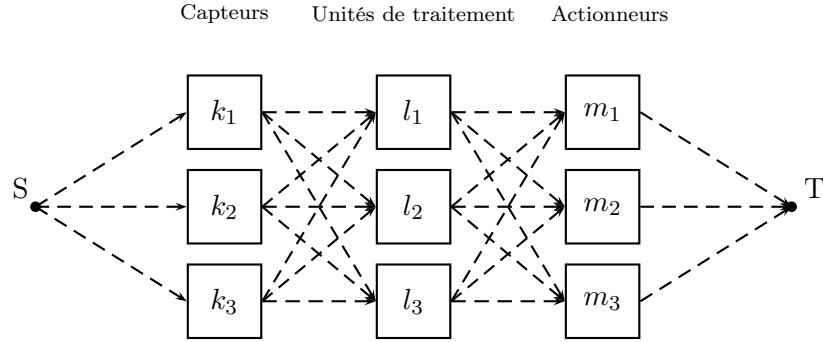


FIG. 4.20 – Schéma de connexion du SIS à concevoir

[ISA-TR84.00.02-2002, 2002].

Notre objectif est de concevoir le SIS dont le SIL i a été imposé au concepteur avec un coût total minimal. En conséquence, il faut choisir les composants de chaque sous système du SIS, ainsi que les connexions entre ces composants qui permettent d'obtenir le SIL exigé avec un coût de conception minimal. Notons que chaque composant ne peut être utilisé qu'une seule fois dans chaque partie du SIS.

4.5.5.1 Formulation mathématique du problème

Ce problème de conception peut être ramené à un problème de minimisation du coût global du SIS sous une contrainte sur la disponibilité moyenne A_{avg} du SIS. La contrainte sur le SIL exigé est transformée en une contrainte sur la disponibilité moyenne du SIS selon le tableau 1.3. Le coût global du SIS est la somme des coûts de ses composants. En outre, nous supposons qu'il y a trois types de composants disponibles sur le marché pour chaque sous-système du SIS. Les probabilités de défaillance et les coûts des composants du SIS sont donnés dans le tableau 4.3. Les coûts des composants représentent les prix moyens de ces composants sur le marché en Euros. La structure générale du SIS ainsi que les connexions éventuelles qui peuvent exister entre les différents composants sont données dans la figure 4.20.

Ainsi, le problème de conception revient à trouver la configuration (ou les configurations) optimale(s) (*i.e.*, le choix des composants du SIS et les connexions entre ces composants) qui permet de :

Statistiques	$C(x)$	$A_{avg}(x)$	SIL
Moyenne	8600	0.913037	1
Maximum	9400	0.927941	1
Minimum	8100	0.906056	1
Écart type	661	0.017	-

TAB. 4.4 – Résultats pour l’obtention d’un SIL1

- Minimiser $C(x)$
- sous la contrainte : $A_{min} \leq A_{avg}(x)$

4.5.5.2 Optimisation par l’AG

Dans cette application, suite à de nombreux essais, les paramètres de l’AG que nous avons utilisé sont :

- Taille de la population : 20.
- Nombre de génération de la population : 300.
- Probabilité de croisement : 0.6.
- Probabilité de mutation : 0.05.

4.5.6 Résultats et analyse

Ayant défini toutes les données nécessaires (probabilités de défaillance et coûts des composants disponibles pour chaque sous-système du SIS) pour traiter le problème d’optimisation, l’étape suivante est l’implémentation de la procédure de conception optimale décrite précédemment (cf. section 4.5) pour différents SIL exigés dans l’exemple. En ce qui concerne l’optimisation, nous avons utilisé le logiciel MATLAB 7.1 et un processeur Intel(R) Pentium de 3.0 GHz. Les différents programmes d’optimisation ont été écrit sous MATLAB. Nous avons répété la procédure d’optimisation de l’AG 100 fois pour chaque SIL exigé.

4.5.6.1 Exigence SIL1

Le tableau 4.4 présente les statistiques relatives aux résultats d’optimisation obtenus. Nous y donnons en particulier les statistiques relatives au coût $C(x)$ et à la disponibilité $A_{avg}(x)$ des architectures optimales représentées par le vecteur x . Nous remarquons clairement que pour tous les essais, la contrainte du SIL1 a été respecté, puisque la disponibilité moyenne minimale obtenue est $A_{min} = 0.906056$ qui est supérieure au seuil $A_{avg} = 0.90$ exigé pour un SIL1. Les figures 4.21 à 4.26 présentent les schémas de connexion et les réseaux de fiabilité des meilleures architectures optimales obtenues. L’intérêt de présenter plusieurs solutions optimales est d’offrir plus de choix aux concepteurs selon d’autres critères non spécifiés dans le cahier des charges. Nous remarquons que les structures obtenues sont des structures séries parallèles. Cependant, dans la suite du mémoire, nous obtiendrons des systèmes ayant une structure plus complexe dès lors que le problème est plus contraint. Cela valide l’application de notre approche à tout type de structure complexe.

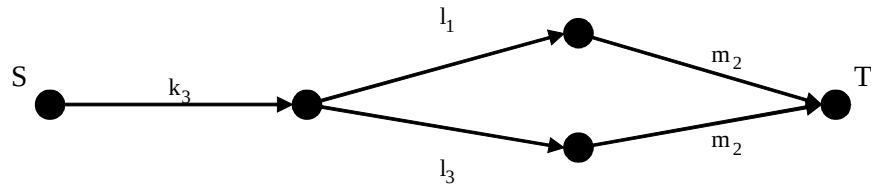


FIG. 4.21 – Architecture obtenue (SIL1 : $A_{avg} = 0.906056, C = 8100$) : Réseau de fiabilité

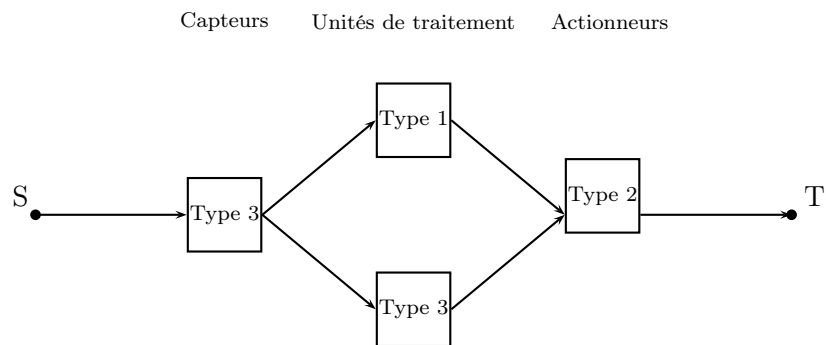


FIG. 4.22 – Architecture obtenue (SIL1 : $A_{avg} = 0.906056, C = 8100$) : schéma de connexion

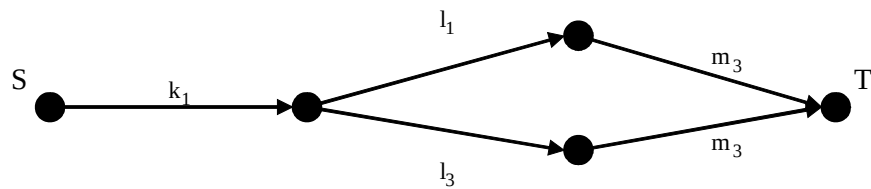


FIG. 4.23 – Architecture obtenue (SIL1 : $A_{avg} = 0.916748, C = 8800$) : Réseau de fiabilité

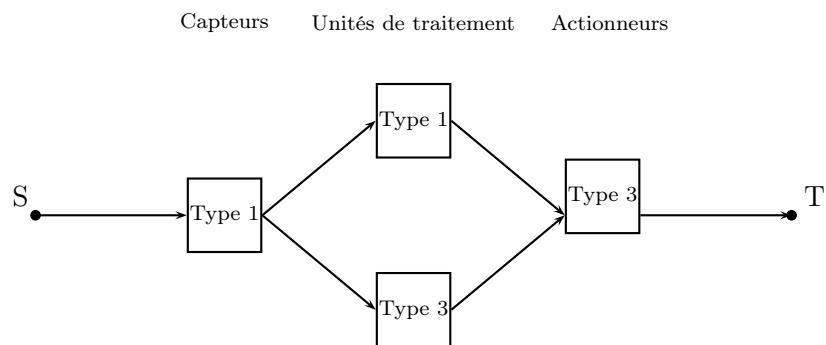


FIG. 4.24 – Architecture obtenue (SIL1 : $A_{avg} = 0.916748, C = 8800$) : schéma de connexion

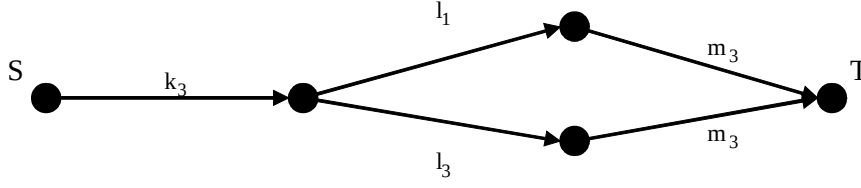


FIG. 4.25 – Architecture obtenue (SIL1 : $A_{avg} = 0.925333$, $C = 8700$) : Réseau de fiabilité

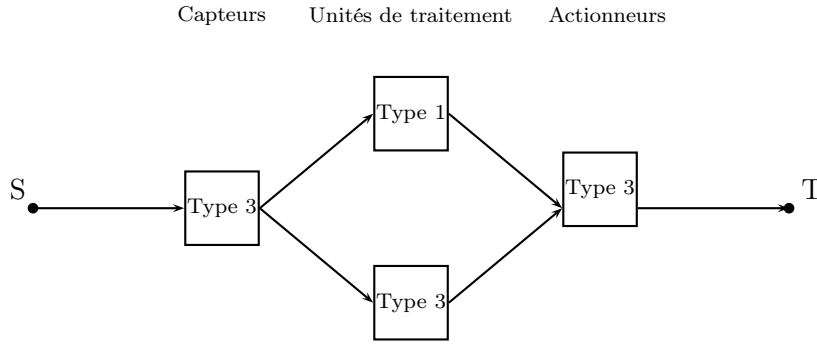


FIG. 4.26 – Architecture obtenue (SIL1 : $A_{avg} = 0.925333$, $C = 8700$) : schéma de connexion

4.5.6.2 Exigence SIL2

Le tableau 4.5 présente les statistiques relatives aux résultats d'optimisation obtenus. Là encore, nous remarquons que pour tous les essais, la contrainte du SIL2 a été respecté puisque la disponibilité moyenne minimale obtenue est $A_{min} = 0.991565$ qui est supérieure au seuil $A_{avg} = 0.990$ exigé pour un SIL2. Les figures 4.27 à 4.32 présentent les schémas de connexion et les réseaux de fiabilité des meilleures architectures optimales obtenues. Nous remarquons que nous obtenons des structures complexes. Ainsi, cet exemple illustre l'intérêt de notre approche pour la conception optimale des systèmes à structure quelconque.

4.5.6.3 Exigences SIL3 et SIL4

Les composants dont nous disposons pour la conception du SIS ne permettront pas l'obtention d'un SIL3 ou d'un SIL4 parce que la disponibilité moyenne maximale que nous pouvons obtenir est 0.9837. Cela ne permettra pas d'atteindre $A_{avg} = 0.999$ nécessaire au moins pour un SIL3. La disponibilité moyenne maximale du SIS est obtenue en utilisant tous les composants

Statistiques	$C(x)$	$A_{avg}(x)$	SIL
Moyenne	18600	0.996039	2
Maximum	20400	0.998104	2
Minimum	16300	0.991565	2
Écart type	1400	0.001376	-

TAB. 4.5 – Résultats pour l'obtention d'un SIL 2

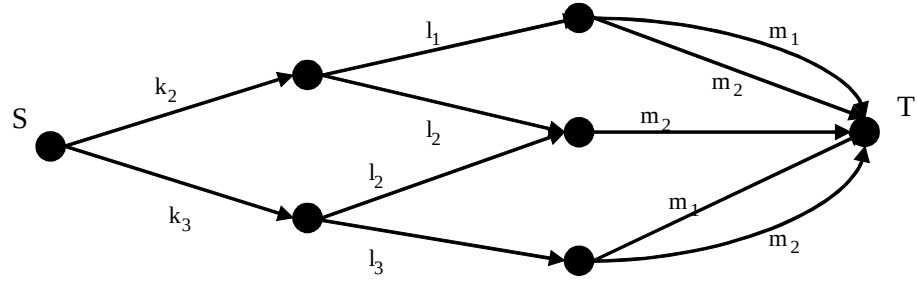


FIG. 4.27 – Architecture obtenue (SIL2 : $A_{avg} = 0.990326$, $C = 14200$) : Réseau de fiabilité

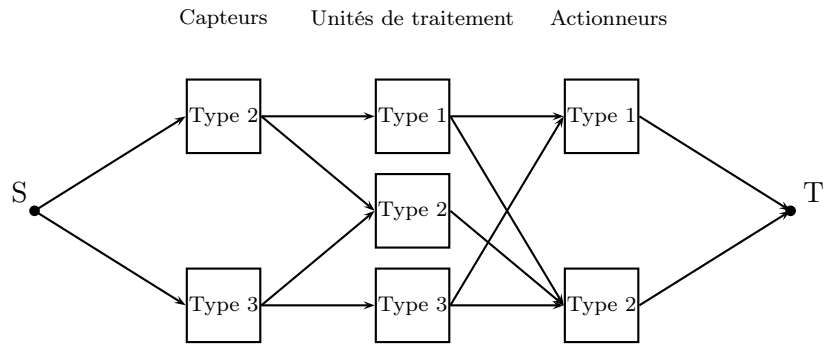


FIG. 4.28 – Architecture obtenue (SIL2 : $A_{avg} = 0.990326$, $C = 14200$) : schéma de connexion

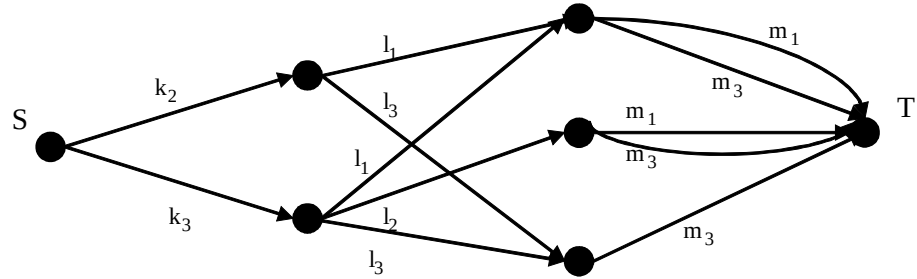


FIG. 4.29 – Architecture obtenue (SIL2 : $A_{avg} = 0.992653$, $C = 14800$) : Réseau de fiabilité

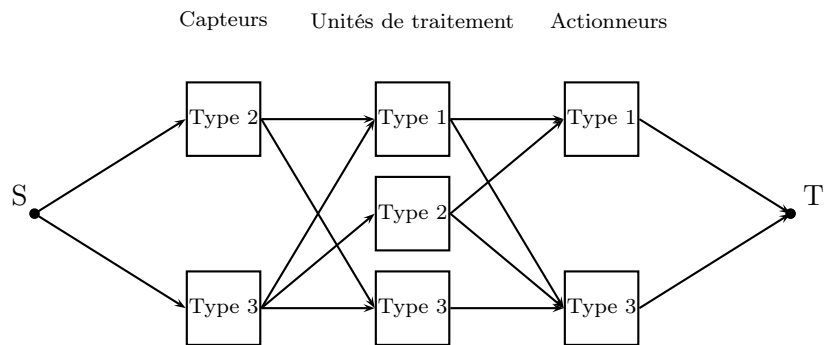


FIG. 4.30 – Architecture obtenue (SIL2 : $A_{avg} = 0.992653$, $C = 14800$) : schéma de connexion

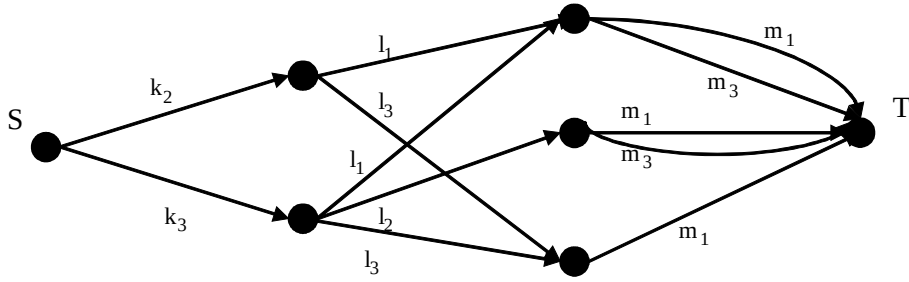


FIG. 4.31 – Architecture obtenue (SIL2 : $A_{avg} = 0.992805$, $C = 14800$) : Réseau de fiabilité

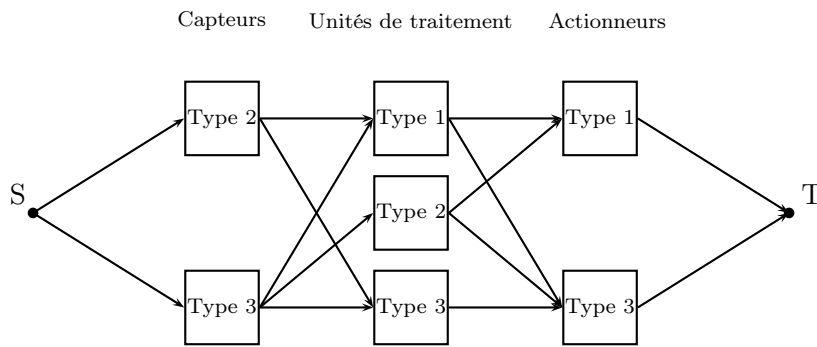


FIG. 4.32 – Architecture obtenue (SIL2 : $A_{avg} = 0.992805$, $C = 14800$) : schéma de connexion

disponibles (les trois composants disponibles pour chaque sous-système du SIS) et toutes les connexions possibles entre ces composants.

Pour l'obtention d'un SIL3 ou d'un SIL4, il faut donc disposer de composants plus fiables que ceux qui nous sont fournis. Il faut donc utiliser des composants avec des disponibilités plus élevées pour augmenter la disponibilité totale du SIS. En effet, si nous proposons d'utiliser des composants dont les probabilités de défaillance et les coûts sont donnés dans le tableau 4.6 alors nous arrivons à obtenir une architecture qui assure l'obtention d'un SIL 3 (cf. figure 4.33 et figure 4.34).

En conclusion, nous pouvons dire que dans l'ensemble des cas, l'utilisation de la méthodologie que nous avons proposée, nous permet d'obtenir une structure optimale du SIS qui respecte les

Composants du SIS	Sous systèmes					
	Capteurs		Unités de traitements		Actionneurs	
	c_1 (unités)	p_1	c_2 (unités)	p_2	c_3 (unités)	p_3
Type 5	2100	0.019	1400	0.04	2500	0.02
Type 6	1500	0.03	2100	0.05	3500	0.03
Type 7	2000	0.0225	1200	0.03	4100	0.04

TAB. 4.6 – Coûts et probabilités de défaillance des composants du SIS n°2

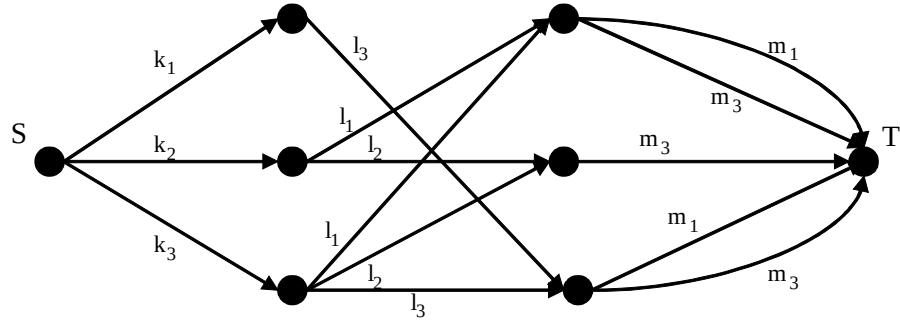


FIG. 4.33 – Architecture obtenue (SIL3 : $A_{avg} = 0.999003$, $C = 16900$) : Réseau de fiabilité

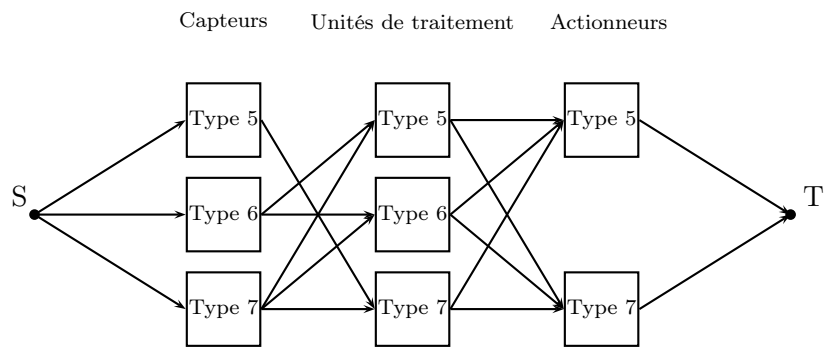


FIG. 4.34 – Architecture obtenue (SIL3 : $A_{avg} = 0.999003$, $C = 16900$) : schéma de connexion

contraintes si les paramètres de fiabilité des composants dont on dispose le permettent. Ainsi, pour obtenir un SIL3 ($A_{avg} \geq 0.999$) ou un SIL4 ($A_{avg} \geq 0.9999$), il faut utiliser d'autres composants plus fiables et qui permettent donc d'augmenter la disponibilité moyenne du SIS. En outre, nous vérifions clairement que notre approche peut être appliquée à tout type de systèmes ayant une structure en couches et propose des architectures optimales complexes et non forcément séries parallèles.

4.6 Conception optimale à paramètres flous

4.6.1 Introduction

Dans les problèmes de conception des SIS, nous rencontrons la même problématique que celle de l'évaluation des SIL : la prise en compte des connaissances imparfaites relatives aux données de fiabilité des composants. La validité des résultats d'optimisation dépend donc de la prise en compte de l'imperfection des connaissances utilisées. Dans un premier temps, nous modelisons les probabilités de défaillance imprécises des composants disponibles sur le marché par des nombres flous. Dans un deuxième temps, le modèle de conception des SIS présenté dans la section précédente a été adapté à la prise en compte des probabilités de défaillance floues. Nous avons gardé le même AG et nous avons changé la méthode de calcul de la disponibilité moyenne A_{avg} du SIS en utilisant l'arithmétique des nombres flous. La dernière partie présente les conclusions issues de l'utilisation de cette approche.

4.6.2 Méthodologie générale

Par rapport à la méthodologie de conception présentée dans la section précédente (cf. organigramme 4.19), deux étapes ont changé :

- Proposition de différents types de composants pour la conception du SIS :

Certaines bases de données de fiabilité [IEEE, 1984; CCPS, 1991; CCPS, 2002] fournissent les bornes inférieures et supérieures, les valeurs moyennes et les facteurs d'erreurs des données de fiabilité des composants. C'est pourquoi, nous utilisons ces valeurs pour construire les fonctions d'appartenance des taux de défaillance des composants disponibles sur le marché. Puis, nous évaluons leurs probabilités de défaillance à partir de leurs taux de défaillance (cf. chapitre 3).

- Calcul de la disponibilité moyenne A_{avg} du SIS :

A l'aide de l'opérateur arithmétique flou d'addition, le nombre flou qui représente la disponibilité instantanée du SIS est calculé à partir des nombres flous qui représentent les disponibilités des liens minimaux l_i du réseau de fiabilité. Les disponibilités des liens sont calculées à partir des probabilités de défaillance floues des composants :

$$\mu_{\widetilde{A}(t)}(x) = \sum_{i=1}^n \mu_{P(\widetilde{l}_i)(t)}(x) \quad (4.39)$$

Ensuite, nous procédons à la défuzzification de $\mu_{\widetilde{A}(t)}$ par la méthode du centre de gravité pour obtenir une valeur précise de la disponibilité instantanée :

TAB. 4.7 – Paramètres flous des probabilités de défaillance des capteurs

Capteurs	a_i	m_i	b_i
Type 1	0.017	0.039	0.059
Type 2	0.065	0.07	0.08
Type 3	0.02	0.03	0.05

TAB. 4.8 – Paramètres flous des probabilités de défaillance des unités de traitement

Unités de traitement	a_i	m_i	b_i
Type 1	0.086	0.09	0.092
Type 2	0.048	0.05	0.052
Type 3	0.065	0.07	0.076

$$A(t) = defuzz(\mu_{\widetilde{A(t)}}(x)) \quad (4.40)$$

Enfin, la disponibilité moyenne A_{avg} du SIS est obtenue en intégrant la disponibilité instantanée sur la période d'étude T :

$$A_{avg} = \frac{1}{T} \int_0^T A(t) dt \quad (4.41)$$

Avec :

n : le nombre de liens minimaux du réseau de fiabilité.

$P(l_i)(t)$: la disponibilité instantanée floue du lien minimal l_i .

4.6.3 Cas d'application

Dans cette partie, nous reprenons l'application précédente (cf. section 4.5.5). Le seul changement est que nous sommes en présence de composants qui ont des probabilités de défaillance imprécises. En outre, nous gardons les mêmes paramètres de l'AG utilisé précédemment. Dans les figures 4.35, 4.36 et 4.37, nous représentons les nombres flous triangulaires représentant les probabilités imprécises de défaillance des composants (capteurs, unités de traitement et actionneurs) susceptibles d'être utilisés dans le SIS. Ces probabilités ont été calculées à partir des taux de défaillance imprécis des composants. Les paramètres de distributions des probabilités de défaillance sont donnés dans les tableaux 4.7, 4.8 et 4.9. Il faut signaler que ce sont les mêmes composants qui ont été utilisés dans l'approche classique, mais avec une imprécision concernant leurs probabilités de défaillance. Il faut noter que les valeurs modales ($Pi(\alpha)/\mu_{Pi(\alpha)} = 1$) des probabilités floues sont égales aux valeurs précises utilisées dans la conception optimale classique.

TAB. 4.9 – Paramètres flous des probabilités de défaillance des actionneurs

Actionneurs	a_i	m_i	b_i
Type 1	0.086	0.10	0.11
Type 2	0.057	0.06	0.061
Type 3	0.038	0.04	0.042

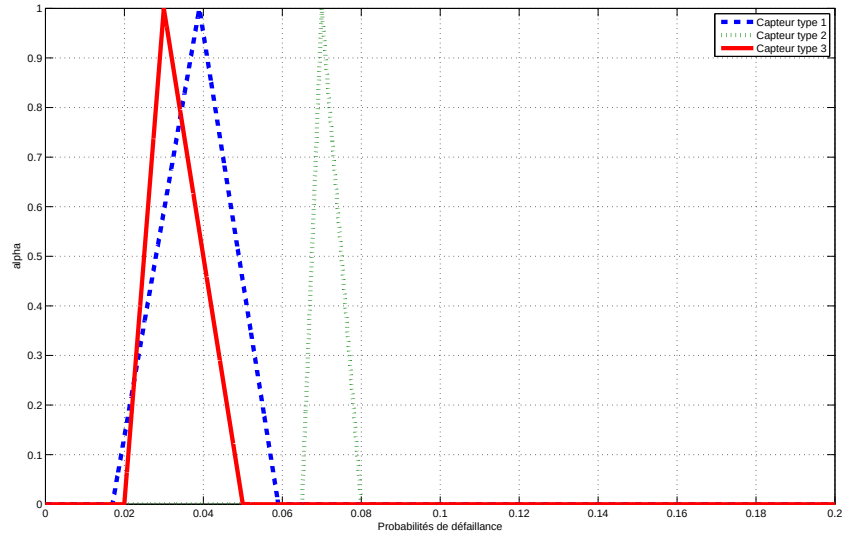


FIG. 4.35 – Probabilités de défaillance floues des capteurs

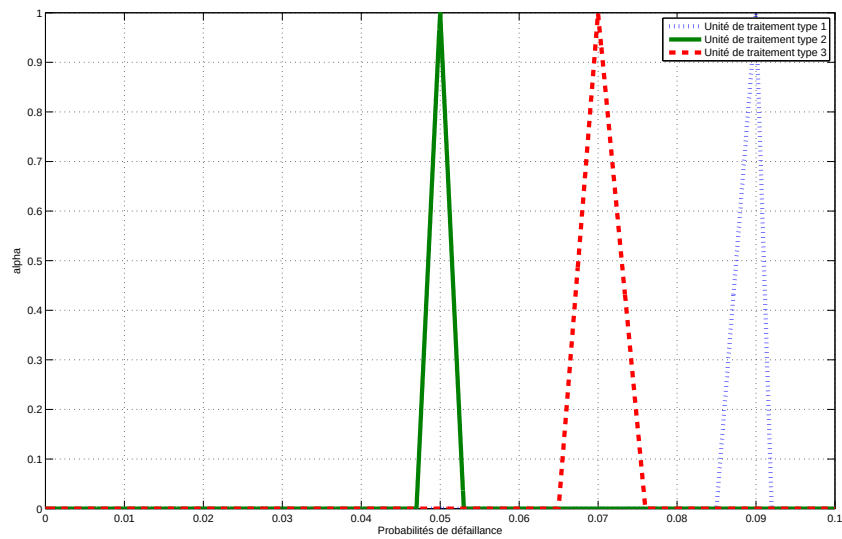


FIG. 4.36 – Probabilités de défaillance floues des unités de traitement

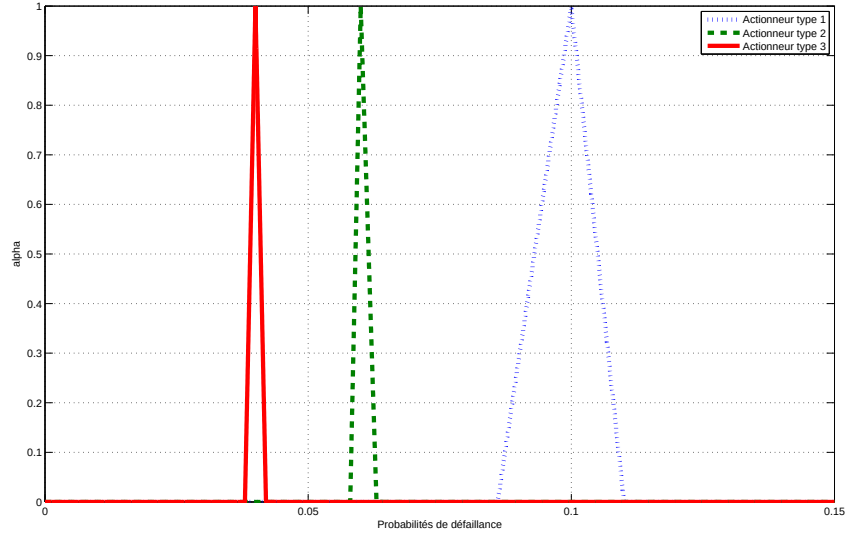


FIG. 4.37 – Probabilités de défaillance floues des actionneurs

Statistiques	$C(x)$	$A_{avg}(x)$	SIL
Moyenne	9000	0.923	1
Maximum	9400	0.931	1
Minimum	8800	0.9169	1
Écart type	740	0.0012	-

TAB. 4.10 – Approche floue : résultats pour l'obtention d'un SIL1

4.6.4 Résultats et analyse

4.6.4.1 Exigence SIL1

Le tableau 4.10 présente les statistiques relatives aux résultats obtenus. Nous remarquons que pour tous les essais, la contrainte du SIL1 a été respecté ($A_{min} = 0.9169$). En outre, la durée de l'optimisation est beaucoup plus longue que celle de l'approche classique (en moyenne 10 fois plus). Les figures 4.38 à 4.41 présentent les schémas de connexion et les réseaux de fiabilité des meilleures architectures optimales obtenues. Nous remarquons que nous obtenons des architectures semblables à celles obtenues dans l'approche classique. En effet, comme l'imprécision des probabilités de défaillance des composants est faible, il est tout à fait normal que pour une disponibilité moyenne faible (SIL1), nous retrouvions les mêmes architectures. Donc pour une exigence supérieure (SIL2), nous nous attendons à ce que les architectures dans les approches floues et classiques ne soient pas les mêmes.

4.6.4.2 Exigence SIL2

Le tableau 4.11 présente les statistiques relatives aux résultats obtenus. Les figures 4.42, 4.43, 4.44 et 4.46 présentent les schémas de connexion et les réseaux de fiabilité des meilleures architectures optimales obtenues. La meilleure configuration du SIS est obtenue avec un coût de

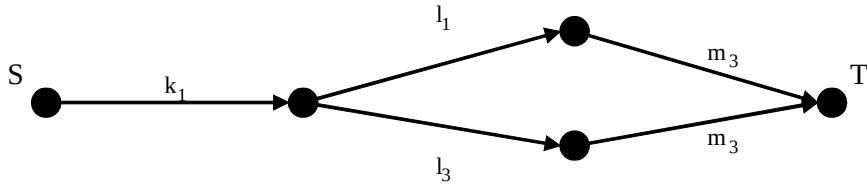


FIG. 4.38 – Architecture obtenue (SIL1 : $A_{avg} = 0.916904, C = 8800$) : Réseau de fiabilité

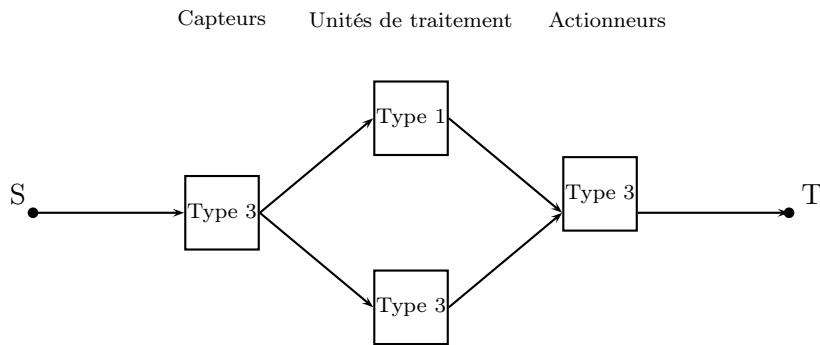


FIG. 4.39 – Architecture obtenue (SIL1 : $A_{avg} = 0.916904, C = 8800$) : schéma de connexion

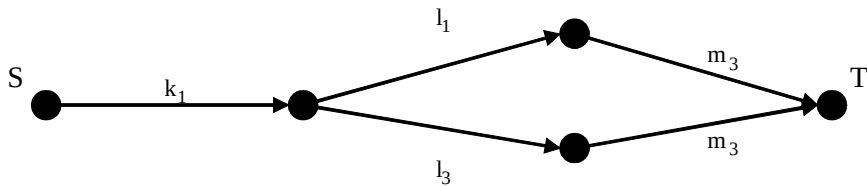


FIG. 4.40 – Architecture obtenue (SIL1 : $A_{avg}=0.912206, C=8800$) : Réseau de fiabilité

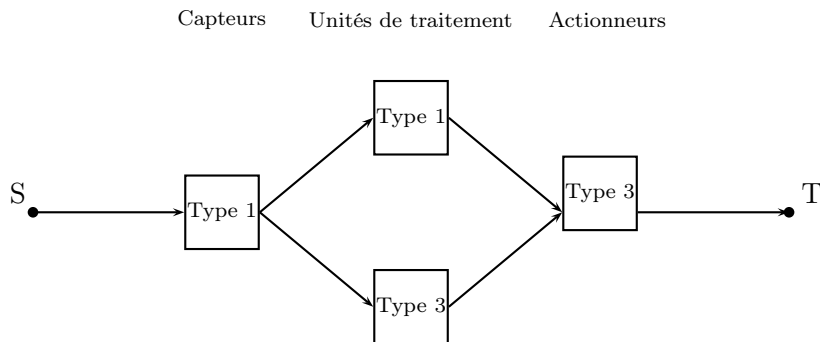
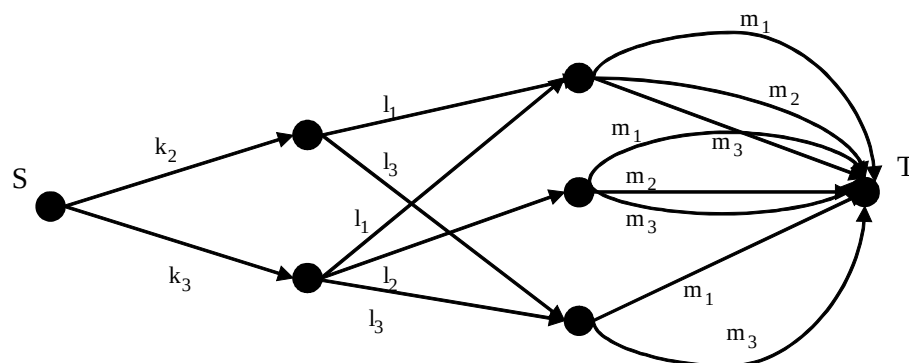


FIG. 4.41 – Architecture obtenue (SIL1 : $A_{avg}=0.912206, C=8800$) : schéma de connexion

Statistiques	$C(x)$	$A_{avg}(x)$	SIL
Moyenne	18700	0.9941	2
Maximum	20400	0.9991	2
Minimum	18300	0.9904	2
Écart type	1840	0.0033	-

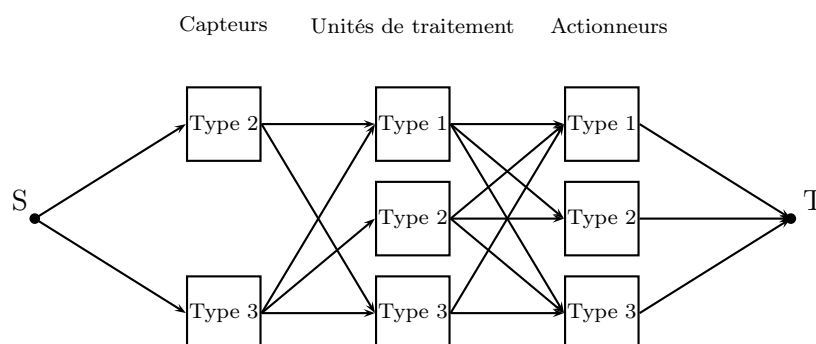
TAB. 4.11 – Approche floue : résultats pour l'obtention d'un SIL2

FIG. 4.42 – Architecture obtenue (SIL2 : $A_{avg}=0.990437, C=18300$) : Réseau de fiabilité

18300 unités et une disponibilité moyenne égale à 0.9904. Nous vérifions ainsi que nous n'obtenons pas des architectures semblables à celles obtenues dans l'approche classique (cf. figures 4.31 et 4.30). Ce résultat est dû à la valeur élevée de la disponibilité moyenne exigée (SIL2).

4.6.4.3 Exigence SIL3

Le tableau 4.12 présente les statistiques relatives aux résultats obtenus. Les figures 4.46 et 4.47 présentent les schémas de connexion et les réseaux de fiabilité des meilleures architectures optimales obtenues. La meilleure configuration du SIS est obtenue avec un coût de 20400 unités et une disponibilité moyenne égale à 0.99906. Nous remarquons que nous n'obtenons pas d'architectures semblables à celles obtenues dans l'approche classique. Ce résultat est également dû

FIG. 4.43 – Architecture floue obtenue (SIL2 : $A_{avg}=0.990437, C=18300$) : schéma de connexion

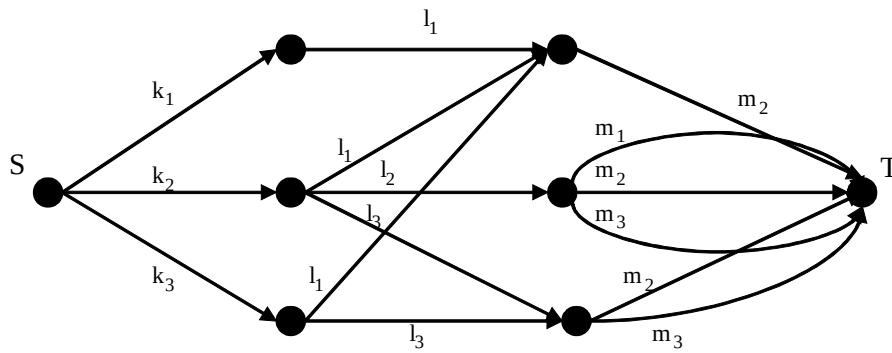


FIG. 4.44 – Architecture obtenue (SIL2 : $A_{avg}=0.994834, C=20400$) : Réseau de fiabilité

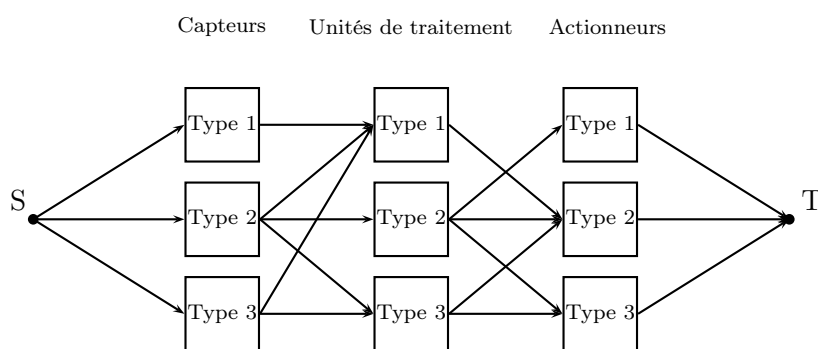
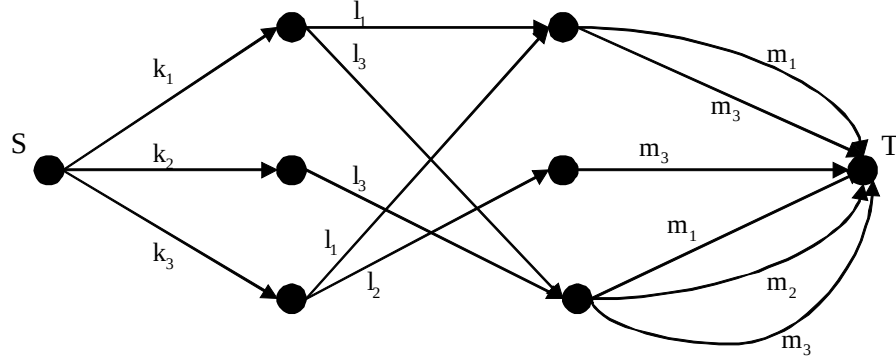


FIG. 4.45 – Architecture obtenue (SIL2 : $A_{avg}=0.994834, C=20400$) : schéma de connexion

Statistiques	$C(x)$	$A_{avg}(x)$	SIL
Moyenne	20400	0.99909	3
Maximum	20400	0.9991	3
Minimum	20400	0.99906	3
Écart type	1010	0.0024	-

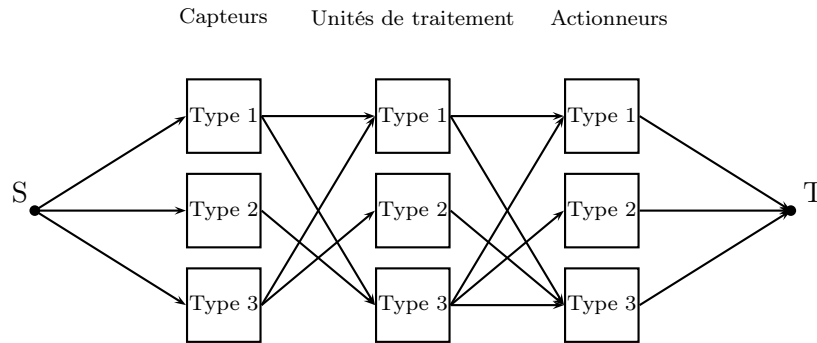
TAB. 4.12 – Approche floue : résultats pour l’obtention d’un SIL3


FIG. 4.46 – Architecture obtenue (SIL3 : $A_{avg} = 0.99906, C = 20400$) : Réseau de fiabilité

à la valeur élevée de la disponibilité moyenne exigée (SIL3).

4.6.4.4 Exigence SIL4

Les composants dont nous disposons pour la conception du SIS ne permettront pas l’obtention d’un SIL4 parce que la disponibilité moyenne maximale que nous pouvons obtenir est 0.9991. Cela ne permettra pas d’atteindre $A_{avg} = 0.9999$ nécessaire pour un SIL4. Dans l’approche floue, la disponibilité moyenne maximale du SIS est obtenue aussi en utilisant tous les composants disponibles (les trois composants disponibles pour chaque sous-système du SIS) et toutes les connexions possibles entre ces composants.


FIG. 4.47 – Architecture obtenue (SIL3 : $A_{avg} = 0.99906, C = 20400$) : schéma de connexion

Pour obtenir un SIL4, il faut donc disposer, comme pour l'approche classique, de composants plus fiables que ceux qui nous sont fournis.

4.7 Conclusion et perspectives

Dans ce chapitre, nous avons traité le problème de conception optimale des SIS pour le respect du SIL exigé. Il est très important de signaler que notre approche peut être appliquée à tout type de système. Nous avons commencé par dresser un état de l'art des méthodes d'allocation de fiabilité et de redondance. Ensuite, deux méthodes de modélisation des SIS ont été introduites : les BdF et les réseaux de fiabilité. Nous avons utilisé un AG pour l'optimisation de la structure des SIS et l'obtention de configurations optimales des SIS.

Le problème de conception optimale des SIS a été ramené à un problème de minimisation du coût global du SIS sous contrainte du SIL exigé qui est exprimé en fonction de la disponibilité moyenne du SIS. La conception optimale du SIS se fait donc par le choix d'un nombre et types de composants dans chaque sous-système du SIS ainsi que les connexions entre ces composants afin de garantir un coût global minimal et sous contrainte du SIL exigé.

Une application référencée dans le document [ISA-TR84.00.02-2002, 2002] a servi de support à l'illustration de la méthodologie proposée. Nous avons déterminé les configurations optimales pour les SIL exigés sous forme de BdF. Les résultats de l'application étudiée valident la méthodologie que nous avons proposée et montrent son importance pour la conception des systèmes complexes en général et des SIS en particulier. Nous avons obtenu pour chaque SIL des structures optimales que ce soit pour l'approche classique ou floue. En outre, nous avons remarqué que la plupart de ces structures sont complexes. Cela démontre que notre approche n'est pas restreinte uniquement aux systèmes séries parallèles et parallèles séries à condition que les systèmes aient une structure en couches.

Conclusion générale

Les travaux de recherche présentés dans cette thèse ont porté sur la problématique de la prise en compte des incertitudes relatives aux données de fiabilité des composants pour l'évaluation de la sûreté de fonctionnement des systèmes en général et des SIS en particulier. En outre, une méthodologie d'aide à la conception des SIS a été proposée. Elle est basée sur l'optimisation du choix des composants et de la structure des SIS pour le respect du SIL exigé.

Bilan

Nous avons d'abord présenté la problématique de la gestion de risque dans les entreprises. Nous avons repris les notions clés comme la sécurité, le risque et le danger et, nous avons cherché à faire le lien entre ces différents termes. Ensuite, nous avons introduit la notion de sécurité fonctionnelle et, nous avons montré comment il était possible aujourd'hui de réaliser des fonctions de sécurité à l'aide de système E/E/PE. Nous avons évoqué les difficultés liées à l'évaluation de la sûreté de fonctionnement de ces systèmes sécuritaires. Nous avons présenté, par la suite, les méthodes de sûreté de fonctionnement définies dans le cadre de la théorie de fiabilité PROBIST floue (hypothèse probabiliste floue et hypothèse d'états binaires). Nous avons d'abord introduit les notions de bases concernant les théories de représentation et d'analyse des informations imparfaites. Puis, nous avons dressé un état de l'art de l'emploi de ces théories par les principales méthodes de sûreté de fonctionnement en présence d'incertitudes.

Nous avons proposé l'utilisation des arbres de défaillance flous pour évaluer le PFD_{avg} des SIS et déduire le SIL des SIF exigées pour ces SIS. Ainsi, en utilisant la méthode de Soman [Soman and Misra, 1993] basée sur les α -coupes, nous avons obtenu un nombre flou du PFD_{avg} du SIS. Cette probabilité floue a mis en évidence l'existence d'incertitude concernant le SIL du SIS. L'utilisation des mesures de possibilité et de nécessité nous a permis de caractériser l'incertitude des fonctions d'appartenance des PFD_{avg} des SIS. La quantification de cette incertitude a souligné le fait que plusieurs SIL étaient probables pour un même SIS à cause de l'imprécision des taux de défaillance de ses composants. En outre, l'introduction de nouveaux facteurs d'importance a permis d'identifier les composants dont l'incertitude de la probabilité de défaillance contribue significativement à l'incertitude entachant la probabilité de défaillance du SIS et donc du SIL. Ainsi, nous avons apporté une aide à la décision aux fiabilistes, pour réduire d'une manière efficace les incertitudes dans la détermination des SIL en se concentrant principalement sur les composants présentant le plus grand facteur d'incertitude flou ou le plus grand facteur d'importance flou. L'application de l'approche proposée ainsi que les facteurs d'importance flous à un exemple applicatif a montré que les résultats obtenus sont très satisfaisants. En effet, nous avons caractérisé l'incertitude des résultats obtenus et nous avons réussi à réduire efficacement l'incertitude du SIL obtenu grâce aux études de sensibilité.

Enfin, nous avons ramené le problème de conception des SIS à un problème de minimisation du coût global du SIS sous contrainte du SIL exigé qui est exprimé en fonction de la disponibilité moyenne du SIS. La conception optimale du SIS se fait donc par le choix d'un certain nombre et type de composants dans chaque sous-système du SIS qui garantissent un coût global minimal et le SIL exigé. Comme l'a souligné Elegbede *et al.* [Elegbede, 2000], les problèmes d'allocation de redondance ont traité principalement les systèmes séries parallèles et parallèles séries. Cette restriction est due au fait que pour ce type de systèmes nous disposons d'une expression explicite de la fiabilité totale qui facilite le processus d'optimisation. Cela n'est pas le cas pour des structures plus complexes. C'est pourquoi, nous avons proposé une nouvelle approche qui est appliquée à tout type de structure pouvant être modélisée par des réseaux de fiabilité. Les résultats de l'exemple obtenus valident la méthodologie que nous avons proposée et montre son utilité pour la conception des SIS à structure complexe. En particulier, nous avons obtenu des structures complexes pour certains SIL. En outre, La méthodologie de conception proposée a été adaptée aux situations où nous sommes en présence de composants qui ont des taux de défaillance imprécis. Dans l'approche floue, nous avons également réussi à obtenir des architectures optimales pour les différents SIL exigés.

Perspectives

L'ensemble des résultats obtenus dans ce mémoire se situe dans le cadre défini dans l'introduction générale : composants non réparables, taux de défaillance constants, absence de causes communes de défaillances. Dans ce qui suit, nous allons revenir sur ces hypothèses en proposant et en mettant en évidence quelques perspectives d'extension des résultats obtenus.

- La prise en compte des paramètres figurant dans les normes IEC 61508 [IEC61508, 1998] et IEC 61511 [IEC61511, 2000] :
 - Facteur β de défaillances de cause commune. Dans nos travaux, il est aisé de prendre en compte ce facteur en ajoutant simplement le même terme dans les liens du réseau de fiabilité qui présentent des défaillances de cause commune.
 - Modes de défaillances : défaillances dangereuses λ^D (détectables ou pas) et défaillances sûres λ^S (détectables ou pas).
 - Couverture du diagnostic (DC) : Rapport du taux de défaillances dangereuses détectées par les tests automatiques sur le taux de défaillances dangereuses.
- L'étude de la faisabilité des architectures optimales obtenues :
 - Association de composants de constructeurs différents.
 - Connexions et bus de communication utilisés dans la structure du SIS. En effet, si nous considérons des connexions point à point, il faut ajouter le coût de ces connexions. Par contre, si nous utilisons un bus de communication, il n'y a pas de surcoût. Nous serons alors en présence de modes communs de défaillances.
- L'exploitation des architectures optimales obtenues pour le calcul du temps moyen de bon fonctionnement du SIS qui tient compte des défaillances fugitives au niveau des connexions.

-
- L'intégration de l'aspect de la maintenance des composants des SIS dans la méthode d'aide à la décision que nous avons proposée. En particulier, il va falloir tenir compte des taux de réparations des composants et étudier l'influence des périodes de tests des composants sur les résultats obtenus.
 - La proposition de nouveaux facteurs d'importance flous inspirés du facteur d'importance de Lambert ou du facteur d'importance de Vesely-Fussell.
 - Enfin, il serait intéressant d'exploiter d'autres méthodes pour l'évaluation de la sûreté de fonctionnement des SIS en présence d'informations imparfaites comme les réseaux de Petri ou les approches markoviennes.

Bibliographie

- [Abraham, 1979] J. A. Abraham. An improved algorithm for network reliability. *IEEE Transactions on Reliability*, 28 :58–61, 1979.
- [AGREE, 1957] AGREE. *Reliability of military electronic equipment*. Advisory Group on Reliability of Electronic Equipment, 1957.
- [Allain-Morin and Cloarec, 1994] G. Allain-Morin and J-M. Cloarec. L'allocation d'objectifs de sûreté de fonctionnement en phase de spécification et de conception. In *9ème Colloque International de Fiabilité et Maintenabilité, Lambda mu 9 et ESREL 94*, pages 43–52, La Baule, France, 1994.
- [Alven, 1964] W. H. Von Alven. *Reliability engineering*. Englewood Cliff, Prentice Hall, 1964.
- [ANSI/ISA-S84.01-1996, 1996] ANSI/ISA-S84.01-1996. *Application of Safety Instrumented Systems for the process control industry*. Instrumentation Society of America (ISA), 1996.
- [Aubry, 2003] J-F. Aubry. *Méthode de l'arbre des causes (MAC)*. Ecole Nationale Supérieur d'Electricité et de Mécanique de Nancy (ENSEM), 2003.
- [Baea et al., 2004] H. Baea, V. Grandhi, and A. Canfield. Epistemic uncertainty quantification techniques including evidence theory for large-scale structures. *Computers and Structures*, 82 :1101–1112, 2004.
- [Barlow and Proschan, 1973] B. E. Barlow and F. Proschan. *Importance of system components and fault tree analysis*. Operations Research Center, Univ. of California, Berkeley, 1973.
- [Baudrit, 2005] C. Baudrit. *Représentation et propagation de connaissances imprécises et incertaines : application à l'évaluation des risques liés aux sites et aux sols pollués*. PhD thesis, Université de Toulouse III Paul Sabatier, France, 2005.
- [Beasley et al., 1993] D. Beasley, D.R. Bull, and R.R. Martin. An overview of genetic algorithms : Part 1, fundamentals. *University Computing*, 15 :58–59, 1993.
- [Beckman, 2000] L. Beckman. Expanding the applicability of ISA TR84.02 in the field. *ISA Transactions*, 39 :357–361, 2000.
- [Bellman and Zadeh, 1970] R. E. Bellman and L. Zadeh. Decision-making in a fuzzy environment. *Management Science Series*, B17 :141–164, 1970.
- [Berge, 1958] C. Berge. *Théorie des graphes et ses applications*. Dunod, France, 1958.
- [Beugin, 2006] J. Beugin. *Contribution à l'évaluation de la sécurité des systèmes complexes de transport guidé*. PhD thesis, Université de Valenciennes et du Hainaut-Cambrésis, 2006.
- [Bhattacharyya, 1999] M. Bhattacharyya. Fuzzy markovian decision process. *Fuzzy Sets and Systems*, pages 273–282, 1999.
- [Bhimavarapu et al., 1997] K. Bhimavarapu, L. Moore, and P. Stavrianidis. Performance based safety standards : an integrated risk assessment program. *ISA TECH*, 1, 1997.

- [Bier, 1983] V. Bier. *A measure of uncertainty importance for components in fault trees*. PhD thesis, MIT, Cambridge, Jan 1983.
- [Birnbaum, 1969] Z. W. Birnbaum. On the importance of different components in a multicomponent system. In *Multivariate Analysis II*. P. R. Krishnaiah, Ed, N. Y :Academic, 1969.
- [Blanton, 1957] E. Blanton. Reliability-prediction technique for use in design of complex systems. *IRE National Convention Record*, 10 :68–79, 1957.
- [Bouchon-Meunier, 1995] B. Bouchon-Meunier. *La logique floue et ses applications*. Vie artificielle. 1995.
- [Boyer, 2001] M. Boyer. *Contribution à la modélisation des systèmes à temps contraint et application au multimédia*. PhD thesis, Université Paul Sabatier, Toulouse, France, 2001.
- [Bugarin and Barro, 1994] A. Bugarin and S. Barro. Fuzzy reasoning supported by petri nets. *LNCS Advances in Petri nets*, 2 :135–150, 1994.
- [Cai *et al.*, 1991] K. Y. Cai, C. Y. Wen, and M. L. Zhang. Fuzzy variables as a basis for a theory of fuzzy reliability in the possibility context. *Fuzzy Sets and Systems*, 42 :145–172, 1991.
- [Cai, 1993] K. Y. Cai. System failure engineering and fuzzy methodology : An introductory overview. *Fuzzy Sets and Systems*, 83 :113–133, 1993.
- [Camargo, 1998] H. Camargo. The design of knowledge bases through hierarchical high level fuzzy petri nets. *Fuzziness in Petri Nets*, 1998.
- [Cao, 1993] T. Cao. A fuzzy petri nets approach to reasoning about uncertainty in robotic system. In *Proceedings of IEEE International Conference on Robotics and Automation*, 1993.
- [Cardoso *et al.*, 1990] J. Cardoso, R. Valette, and D. Dubois. Petri nets with uncertain markings. *LNCS Advances in Petri nets*, 483 :64–78, 1990.
- [Cardoso, 1990] J. Cardoso. *Sur les Réseaux de Petri avec Marquages Flous*. PhD thesis, Laboratoire d'Analyse et d'Architecture des Systèmes, 1990.
- [Cardoso, 1999] J. Cardoso. Time fuzzy petri nets. *Fuzziness in Petri Nets : Studies in Fuzziness and Soft Computing*, 1999.
- [CCPS, 1991] CCPS. *Guidelines for process equipment reliability data with data tables*. Center for chemical process safety of the American institute of chemical engineers, 1991.
- [CCPS, 1993] CCPS. *Guidelines for safe automation of chemical processes*. 1993.
- [CCPS, 2002] CCPS. *Offshore reliability data handbook, 4th Edition*. 2002.
- [Chappelle, 2005] O. Chappelle. Active learning for parzen windows classifier. *AI and Statistics*, pages 49–56, 2005.
- [Chatterjee, 1974] P. Chatterjee. Fault tree analysis : Reliability theory and systems safety analysis. *Operations Research Center*, pages 34–74, 1974.
- [Chen *et al.*, 1990] S-M. Chen, J-S. Ke, and J-F. Chang. Knowledge representation using fuzzy petri nets. *IEEE Transactions on knowledge and data engineering*, 2, 1990.
- [Chern, 1992] M-S. Chern. On the computational complexity of reliability redundancy allocation in series system. *Operation Research letters*, 11 :309–315, 1992.
- [Cogis and Robert, 2003] O. Cogis and C. Robert. *Théorie des graphes : Au delà des ponts de Königsberg, Problèmes, théorèmes, algorithmes*. Vuibert, 2003.
- [Collet, 1996] J. Collet. Some remarks on rare-event approximation. *IEEE Transactions on Reliability*, 45 :106–108, 1996.

-
- [Deb *et al.*, 2006] K. Deb, A. Sinha, and S. Kukkonen. Multi-objective test problems, linkages, and evolutionary methodologies. In *genetic and evolutionary computation conference, GECCO 2006, Seattle, Washington, USA*, 2006.
- [Dempster, 1967] A. P. Dempster. Upper and lower probabilities induced by a multivalued mapping. *Annals of Mathematical Statistics*, 38 :325–339, 1967.
- [Desroches *et al.*, 2003] A. Desroches, A. Leroy, and F. Vallée. *La gestion des risques : principes et pratiques*. Lavoisier, France, 2003.
- [Desroches, 1995] A. Desroches. *Concepts et méthodes probabilistes de base de la sécurité*. Lavoisier, France, 1995.
- [Desroches, 2005] A. Desroches. Les invariants de l’analyse préliminaire des risques. In *Qualita 2005, France*, 2005.
- [Dhillon and Yang, 1997] B.S. Dhillon and N. Yang. Comparisons of block diagram and markov method system reliability and mean time to failure results for constant and non constant unit failure rates. *Microelectronics and Reliability*, 37 :505–509, 1997.
- [DIN19250, 1990] DIN19250. *Control Technology ; Fundamental Safety Aspects To Be Considered for Measurement and Control Equipment*. Automation, Software and Information Technology : ASI, 1990.
- [Dubois and Prade, 1988] D. Dubois and H. Prade. Possibility theory : An approach to computerized processing of uncertainty. *Plenum Press*, 1988.
- [Dubois and Prade, 1994] D. Dubois and H. Prade. La fusion d’informations imprécises. *Traitement du Signal*, pages 447–458, 1994.
- [Dubois *et al.*, 1993] D. Dubois, H. Prade, and S. Sandri. *On Possibility/Probability transformations*. Kluwer Academic Publishers, r. lowen and m. roubens edition, 1993.
- [Elegbede *et al.*, 2003] C. Elegbede, C. Chengbin, K.H. Adjallah, and F. Yalaoui. Reliability allocation through cost minimization. *IEEE Transactions on Reliability*, 52 :106–111, 2003.
- [Elegbede, 2000] C. Elegbede. *Contribution aux méthodes d’allocation d’exigences de fiabilité aux composants de systèmes*. PhD thesis, Université de Technologie de Troyes, 2000.
- [EN50126, 1999] EN50126. *Railway applications. The specification and demonstration of reliability, availability, maintainability and safety (RAMS)*. 1999.
- [EN50128, 2001] EN50128. *Railway applications. Communications, signalling and processing systems. Software for railway control and protection systems*. 2001.
- [EN50129, 1998] EN50129. *Safety related electronic systems for signalling*. 1998.
- [Esogbue and Bellman, 1984] A. O. Esogbue and R. E. Bellman. Fuzzy dynamic programming and its extensions. *TIMS/Studies in Management Sciences*, 20 :147–167, 1984.
- [Exida, 2005] Exida. *Safety Equipment Reliability Handbook, 2nd Edition*. 2005.
- [Farmer, 1967] F. R. Farmer. Siting criteria : a new approach. *Atom*, 128 :152–166, 1967.
- [Feng and Wu, 2001] J. Feng and M.D. Wu. The profust fault tree and its analysis. *Journal of National University of Defense Technology*, 23 :85–88, 2001.
- [Fussell, 1975] J. B. Fussell. How to hand-calculate system reliability characteristics. *IEEE Transactions on Reliability*, 24 :169–174, 1975.
- [Goble and Cheddie, 2006] W. M. Goble and H. Cheddie. *Safety Instrumented Systems Verification- Practical Probabilistic Calculations*. ISA, 2006.

- [Goldberg, 1994] D. Goldberg. *Algorithmes génétiques*. Addison-Wesley, France, 1994.
- [Gomes and Steiger-Garçao, 1995] L. Gomes and A. Steiger-Garçao. Programmable controller design based on a synchronized coloured petri net model and integrating fuzzy reasoning. *Lecture Notes in Computer Science*, (218-237), 1995.
- [Gopal et al., 1978] K. Gopal, K.K. Aggarwal, and J.S. Gupta. An improved algorithm for reliability optimization. *IEEE Transactions on Reliability*, 27 :325–328, 1978.
- [Gouriveau, 2003] R. Gouriveau. *Analyse de risques, formalisation des connaissances et structuration des données pour l'intégration des outils d'étude et de décision*. PhD thesis, Institut National Polytechnique de Toulouse, 2003.
- [Grabisch et al., 1994] M. Grabisch, H. T. Nguyen, and E. A. Walker. *Fundamentals of Uncertainty Calculi with Applications to Fuzzy Inference*. Kluwer Academic, 1994.
- [Gruhn and Cheddie, 2006] P. Gruhn and H.L. Cheddie. *Safety Instrumented Systems : Design, analysis and justification*. ISA, 2006.
- [Guiochet, 2003] J. Guiochet. *Maîtrise de la sécurité des systèmes de la robotique de service*. PhD thesis, Institut National des Sciences Appliquées de Toulouse, 2003.
- [Guo and Yang, 2006] H. Guo and X. Yang. A simple reliability block diagram method for safety integrity verification. *Reliability Engineering and System Safety*, 92 :1267–1273, 2006.
- [Guth, 1991] M. A. Guth. A probability foundation for vagueness and imprecision in fault tree analysis. *IEEE Transactions on Reliability*, 40 :563–570, 1991.
- [Haire et al., 1985] M.J. Haire, J.G. Maltese, and R.G. Sohmer. A system availability "top down" apportionment method. In *Proceedings of the Annual Reliability and Maintainability Symposium*, 1985.
- [Hanna, 1994] M. Hanna. Determination of product quality from an fms cell using petri nets. In *Proceedings of IEEE International Conference on Systems, Man and Cybernetics*, volume 2, 1994.
- [Hauge et al., 2006] S. Hauge, H. Langseth, and T. Onshus. *Reliability Data for Safety Instrumented Systems, PDS Data Handbook*. Sintef, 2006.
- [Helton and Oberkampf, 2004] J. C. Helton and W. L. Oberkampf. Alternative representations of epistemic uncertainty. *Reliability Engineering and System Safety*, 85 :1–10, 2004.
- [Holland, 1975] J. H. Holland. *Adaptation In Natural And Artificial Systems*. University of Michigan Press, 1975.
- [Hsieh, 2003] Y-C. Hsieh. A linear approximation for redundant reliability problems with multiple component choices. *Computers and industrial engineering*, 44 :91–103, 2003.
- [Huang et al., 2004] Hong Zhong Huang, Xin Tong, and Ming J. Zuo. Posbist fault tree analysis of coherent systems. *Reliability Engineering and System Safety*, 84 :141–148, 2004.
- [IEC61061, 1998] IEC61061. *Stratifiés de bois densifiés, non imprégnés, à usages électriques*. International Electrotechnical Commission (IEC), 1998.
- [IEC61508, 1998] IEC61508. *Functional safety of Electrical/Electronic/Programmable Electronic (E/E/PE) safety related systems*. International Electrotechnical Commission (IEC), 1998.
- [IEC61511, 2000] IEC61511. *Functional safety : Safety Instrumented Systems for the process industry sector*. International Electrotechnical Commission (IEC), 2000.
- [IEC61513, 2001] IEC61513. *Centrales nucléaires : Instrumentation et contrôle commande des systèmes importants pour la sûreté, Prescriptions générales pour les systèmes*. International Electrotechnical Commission (IEC), 2001.

-
- [IEC62061, 2005] IEC62061. *Sécurité des machines : Sécurité fonctionnelle des systèmes de commande électriques, électroniques et électroniques programmables relatifs à la sécurité*. International Electrotechnical Commission (IEC), 2005.
- [IEEE, 1984] IEEE. *IEEE guide to the collection and presentation of electrical, electronic, sensing component, and mechanical equipment reliability data for nuclear-power generating station*. IEEE-std-500, 1984.
- [Innal et al., 2005] F. Innal, Y. Dutuit, and M. Djebabra. Analyse critique des formules de base données dans la norme internationale cei 61508-6. In *Proceedings of the QUALITA 2005 Conference, Bordeaux, France, 2005*.
- [Innal et al., 2006] F. Innal, Y. Dutuit, and A. Rauzy. Quelques interrogations et commentaires relatifs a la norme cei 61508. In *Proceedings of the Lambda Mu 2006 Conference, Lille, France, 2006*.
- [ISA-TR84.00.02-2002, 2002] ISA-TR84.00.02-2002. *Safety Instrumented Fonctions (SIF), Safety Integrity Level (SIL), Evaluation Techniques*. Instrumentation Society of America (ISA), 2002.
- [ISA, 2002a] ISA, 2002b, *Safety Instrumented Fonctions (SIF), Safety Integrity Level (SIL), Evaluation Techniques Part 2 : Determining the SIL of a SIF via Simplified Equations, ISA-TR84.00.02-2002*. 2002.
- [ISA, 2002b] ISA, 2002d, *Safety Instrumented Fonctions (SIF), Safety Integrity Level (SIL), Evaluation Techniques Part 4 : Determining the SIL of a SIF via Markov Analysis, ISA-TR84.00.02-2002*. 2002.
- [ISO, 1999] ISO. *Aspects liés à la sécurité : Principes directeurs pour les inclure dans les normes*. Organisation internationale de normalisation, 1999.
- [ISO, 2002] ISO. *Management du risque : Vocabulaire, Principes directeurs pour l'utilisation dans les normes*. Organisation internationale de normalisation, 2002.
- [Kaufman and Gupta, 1991] A. Kaufman and M. M. Gupta. *Introduction to Fuzzy Arithmetic Theory and Application*. Van Nostrand Reinhold Company, New York, 1991.
- [Kaufmann et al., 1975] A. Kaufmann, D. Grouchko, and R. Cruon. *Modèles mathématiques pour l'étude de la fiabilité des systèmes*. Masson et Cie, France, 1975.
- [Kaufmann, 1968] A. Kaufmann. *Introduction à la combinatoire en vue des applications*. Dunod, France, 1968.
- [Kaufmann, 1972] A. Kaufmann. *Méthodes et modèles de la recherche opérationnelle*. Dunod, France, 1972.
- [Khansa, 1997] W. Khansa. *Réseaux de Petri p-temporels : contribution à l'étude des systèmes à événements discrets*. PhD thesis, Université de Savoie, Annecy, France, 1997.
- [Kim,]
- [Kleinerman and Weiss, 1954] M. M. Kleinerman and G. H. Weiss. On the reliability of networks. In *Proceedings of the National Electronics Conference*, volume 10, pages 128–136, 1954.
- [Kletz, 1999] T. A. Kletz. *HAZOP and HAZAN : Identifying and assessing process industry hazards, 4th edition*. Institution of chemical engineers, 1999.
- [Kohlas, 1997] J. Kohlas. Allocation of arguments and evidence theory. *Theoretical Computer Science*, 171 :221–246, 1997.
- [Kolmogorov, 1960] A. N. Kolmogorov. *Foundations of the Theory of Probability*. Chelsea Pub Co; 2nd edition, 1960.

- [Kruse *et al.*, 1991] R. Kruse, E. Schwecke, and J. Heinsohn. *Uncertainty and Vagueness in Knowledge-Based Systems*. Springer-Verlag, 1991.
- [Kuo *et al.*, 1978] W. Kuo, C.L. Hwang, and F.A. Tillman. A note on heuristic methods in optimal system reliability. *IEEE Transactions on Reliability*, 27 :320–324, 1978.
- [Kuo *et al.*, 2001] W. Kuo, V.R. Prasad, F.A. Tillman, and C.L. Hwang. *Optimal reliability design : fundamentals and applications*. Cambridge University Press, 2001.
- [Kurse *et al.*, 1987] R. Kurse, R. Buck-Emden, and R. Cordes. Processor power considerations- an application of fuzzy markov chains. *Fuzzy Sets and Systems*, 21 :289–299, 1987.
- [Lambert, 1975] H. E. Lambert. *Fault tree for decision making in system analysis*. PhD thesis, Lawrence Livermore Laboratory, 1975.
- [Lamy *et al.*, 2006] P. Lamy, E. Levrat, and J-J. Paques. Méthodes d'estimation des risques machines : analyse bibliographique. In *Lambda Mu, Lille, France*, 2006.
- [Levitin, 2007] G. Levitin. Block diagram method for analyzing multi-state systems with uncovered failures. *Reliability Engineering and System Safety*, 92 :727–734, 2007.
- [Li and Aggarwal, 2000] F. Li and R. K. Aggarwal. Fast and accurate power dispatch using a relaxed genetic algorithm and a local gradient technique. *Expert Systems with Applications*, 19 :159–165, 2000.
- [Li, 1996] J.P. Li. A bound heuristic algorithm for solving reliability redundancy optimization. *Microelectronics and Reliability*, 36 :335–339, 1996.
- [Liang and Wang, 1993] G-S. Liang and M-J. J. Wang. Fuzzy fault-tree analysis using failure possibility. *Microelectronics and Reliability*, 33 :583–597, 1993.
- [Lipp and Gunter, 1994] H-P. Lipp and R. Gunter. A study of fuzzy petri nets concepts. In *Proceedings European Congress on Fuzzy and Intelligent Technologies, Aachen, Germany*, 1994.
- [Ludovic, 1994] M. Ludovic. *Audit de sécurité par algorithmes génétiques*. PhD thesis, Université de Rennes 1, France, 1994.
- [Lutton, 1999] E. Lutton. *Algorithmes génétiques et Fractales*. Habilitation à diriger des recherches, Université Paris XI Orsay, France, 1999.
- [Marson *et al.*, 1989] M. Ajmone Marson, G. Balbo, A. Bobbio, G. Chiola, G. Conte, and A. CUMANI. The effects of execution policies on semantics and analysis of stochastic petri nets. *IEEE Transactions on Software Eng.*, 15 :832–846, 1989.
- [Matarazzo and Munda, 2001] B. Matarazzo and G. Munda. New approach for the comparison of l-r fuzzy numbers : a theoretical and operational analysis. *Fuzzy Sets and Systems*, 118 :407–418, 2001.
- [Mazouni *et al.*, 2007] M-H. Mazouni, D. Bied Charreton, and J-F. Aubry. Proposal of a generic methodology to harmonize preliminary hazard analyses for guided transport. In *IEEE(SMC), San Antonio-TX, USA*, 2007.
- [McCormick, 1981] N. J. McCormick. *Reliability and Risk Analysis*. 1981.
- [Mearns, 1965] A. B. Mearns. Fault tree analysis : the study of unlikely events in complex systems. In *Boeing/UW System safety symposium*, 1965.
- [Merlin, 1974] P. Merlin. *A study of the recoverability of computer system*. PhD thesis, Dep. Comput. Sci., Univ., Californie, Irvine, 1974.
- [Michalewicz and Fogel, 2000] Z. Michalewicz and D.B. Fogel. *How to solve it : Modern heuristics*. Springer-Verlag, 2000.

-
- [MIL217B, 1974] MIL217B. *Reliability Prediction of Electronic Equipment*. Dept. of Defence, USA, 1974.
- [Ministère de l'écologie, 2007] Ministère de l'écologie. Les risques majeurs. *Prime.net*, 2007.
- [Misra and Weber, 1989] K. B. Misra and G. G. Weber. A new method for fuzzy fault tree analysis. *Microelectronics and Reliability*, 29 :195–216, 1989.
- [Murata *et al.*, 1999] T. Murata, T. Suzuki, and S. M. Shatz. Fuzzy-timing high-level petri nets (fthns) for time-critical systems. *Fuzziness in Petri Nets : Studies in Fuzziness and Soft Computing*, pages 88–114, 1999.
- [Nakagawa and Miyazaki, 1981] Y. Nakagawa and S. Miyazaki. An experimental comparison of the heuristic methods for solving reliability optimization problems. *IEEE Transactions on Reliability*, 30 :181–184, 1981.
- [Nakagawa and Nakashima, 1977] Y. Nakagawa and K. Nakashima. A heuristic method for determining optimal reliability allocation. *IEEE Transactions on Reliability*, 26 :156–161, 1977.
- [of Defence, 1998] U.S Department of Defence. *MIL-HDBK-338B Electronic reliability design handbook*. 1998.
- [of European Communities, 1992] Commission of European Communities. *Benchmark exercise on major hazard analysis*. 1992.
- [OHSAS18001, 1999] OHSAS18001. *Système de management de la santé et de la sécurité au travail - Spécification* -. BSI, Afnor, 1999.
- [Onisawa, 1988] T. Onisawa. A representation of human reliability using fuzzy concepts. *Information Sciences*, 45 :153–173, 1988.
- [OREDA, 2002] OREDA. *Offshore reliability data handbook, 4th Edition*. 2002.
- [Ouarzizi, 1997] A. E. Ouarzizi. Line fitting in noisy data using genetic algorithm. In *3ème Conférence Internationale sur le Contrôle Qualité par Vision Artificielle QCAV'97*, pages 214–217, 1997.
- [P. V. Suresh and Raj, 1996] A. K. Babar P. V. Suresh and V. V. Raj. Uncertainty in fault tree analysis : A fuzzy approach. *Fuzzy Sets and Systems*, 83 :135–141, 1996.
- [Pages and Gondran, 1980] A. Pages and M. Gondran. *Fiabilité des systèmes*. Editions Eyrolles, France, 1980.
- [Painton and Campbell, 1995] L. Painton and J. Campbell. Genetic algorithm in optimization of system reliability. *IEEE Transactions on Reliability*, 44(2) :172–180, 1995.
- [Pan and Tai, 1988] Z. J. Pan and Y. C. Tai. Variance importance of system component by monte carlo. *IEEE Transactions on Reliability*, R-37 :521–523, 1988.
- [Parzen, 1962] E. Parzen. On estimation of a probability density function and mode. *Annals of the Institute of Statistical Mathematics*, 33 :1065–1076, 1962.
- [Pedrycz and Gomide, 1994] W. Pedrycz and F. Gomide. A generalized fuzzy petri net model. *IEEE Transactions on Fuzzy Systems*, 2 :295–301, 1994.
- [Peguet and Boyer, 2006] F-X. Peguet and A. Boyer. Mise en oeuvre et interpretation de la norme CEI 61508 par la société LGM. In *Proceedings of the Lambda Mu 2006 Conference, Lille, France*, 2006.
- [Peterson, 1977] J. L. Peterson. Petri nets. *Computing Surveys*, 1 :153–170, 1977.
- [Petri, 1962] C. A. Petri. Kommunikation mit automaten. Technical report, Schriften des Instituts für instrumentelle Mathematik. Universität Bonn, 1962.

- [Prade and Dubois, 1983] H. Prade and D. Dubois. Unfair coins and necessity measures : toward a possibilistic interpretation of histograms. *Fuzzy Sets and Systems*, 10 :15–20, 1983.
- [Prade and Dubois, 1985] H. Prade and D. Dubois. A review of fuzzy set aggregation connectives. *Information sciences*, 36 :85–121, 1985.
- [Périllon, 1998] P. Périllon. *L'analyse des risques : méthode MOSAR*. EDF, 1998.
- [Quinquis and Radoi, 1998] A. Quinquis and E. Radoi. Comparaison de méthodes de classification : application à l'identification d'objets ferromagnétiques sous-marin. *Traitement du signal*, 15 :181–195, 1998.
- [Racoceano, 2006] D. Racoceano. *Contribution à la surveillance des Systèmes de Production en utilisant les techniques de l'intelligence artificielle*. Habilitation à diriger des recherches, Université de FRANCHE COMTÉ de Besançon, France, 2006.
- [Ramirez-Marquez et al., 2004] J. Ramirez-Marquez, D. Coit, and A. Konak. Reliability optimization of series-parallel systems using a max-min approach. *IIE Transactions*, 36 :891–898, 2004.
- [Rao, 1996] S. S. Rao. *Engineering optimization-theory and practice, 3rd edition*. John Wiley and Sons, New York, 1996.
- [Rauzy et al., 2003] A. Rauzy, E. Châtelet, Y. Dutuit, and C. Bérenguer. A practical comparison of methods to assess sum-of-products. *Reliability Engineering and System Safety*, 79 :33–42, 2003.
- [Ribaric et al., 1998] S. Ribaric, B. Dalbelo-Basic, and D. Tomac. Program implementation of reasoning using fuzzy and fuzzy time petri nets. In *20th Int. Conf. on Information Technology Interfaces (ITI'98), Croatia*, pages 1119–1124, 1998.
- [Ruegg, 1985] A. Ruegg. *Probabilités et statistique*. Lausanne, Presse Polytechniques Romandes, 1985.
- [Sallak et al., 2005] M. Sallak, C. Simon, and J-F. Aubry. Impact de l'imprécision des taux de défaillances dans les arbres de défaillances et facteurs d'importance flous. In *Journées Doctorales Modélisation, Analyse et Conduite des Systèmes dynamiques, JDMACS 2005, CDROM papier n°83, Lyon, France*, 2005.
- [Sallak et al., 2006a] M. Sallak, C. Simon, and J-F. Aubry. Aide à la décision dans la réduction de l'incertitude des SIL : une approche floue/possibiliste. *e-STA, Revue des Sciences et Technologies de l'Automatique*, 2006.
- [Sallak et al., 2006b] M. Sallak, C. Simon, and J-F. Aubry. Evaluating safety integrity level in presence of uncertainty. In *KONBiN 2006, The 4th International Conference on Safety and Reliability*, Krakow, Poland, 2006.
- [Sallak et al., 2006c] M. Sallak, C. Simon, and J-F. Aubry. On the use of a new possibilist importance measure to reduce safety integrity level uncertainty. In *KONBiN 2006, The 4th International Conference on Safety and Reliability*, Krakow, Poland, 2006.
- [Sallak et al., 2006d] M. Sallak, C. Simon, and J-F. Aubry. Optimal design of safety instrumented systems. In *ACD, Workshop on Advanced Control and Diagnosis, CDROM papier n°30, Nancy, France*, 2006.
- [Sallak et al., 2007a] M. Sallak, C. Simon, and J-F. Aubry. A fuzzy probabilistic approach for determining safety integrity level. *IEEE Transactions on Fuzzy Systems*, 2007.
- [Sallak et al., 2007b] M. Sallak, C. Simon, and J-F. Aubry. Optimal design of safety instrumented systems : A graph reliability approach. In *7ème édition du congrès international pluridisciplinaire, Qualita2007, Tanger, Maroc*, 2007.

-
- [Sandri, 1991] S. Sandri. *La combinaison de l'information incertaine et ses aspects algorithmiques*. PhD thesis, Université Paul Sabatier, Toulouse, France, 1991.
- [Sawer and Rao, 1994] J. P. Sawer and S. S. Rao. Fault tree analysis of fuzzy mechanical systems. *Microelectronics and Reliability*, 34 :653–667, 1994.
- [Scarpelli and Gomide, 1994] H. Scarpelli and F. Gomide. High level fuzzy petri nets and backward reasoning. In *Proceedings IPMU, Paris, France*, 1994.
- [Shafer, 1976] G. Shafer. *A Mathematical Theory of Evidence*. Princeton University Press, 1976.
- [Sharma and Venkateswaran, 1971] J. Sharma and K.V. Venkateswaran. A direct method for maximizing the system reliability. *IEEE Transactions on Reliability*, 20 :256–259, 1971.
- [Shia and Shenb, 2007] Z-J. Shia and J. Shenb. Convergence of liu storey conjugate gradient method. *European Journal of Operational Research*, 182 :552–560, 2007.
- [Simon and Weber, 2007] C. Simon and P. Weber. Bayesian networks inference algorithm to implement dempster shafer theory in reliability analysis. *Reliability Engineering and System Safety*, 2007.
- [Simon *et al.*, 2006] C. Simon, M. Sallak, and J-F. Aubry. Allocation de SIL par agrégation d’avis d’experts. In *15ème Colloque National de Maîtrise des Risques et Sécurité de Fonctionnement, Lambda-mu 15, Octobre 2006, Lille, France*, 2006.
- [Simon *et al.*, 2007] C. Simon, M. Sallak, and J-F. Aubry. SIL allocation of SIS by aggregation of experts opinions. In *ESREL, Safety and Reliability Conference, Stavanger, Norvège*, 2007.
- [Singer, 1990] D. Singer. A fuzzy set approach to fault tree and reliability analysis. *Fuzzy Sets and Systems*, 34 :145–155, 1990.
- [Siu, 1996] N. Siu. Risk assessment of dynamic systems : an overview. *Reliability Engineering and System Safety*, 43, 1996.
- [Smets and Kennes, 1994] P. Smets and R. Kennes. The transferable belief model. *Artificial Intelligence*, 66 :191–243, 1994.
- [Soman and Misra, 1993] K.P. Soman and K.B. Misra. Fuzzy fault tree analysis using resolution identity. 1, page 193, 1993.
- [Stavrianidis and Bhimavarapu, 1998] P. Stavrianidis and K. Bhimavarapu. Safety Instrumented Functions and Safety Integrity Levels (SIL). *ISA Transactions*, 37 :337–351, 1998.
- [Stavrianidis, 1995] P. Stavrianidis. Improving management of technological risk : a process safety compliance framework. *Risk and Safety Assessment Conference, Hawaii*, 1995.
- [Summers, 2002] A. E. Summers. Viewpoint on ISA TR84.0.02 : simplified methods and fault tree analysis. *ISA Transactions*, 39 :125–131, 2002.
- [Symeonaki and Stamou, 2004] M. A. Symeonaki and G. B. Stamou. Theory of markov systems with fuzzy states. *Fuzzy Sets and Systems*, 143 :427–445, 2004.
- [Tanaka *et al.*, 1983] H. Tanaka, L. T. Fan, F. S. Lai, and K. Toguchi. Fault tree analysis by fuzzy probability. *IEEE Transactions on Reliability*, 32 :453–457, 1983.
- [Tanrioven *et al.*, 2004] M. Tanrioven, Q. H. Wub, D. R. Turner, C. Kocatepe, and J. Wang. A new approach to real time reliability analysis of transmission system using fuzzy markov model. *Electrical Power and Energy Systems*, 26 :821–832, 2004.
- [Tanzi and Delmer, 2003] T. Tanzi and F. Delmer. *Ingénierie du risque*. Lavoisier, France, 2003.
- [Valette and Atabakhche, 1987] R. Valette and H. Atabakhche. Petri nets for sequence constraint propagation in knowledge based approaches. *Concurrency and Petri Nets*, (555-569) :555 – 569, 1987.

- [Valette and Kunzle, 1994] R. Valette and L. A. Kunzle. Réseaux de petri pour la détection et le diagnostic. In *Journées Nationales : Sûreté, Surveillance, Supervision, GdR Automatique, Paris, France*, 1994.
- [Valette *et al.*, 1989] R. Valette, J. Cardoso, and D. Dubois. Monitoring manufacturing system by means of petri nets with imprecise markings. In *IEEE International Symposium on Intelligent Control, Albany N.Y., USA*, 1989.
- [Vesely, 1981] W. E. Vesely. Fault tree handbook. Technical report, U.S Nuclear Regulatory Commission, Washington, 1981.
- [Villemeur, 1987] A. Villemeur. *Evaluation de la fiabilité, disponibilité et maintenabilité des syst ?mes réparables : la méthode de l'Espace des Etats*. Eyrolles, 1987.
- [Villemeur, 1998] A. Villemeur. *Sûreté de fonctionnement des syst ?mes industriels*. Eyrolles, 1998.
- [Vladimir, 2003] F. Vladimir. Fine-grained tournament selection operator in genetic algorithms. *Computing and Informatics*, 22 :143–162, 2003.
- [Walley, 1991] P. Walley. *Statistical reasoning with imprecise probabilities*. Chapman and Hall, 1991.
- [Wang *et al.*, 2004] Y. Wang, H. H. West, and M. S. Mannan. The impact of data uncertainty in determining Safety Integrity Level. *Process Safety and Environmental Protection*, 82 :393–397, 2004.
- [Watson, 1961] H. A. Watson. *Launch control safety study*, volume 1. Bell labs, 1961.
- [Wu, 1998] W. Wu. *Synthèse d'un contrôleur flou par Algorithme Génétique : Application au réglage dynamique des paramètres d'un système*. PhD thesis, Université de Lille 1, 1998.
- [Yalaoui *et al.*, 2005] A. Yalaoui, E. Chatelet, and C. Chengbin. A new dynamic programming method for reliability and redundancy allocation in a parallel-series system. *IEEE Transactions on Reliability*, 54 :254–261, 2005.
- [Zadeh, 1965] L. Zadeh. Fuzzy sets. *Information and Control*, 8 :338–353, 1965.
- [Zadeh, 1968] L. Zadeh. Probability measures of fuzzy events. *J. Math. Anal. Appl*, 23 :421–427, 1968.
- [Zadeh, 1978] L. Zadeh. Fuzzy sets as a basis for a theory of possibility. *Fuzzy Sets and Systems*, 1 :3–28, 1978.
- [Zadeh, 1998] L. Zadeh. Maximizing sets and fuzzy Markov algorithms. *IEEE Trans. on Systems, Man, and Cybernetics*, 28 :9–15, 1998.
- [Zimmermann, 1983] H.J. Zimmermann. Using fuzzy sets in operational research. *European Journal of Operational Research*, 13 :201–216, 1983.
- [Zouhal, 1997] L. Zouhal. *Contribution à l'application de la théorie des fonctions de croyance en reconnaissance des formes*. PhD thesis, Université de Technologie de Compiègne, France, 1997.

AUTORISATION DE SOUTENANCE DE THESE
DU DOCTORAT DE L'INSTITUT NATIONAL
POLYTECHNIQUE DE LORRAINE

o0o

VU LES RAPPORTS ETABLIS PAR :

Monsieur Eric CHATELET, Professeur, Université de Savoie, Troyes

Monsieur Laurent FOULLOY, Professeur, Polytech'Savoie, Annecy-le-vieux

Le Président de l'Institut National Polytechnique de Lorraine, autorise :

Monsieur SALLAK Mohamed

à soutenir devant un jury de l'INSTITUT NATIONAL POLYTECHNIQUE DE LORRAINE,
une thèse intitulée :

**"Evaluation de paramètres de sûreté de fonctionnement en présence d'incertitudes et
aide à la conception : application aux Systèmes Instrumentés de Sécurité"**

en vue de l'obtention du titre de :

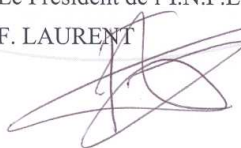
DOCTEUR DE L'INSTITUT NATIONAL POLYTECHNIQUE DE LORRAINE

Spécialité : « **Automatique et traitement du signal** »

Fait à Vandoeuvre, le 02 octobre 2007

Le Président de l'I.N.P.L.,

F. LAURENT



NANCY BRABOIS
2, AVENUE DE LA
FORET-DE-HAYE
BOITE POSTALE 3
F - 54501
VANDŒUVRE CEDEX

Résumé

L'introduction de systèmes instrumentés dédiés aux applications de sécurité impose l'évaluation de leur sûreté de fonctionnement. On utilise généralement pour cela les bases de données de fiabilité génériques. Cependant, le retour d'expérience pour ces systèmes qui présentent en général des défaillances rares est insuffisant pour valider les résultats obtenus. En outre, la collecte de données de fiabilité et leur extrapolation à d'autres composants introduisent des incertitudes. Les travaux de cette thèse portent sur la problématique de la prise en compte des incertitudes relatives aux données de fiabilité des composants pour l'évaluation de la sûreté de fonctionnement des systèmes par le formalisme des sous ensembles flous. La méthodologie proposée est appliquée à l'évaluation des probabilités de défaillance des Systèmes Instrumentés de Sécurité (SIS) en présence de données de fiabilité imprécises. Nous introduisons deux nouveaux facteurs d'importance pour aider le concepteur. En outre, nous proposons une méthodologie d'aide à la conception des SIS basée sur la modélisation par réseaux de fiabilité et l'optimisation par des algorithmes génétiques de la structure des SIS pour le respect des niveaux d'intégrité de sécurité (SIL) exigés.

Mots-clés: Systèmes Instrumentés de Sécurité (SIS), Niveaux d'Intégrité de Sécurité (SIL), incertitudes, arbres de défaillance flous, facteurs d'importance, conception optimale, algorithmes génétiques, réseaux de fiabilité.

Abstract

The use of safety related systems imposes to evaluate their dependability. Laboratory data and generic data are often used to provide failure data of safety components to evaluate their dependability parameters. However, due to the lower solicitation of safety systems in plant, safety components have not been operating long enough to provide statistical valid failure data. Furthermore, measuring and collecting failure data have uncertainty associated with them, and borrowing data from laboratory and generic data sources involve uncertainty as well. Our contribution is to propose a fuzzy approach to evaluate dependability parameters of safety systems when there is an uncertainty about dependability parameters of systems components. This approach is applied to determine the failure probability on demand of Safety Instrumented Systems (SIS) in presence of uncertainty. Furthermore, we present an optimal design of SIS by using reliability graphs and genetic algorithms to identify the choice of components and design configuration in a SIS to meet the required SIL.

Keywords: Safety Instrumented Systems, Safety Integrity Level, uncertainty, fuzzy fault tree, importance measures, optimal design, genetic algorithms, reliability graphs.

