

THESE

présentée à

L'UNIVERSITE DES SCIENCES ET
TECHNOLOGIES DE LILLE

Pour obtenir le titre de

DOCTEUR

En Productique : Automatique et Informatique Industrielle

par

Etienne BAILLY



ETUDE DE LA REGULATION DES LIGNES DE
METRO AUTOMATISEES : APPROCHE PAR
LA LOGIQUE FLOUE

Soutenue publiquement le 11 juillet 1996 devant la commission d'examen :

P. VIDAL	Président du jury
A. TITLI	Rapporteur
R. LAURENT	Rapporteur
A. M. DESODT	Co directeur de recherche
D. JOLLY	Co directeur de recherche
P. BORNE	Examinateur
L. BARANES	Invité

Remerciements

Ce travail a été réalisé au Centre d'Automatique de Lille à l'Université des Sciences et Technologies de Lille, sous la direction du Professeur Pierre Vidal, qu'il soit remercié pour son accueil au sein de son équipe de recherche.

Je me dois de remercier aussi le professeur André Titli du Laboratoire d'Analyse et d'Architecture des Systèmes de Toulouse ainsi que le professeur Robert Laurent de l'IUT de Génie Mécanique et Productique d'Orléans pour l'honneur qu'ils me font en acceptant de juger ce travail et d'en être les rapporteurs.

Ce travail n'existerait sans doute pas sans l'aide de Daniel Jolly et Anne Marie Desodt, codirecteurs de l'équipe Système Homme Machine et Décision du Centre d'automatique de Lille. Ils ont su par des réunions hebdomadaires co-diriger mes travaux en me donnant sans cesse des réponses à mes questions, des nouvelles voies de réflexion et de nombreuses solutions aux problèmes posés. Leurs disponibilités et leurs compétences m'ont permis d'avancer aussi bien du point de vue professionnel que personnel.

Je remercie Pierre Borne, Directeur Scientifique de l'Ecole Centrale de Lille, Vice-Président d'IEEE/SMC et de l'IMACS, pour l'intérêt qu'il me témoigne en acceptant de participer à ce jury.

Monsieur Régis Lardennois, directeur de la recherche à Matra Transport, pour me faire l'honneur de s'intéresser à mes travaux.

Monsieur Lionel Baranes, directeur général adjoint de l'INRETS pour avoir su montrer son intérêt à ces travaux et en faisant partie de ce jury.

Je remercie aussi Saïd Hayat, chargé de recherche à l'INRETS pour avoir su m'accueillir à l'INRETS et de m'avoir ouvert au monde du transport. Il a su par ses conseils orienter mes travaux et m'a permis d'arriver à ces résultats.

Joaquim Rodriguez, chargé de recherche à l'Inrets, par ses nombreuses explications sur la simulation et le langage orienté objet, m'a initié un domaine où j'étais plus que novice.

La région Nord Pas de Calais qui s'est associée à l'INRETS pour me donner les moyens nécessaires pour réaliser ces travaux ainsi que le Groupement Régional de la Recherche sur les Transports pour m'avoir fait confiance en vue de la réalisation de ces travaux.

C'est à l'INRETS que j'ai passé la plus grande partie de mon temps. Je me dois de remercier l'ensemble du personnel de Villeneuve d'Ascq pour sa sympathie ; plus particulièrement Christian, Michel, Bernard, Daniel et Jean Yves pour ne citer que certains.

Je remercie surtout mes parents pour leur support, leur compréhension et leur dévouement durant ces longues années et d'avoir toujours cru en moi.

Enfin, je terminerai par Patrick qui m'a supporté durant 4 ans dans le même bureau et qui est devenu bien plus qu'une relation professionnelle. Et surtout la personne rencontrée à l'INRETS qui partage désormais ma vie, Claire ma femme.

Table des matières

Introduction générale	7
Chapitre 1. Fonctionnement d'un système de transport automatisé	9
1. Introduction	9
2. Description d'un système de transport automatisé à haute densité.....	11
2.1. Conduite des trains	12
2.1.1. Contraintes générales	12
2.1.2. Contraintes circonstancielles	12
2.1.3. Données de programme du pilotage.....	13
2.1.4. Cas du VAL de Lille	15
2.2. La Régulation.....	16
2.2.1. Objectifs.....	16
2.2.2. Approche cinématique de la régulation.....	17
2.2.3. Approche de séquentialisation temporelle	20
2.2.4. Régulation dans le cas du canton fixe.....	23
2.3. Problématique du service partiel	24
2.3.1. Services partiels envisageables	25
2.3.2. Service partiel ou lignes imbriquées.....	26
2.3.3. Actions envisageables au niveau des terminus partiels.....	27
2.3.4. Conclusion	28
3. Qualité de service d'un système de transport automatisé.....	28
3.1. Critères de qualité.....	29
3.1.1. Un faible temps d'attente en station	29
3.1.2. Intervalle attractif en heures creuses.....	30
3.1.3. Adaptation facile de l'offre et de la demande.....	30
3.1.4. Vitesse commerciale élevée.....	31
3.2. Indicateurs de la qualité de service.....	31
3.3. Paramètres introduisant incertitude et imprécision	33
4. Conclusion	34

Chapitre 2. Apport de la logique floue dans la commande de systèmes complexes	37
1. Introduction	37
2. La logique floue et ses applications à la commande floue et aux systèmes à base de connaissances.....	38
2.1. Principes de base de la logique floue	38
2.2. La commande floue.....	40
2.2.1. Typologie des contrôleurs flous	40
2.2.2. Bases de règles et définitions	42
2.2.3. Interface de fuzzification.....	42
2.2.4. Inférence floue.....	42
2.2.5. Défuzzification.....	43
2.3. Représentation de l'incertain.....	44
2.3.1. Représentation statistique.....	44
2.3.2. Représentation probabiliste.....	45
2.3.3. Représentation possibiliste	45
3. Décomposition hiérarchique.....	47
3.1. Décomposition des systèmes complexes.....	47
3.2. Extension aux systèmes flous	49
3.3. Décomposition hiérarchique de système de gestion d'une ligne de métro	51
Chapitre 3. Application de la logique floue à la régulation du VAL de Lille.....	53
1. Introduction	53
2. Commande floue.....	53
2.1. Définition du modèle	53
2.1.1. Structure du module de commande.....	57
2.1.2. Les fonctions d'appartenance	58
2.1.3. Les règles.....	60
2.2. Simulation sur la ligne 2 du VAL de Lille.....	61
2.2.1. Mesure des retards des rames amont et aval.....	62
2.2.2. Calcul du retard estimé.....	63
2.2.3. Résultats.....	65
2.2.4. Répercussion du résultat de l'inférence floue sur les rames.....	66
2.3. Conclusion.....	67
3. Supervision floue.....	67
3.1. Introduction.....	67
3.2. Scénarii de pannes	68

3.2.1. La rame perturbée est arrêtée en station.....	69
3.2.2. La rame perturbée est arrêtée en inter-station.....	69
3.3. Introduction d'une notion de risque	70
3.4. Calcul du "risque" et du "non-risque"	72
3.4.1. Retard de la rame	73
3.4.2. Temps de parcours de la rame	73
3.5. Différents indices de risque.....	74
3.6. Extension de la notion de calcul du risque	76
3.6.1. Si la rame (i-1) est perturbée	76
3.6.2. Si la rame (i-1) est en marche normale.....	76
3.6.2. Si la rame (i-1) est en anticollision ou en sur-stationnement.....	77
3.7. Comparaison entre les deux mesures de distance	77
3.7.1. La méthode de Dubois et Prade	78
3.7.2. La méthode de Jain	79
3.7.3. La méthode des surfaces.....	80
3.8. Application de l'algorithme au service partiel.....	83
4. Conclusion	87
Chapitre 4. Plateforme de simulation et résultats	89
1. Introduction	89
2. Simulation de la ligne 1 du VAL de Lille	90
2.1. Simulation d'une perturbation.....	91
2.2. Comparaison avec l'algorithme VAL.....	93
3. Simulation de la ligne 2 du VAL de Lille	95
3.1. Simulation d'une perturbation.....	96
Comparaison avec l'algorithme VAL	97
3.2. Simulation de deux perturbations	98
Comparaison avec l'algorithme VAL	99
4. Simulation de la régulation avec indice de risque	100
Application au service partiel.....	103
5. Réalisation d'un estimateur de la durée de perturbation d'un train.....	106
5.1. Différents moyens de réaliser une prédiction de la durée de la panne	107
Calcul d'estimateurs	107
5.2. Approche retenue	108
5.2.1. Résolution par la théorie des probabilités.....	111
5.2.2. Résolution par la théorie des possibilités	112

5.2.3. Intervention du temps dans les calculs.....	114
5.2.4. Simulation de l'estimateur de la durée de la panne	115
6. Conclusion	118
Conclusion générale	119
Bibliographie	123
Annexe 1. Indicateurs de qualité de service	127
Annexe 2. Rappels sur les sous ensembles flous.....	133
Annexe 3. Description de la plateforme de simulation	141

Introduction générale

Cette étude a été réalisée grâce à une collaboration entre l'Institut National de REcherche sur les Transports et leur Sécurité (INRETS) dans l'unité de recherche Evaluation des Systèmes de Transport Automatisés et de leur Sécurité (ESTAS) et l'Université des Sciences et Technologies de Lille (USTL) au Centre d'Automatique de Lille (CAL).

Afin de répondre aux nouveaux défis de l'humanité, les concepteurs sont amenés à développer des systèmes de plus en plus complexes. Cette complexité est générée par le fait que l'on ne peut plus négliger des phénomènes jusqu'à présent considérés comme secondaires. Par conséquent, la représentation mathématique de tels systèmes aboutit à la création d'un modèle généralement fortement non linéaire et/ou incertain.

La théorie des sous ensembles flous introduite par L. Zadeh permet de tenir compte de l'imprécision du monde réel. Les commandes floues n'utilisent alors aucun modèle explicite du système. On évite par conséquent les problèmes d'une modélisation mathématique complexe.

Un système de transport est par définition complexe. En effet, il est composé d'une multitude de sous systèmes qui indépendamment les uns des autres sont modélisables, mais il est très difficile de prévoir les conséquences d'une action sur l'ensemble du système. De plus, la présence de l'homme, en tant qu'utilisateur, rend le système global difficilement modélisable.

Ainsi, nous nous sommes intéressés à l'étude de l'apport de la théorie des sous-ensembles flous dans la régulation des lignes de métro automatisées. Cette étude s'intègre dans le programme INRETS "Amélioration de la qualité de l'offre des transports collectifs". Toutes les actions menées à l'INRETS pour améliorer le confort, l'accessibilité des matériels, l'information, l'exploitation, les correspondances, les performances et la sécurité peuvent se rattacher à l'objectif global d'amélioration de la qualité. Ces études étant, bien entendu, menées tout en gardant présent à l'esprit le critère économique. Le but de notre étude s'insère dans l'amélioration des techniques d'exploitation concernant les diverses configurations des lignes de transport en commun.

Un certain nombre de méthodes pour réguler le trafic des rames de métro ont déjà fait leurs preuves sur les lignes automatisées en service actuellement. Le but de notre travail est

donc de comparer ces méthodes avec une nouvelle approche basée sur la logique floue. Cette approche est justifiée par le fait qu'une ligne de métro est un système complexe et par la présence de l'homme qui en tant qu'utilisateur rend l'environnement incertain et imprécis.

Puisqu'il est difficilement concevable d'utiliser une ligne de métro réelle pour valider notre régulation, nous avons utilisé une simulation de lignes de métro programmée en langage orienté objet développé à l'INRETS - ESTAS. Cette simulation permettra de mettre au point notre régulation et de la comparer avec les autres méthodes couramment utilisées.

Dans une première partie, nous décrirons le fonctionnement d'un système de transport automatisé avec ses principales fonctions. Nous dégagerons les faiblesses de la régulation existante en montrant qu'un certain nombre de paramètres ne sont pas pris en compte.

Le chapitre 2 introduira les concepts de base de la théorie des sous-ensembles flous, notamment la commande floue. Dans notre cas, il est nécessaire de rappeler le principe des systèmes à base de connaissances car les données d'exploitation des lignes de métro des années précédentes peuvent nous apporter des informations pertinentes sur les actions à mener lors de perturbations sur ces lignes. Ensuite, après avoir rappelé les principes de hiérarchisation des systèmes complexes développés pour les commandes classiques, nous montrerons l'extension de ces principes aux systèmes complexes flous.

L'application de ces concepts en vue de la régulation floue d'un système de transport automatisé (le Véhicule Automatisé Léger ou VAL) sera exposée au chapitre 3. Ce chapitre présentera les deux niveaux de régulation ainsi que les modèles flous de chaque niveau. Nous montrerons que ces modèles peuvent être étendus en vue de réguler des lignes nouvelles, notamment des lignes utilisant le service partiel.

Enfin, le chapitre 4 décrira la plate-forme de simulation des lignes de métro automatisées utilisée pour les essais de notre nouveau modèle de régulation. Il présentera les résultats de simulation pour les différents niveaux de la régulation floue et comparera ces résultats avec la régulation existante en montrant les avantages qu'apporte la nouvelle régulation. Ces résultats seront obtenus grâce à la simulation de deux lignes : la ligne 1 du VAL, déjà en exploitation depuis 1983 et la ligne 2 du VAL de Lille, prochainement en exploitation. Ensuite, nous montrerons comment il est possible de réaliser un estimateur de durée d'une panne sur la rame afin d'améliorer l'efficacité de notre modèle de régulation.

CHAPITRE 1

FONCTIONNEMENT D'UN SYSTEME DE TRANSPORT AUTOMATISE

1 . Introduction

De nos jours, des millions d'usagers utilisent quotidiennement les systèmes de transport en commun. Dans les zones urbaines à forte densité de population, ils véhiculent un public de plus en plus important et deviennent un service indispensable dans la vie des usagers. L'arrêt même provisoire de ces transports en commun perturbe désormais fortement les usagers. Cet arrêt provoque un absentéisme important sur les lieux de travail et la saturation du trafic routier dans les grandes agglomérations.

Nous pouvons classer ces systèmes de transport en commun en plusieurs catégories :

- Les systèmes de transport non guidés
- Les systèmes de transports guidés notamment le tramway et le métro
- Les systèmes hybrides
- **Les systèmes de transport non guidés**

Le système de transport non guidé le plus développé est le réseau d'autobus. Il a pour principal avantage de permettre une grande souplesse dans le choix des itinéraires. L'absence d'une infrastructure dédiée induit un coût d'exploitation peu important et permet de desservir des zones à faible densité de population où la construction d'une infrastructure importante ne serait pas rentable. Les principaux inconvénients des systèmes de transport non guidés sont leur grande inertie face à un événement instantané, leur totale dépendance vis à vis de la circulation routière, leur gestion du personnel lourde et difficile. Pour utiliser les avantages de ces systèmes en supprimant les inconvénients, les exploitants des grandes agglomérations utilisent ces systèmes de transport en commun en complément d'autres systèmes de transport guidés. Ils

permettent de transporter les usagers des zones peu fréquentées vers celles plus fréquentées où existe un système de transport guidé.

- Les systèmes de transports guidés notamment le tramway et le métro

Ces systèmes doivent être utilisés en site propre et, de ce fait, ne sont pas totalement dépendants de la circulation. Le tramway est plus ou moins dépendant de la circulation automobile puisqu'il emprunte généralement les voies routières. Par contre le métro possède un site propre ; il est le seul moyen de transport urbain à posséder cette caractéristique qui lui donne un coût de construction plus élevé que d'autres modes de transport urbain.

- Les systèmes hybrides

Certaines villes ont choisi un système hybride qui comprend une partie en site propre et une partie en circulation normale (GLT de la ville de Caen). Mais l'avantage de ce système n'a pas été encore complètement démontré.

Le choix d'un système de transport doit dépendre du nombre de personnes susceptibles de l'emprunter quotidiennement. Il est évident qu'un métro ne pourrait être rentable si l'on veut relier deux villages ayant une population peu importante.

Le métro est, donc réservé aux villes possédant une forte densité de population susceptible de l'emprunter tous les jours. On parlera de système de transport en commun à haute densité.

Une ligne de métro peut être comparée à une ligne moderne de chemin de fer. Elle en possède les caractéristiques principales :

- Exploitation en site propre, souterrain ou non.
- Limitation stricte des mouvements autorisés (gestion de sécurité).

Cependant, elle s'en distingue par :

- Des dimensions de réseaux réduites. Un réseau typique reste limité aux dimensions d'une ville. Il est en outre caractérisé par des temps de parcours généralement petits entre les points d'arrêts successifs.

- Des structures de ligne simples. Chaque réseau est constitué d'un ensemble de lignes de structures simples physiquement indépendantes les unes des autres. Sur une ligne de structure simple, un seul itinéraire permet la liaison entre deux points d'arrêt.

- Des fréquences d'exploitation élevées. Elles sont autorisées par les structures des lignes et sont nécessaires pour répondre à la demande des passagers.

- Le comportement des passagers. Ceux-ci arrivent aléatoirement en station et embarquent généralement dans le premier véhicule qui s'y arrête.

- La conduite des véhicules. Elle peut être confiée à un ordinateur. Dans ce cas, on parle de lignes de métro automatisées.

Dans notre étude, nous ne prendrons en compte que les lignes de métro automatisées. Il est à noter que les systèmes d'aide à l'exploitation ne sont pas réservés aux lignes de métro mais qu'ils sont de plus en plus utilisés dans les autres systèmes de transport en commun (tramway et bus).

2. Description d'un système de transport automatisé à haute densité

Les fonctions assurées par un système de transport automatisé ont pour but de répondre à une mission comportant deux aspects principaux :

- Répondre aux besoins d'exploitation : transport des passagers et souplesse d'exploitation. Ce type de besoin est dit fonctionnel.

- Assurer la sécurité des usagers et du matériel : cet aspect regroupe toutes les contraintes dites sécuritaires.

Pour assurer le transport des usagers, les rames ont besoin d'un système de conduite permettant de transporter les passagers en toute sécurité entre deux stations. Ce mouvement des trains sur l'ensemble de la ligne est réalisé par une fonction dite conduite des trains qui est assujettie à un ensemble de contraintes sécuritaires.

2.1. Conduite des trains

La conduite des trains en mode automatique gère le mouvement des trains nécessaire à l'accomplissement de la mission dans le respect d'un certain nombre de contraintes.

Les mouvements de trains correspondent :

- à la gestion du garage des trains,
- au parcours de la ligne avec arrêts en stations,
- aux manoeuvres de terminus.

Les contraintes à respecter sont d'ordre général, local ou circonstanciel.

2.1.1. Contraintes générales

- La sécurité des passagers et des équipements.
- La qualité de service.

Ces contraintes doivent à tout moment être prises en compte mais leur impact doit souvent être défini localement (par exemple lors de l'arrêt en station).

2.1.2. Contraintes circonstancielles

Elles permettent de continuer l'exploitation dans un mode dégradé, à l'apparition d'événements occasionnels. Par exemple, en cas de vent fort, on peut imposer une réduction de vitesse aux trains sur les interstations aériennes. La plupart des contraintes peuvent être traduites par une limitation de vitesse pour une rame donnée, à un instant donné et en un point donné. Cette limite de vitesse, locale et instantanée, constitue la principale variable d'entrée de la conduite des trains. Elle est élaborée, pour chaque train, en fonction de la localisation des rames puis transmise au train. Le train connaissant à tout instant la vitesse maximale qui lui est autorisée, gère son mouvement à l'aide de sa fonction de pilotage. Ce pilotage est lui-même surveillé en permanence par le contrôle sécuritaire de marche qui vérifie le respect des contraintes de sécurité. Le synoptique général de la fonction de conduite des trains est donné par le schéma représenté sur la figure 1.1.

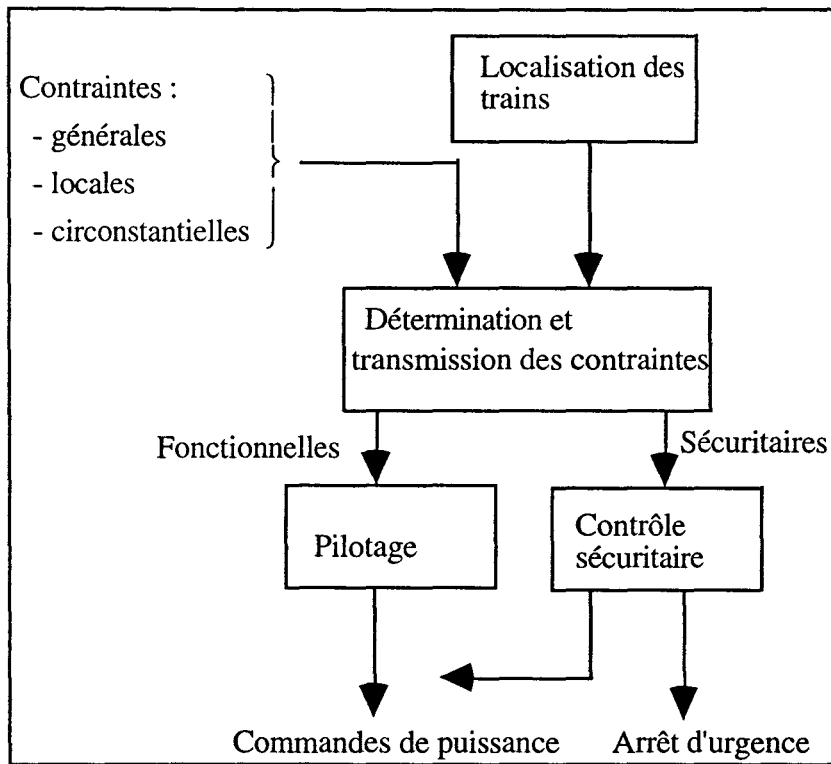


Figure 1.1. Synoptique de la conduite des trains

2.1.3. Données de programme du pilotage

Ces contraintes sont exprimées sous forme d'une limitation de vitesse à l'exception de :

- la limitation de l'accélération et de la décélération à $1,3 \text{ m/s}^2$,
- la limitation du jerk, qui est la variation de l'accélération, à $\pm 0,65 \text{ m/s}^3$.

2.1.3.1. Programme de vitesse limite

On définit, en chaque point de la ligne, une vitesse limite comme la vitesse maximale imposée par les caractéristiques de la ligne. Cette vitesse dépend de :

- la géométrie de la voie,
- la position des stations,
- la position des aiguillages,
- la position des zones de garage.

2.1.3.2. Programme de vitesse d'arrêt sur anticollision

Différentes fonctions d'anticollision sont possibles pour les lignes de métro. La plupart fonctionnent selon le principe du canton. Le canton peut être fixe, mobile ou mobile déformable. Par exemple, les lignes du VAL de Lille sont découpées en cantons fixes. Un canton fixe est une portion de voie sur laquelle on n'autorisera la présence que d'une seule rame à la fois. Cela permet d'éviter qu'une rame n'en percute une autre située devant elle. Le canton mobile est défini comme une zone mobile entre deux rames. Il dépend de la distance séparant les deux rames ainsi que leurs vitesses. La ligne D du métro de Lyon (MAGGALY) a adopté le système du canton mobile déformable [HAY 93] et [BAI 93].

On associe, à chaque canton, un canton amont et un canton aval suivant le sens de marche. L'objet de la fonction de localisation des trains est d'indiquer la présence d'une rame sur un canton. Cette information permet d'interdire au train de quitter son canton si le canton aval est occupé. Lorsque le canton aval est occupé, on définit un programme de vitesse d'arrêt sur anticollision (programme en mode perturbé) lequel provoque l'arrêt des rames sur le canton à une distance suffisante de la limite aval du canton pour assurer en toute sécurité la non-pénétration sur le canton suivant.

Lors de fonctionnements dégradés, la conduite manuelle par un agent d'exploitation est possible. Ce pilotage entraîne la suppression de la fonction d'anticollision définie ci-dessus. Cette fonction est gérée par l'agent pilotant le train et est possible grâce à une limitation de la vitesse à 18 km/h. Les contraintes géométriques sur la voie garantissent que le pilote peut voir suffisamment tôt un obstacle sur la voie de façon à commander à temps l'immobilisation en freinage d'urgence du train. Le programme en mode perturbé n'est utilisé que lorsque le canton suivant est occupé ; il est considéré comme une contrainte de vitesse s'ajoutant aux limitations intégrées dans le programme de vitesse limite. Dès que le canton aval est libéré, le programme en mode perturbé devient sans effet.

2.1.3.3. Vitesse de régulation

Il est nécessaire de distinguer la vitesse maximale autorisée, résultat des contraintes énoncées ci-dessus, de la vitesse nominale d'une rame. Cette disposition permet de donner à une rame retardée par rapport à son horaire nominal une consigne de vitesse supérieure et donc de rattraper son retard. Cette consigne s'appelle vitesse de régulation ; elle est définie à chaque départ de station pour une interstation. Il s'agit d'une consigne fonctionnelle dont le respect concerne la qualité de service mais non la sécurité.

2.1.3.4. Vitesse imposée par le Poste de Commande Centralisé (P.C.C.)

Les opérateurs du P.C.C. ont la possibilité d'envoyer des consignes fonctionnelles de vitesse limite. Cette procédure permet le ralentissement systématique d'une ou de plusieurs rames, lors de travaux en ligne par exemple ou lors de la mise en oeuvre de certains modes dégradés.

2.1.3.5. Profil de vitesse

A son départ d'une station, la rame doit arriver le plus rapidement possible à la prochaine station. Pour cela, elle démarre d'une station avec une vitesse initiale nulle et doit atteindre le plus rapidement possible (compte tenu de ses performances) la vitesse de consigne et tenir cette vitesse le plus longtemps possible. A un moment judicieusement choisi pour permettre l'arrêt à la station suivante, la vitesse de consigne passera à zéro et le véhicule amorcera une phase de décélération. Ce type de pilotage est un pilotage par le jerk (dérivée de l'accélération). A chaque instant, le jerk doit être déterminé pour permettre à la rame d'atteindre la station suivante le plus rapidement possible. Le profil du mouvement de la rame entre deux stations ainsi que le jerk correspondant sont représentés sur la figure suivante.

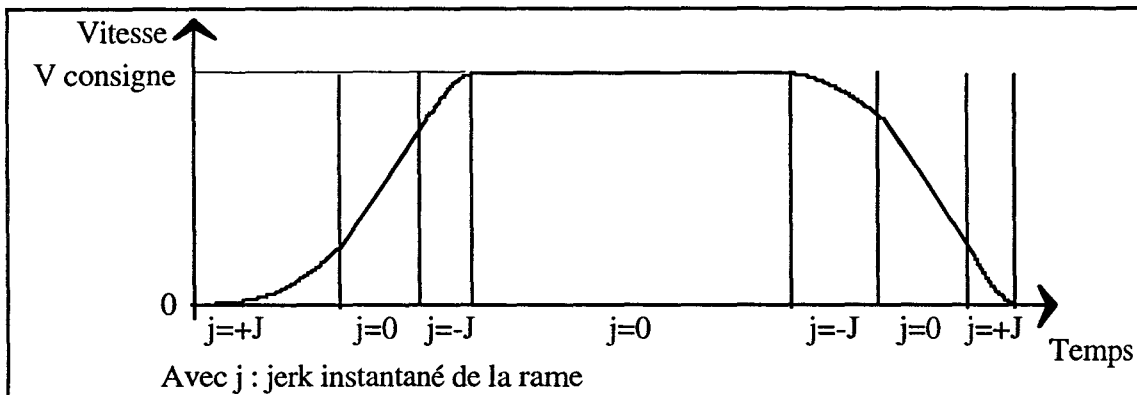


Figure 1.2. Profil type de vitesse sur une interstation

2.1.4. Cas du VAL de Lille

Le VAL de Lille fut le premier mode de transport en automatisme intégral. L'origine du système se situe dans les années 1970-71. Ce système a été développé conjointement par l'EPALE (Etablissement Public d'Aménagement de la Ville Nouvelle de Lille - Est) et par l'Université des Sciences et Technologiques de Lille dans l'équipe du Professeur Gabillard. Il s'agissait de mettre en oeuvre un moyen de transport destiné à relier le centre de Villeneuve d'Ascq à la gare de Lille avec un niveau de service suffisamment élevé pour inciter les habitants

à l'utiliser de préférence à leur véhicule personnel, et ce tout au long de la journée. Il fallait, par conséquent, créer un système doté d'une excellente vitesse commerciale, même aux heures de pointe, et gardant une bonne fréquence aux heures creuses. Compte tenu des trafics modérés à écouler (6000 personnes par sens à l'heure de pointe, 15000 à terme) les fortes fréquences souhaitées permettaient un véhicule de gabarit réduit. Ces fréquences importantes (une rame toutes les minutes à l'heure de pointe) imposaient d'un point de vue économique l'automatisme intégral. Le système devait être en site propre car, pour être compétitif avec la voiture, sa vitesse commerciale devait être supérieure à 30 km/h.

L'ensemble des données définissant la conduite permet de faire circuler les rames d'un point à un autre le plus vite possible et en toute sécurité. L'ensemble de ces mouvements doit être néanmoins coordonné si l'on veut permettre à ces rames d'assurer le transport de personnes. La coordination de tous les mouvements des trains est assurée par la fonction régulation.

2.2. La Régulation

Avant de modifier la régulation existante, il est impératif d'analyser ses objectifs pour y répondre dans le modèle proposé.

2.2.1. Objectifs

En pratique, le système formé par la ligne et les véhicules est soumis à des perturbations qui en altèrent le fonctionnement. L'effet de ces perturbations ne se limite ni aux quais où elles se produisent, ni aux véhicules concernés mais se propage à l'ensemble du réseau. Le comportement du trafic sur une ligne est naturellement instable. Lors d'un incident sur le réseau, l'usager subit un accroissement de la durée du déplacement, un inconfort accru du fait de la distribution inégale des charges entre les véhicules, et un temps d'attente plus ou moins aléatoire. Ce dernier élément constitue le point le plus négatif du transport en commun vis à vis de l'usage du véhicule individuel. Par ailleurs, l'instabilité de fonctionnement entraîne pour l'exploitant un surcoût engendré par la nécessité de prévoir des réserves de temps (temps de battement) pour la régulation, et une perte de capacité générée par les retards des véhicules qui s'accumulent en ligne. Ainsi, le maintien d'un niveau de qualité d'exploitation de transport en commun nécessite de limiter l'incidence des perturbations de manière à offrir à la clientèle un service qui réponde de manière satisfaisante à ses attentes.

Les principaux objectifs de la régulation de trafic, tant quantitatifs que qualitatifs, sont les suivants :

- Amélioration de la fiabilité des horaires (ponctualité) et de la régularité.
- Amélioration de la productivité (ou de la vitesse d'exploitation).
- Amélioration du confort des usagers (diminution des temps d'attente, meilleure répartition du nombre de passagers par véhicule et limitation des accélérations et des décélérations).

Une bonne régulation du trafic consiste à :

- adapter la fréquence de passage des véhicules à la demande journalière,
- faire absorber par l'ensemble des véhicules en ligne à un instant donné une éventuelle perturbation et permettre à la(aux) rame(s) retardée(s) de rattraper l'horaire théorique de référence, ceci afin de garantir un intervalle constant entre les véhicules.

Il y a deux façons différentes de concevoir la régulation de trafic de lignes de métro automatisées :

- La première prend en compte les équations mécaniques du mouvement. A tout instant la position, la vitesse, l'accélération et la dérivée de l'accélération sont connues et elles décrivent l'évolution du système. Ce type de modélisation mène à une simulation très précise car elle fournit à tout instant l'évolution de toutes les rames. Cette régulation précise n'est pas aisée à mettre en pratique à cause du volume d'informations nécessaires.
- La seconde caractérise l'ensemble des mouvements d'un véhicule uniquement par la séquence des instants de départ des quais successifs rencontrés sur la ligne. Ce modèle n'est pas aussi précis que celui utilisant les équations mécaniques des mouvements des véhicules. Par contre, il est beaucoup plus aisé à mettre en oeuvre.

2.2.2. Approche cinématique de la régulation

2.2.2.1. Modélisation de la régulation de trafic

Cette approche repose sur la connaissance, à tout instant et pour chaque rame, des différents paramètres cinématiques qui la caractérisent à savoir sa position, sa vitesse, son accélération et son jerk.

Ces différentes variables sont reliées entre elles par des équations. L'objectif est de maintenir, par des actions de commande, l'ensemble des paramètres cinématiques réels proches des valeurs établies pour un mode de fonctionnement non perturbé. La détermination d'une commande optimale nécessite la description des équations de mouvement des rames.

$$v_i(t) = \frac{d}{dt}(x_i(t))$$
$$m_i \frac{d}{dt}(v_i(t)) = -g_i(v_i(t)) + f_i(t)$$

avec : $x_i(t)$, position du véhicule (i) à l'instant t

$v_i(t)$, vitesse du véhicule (i) à l'instant t

m_i , masse en charge du véhicule (i)

$f_i(t)$, force appliquée au véhicule (i) résultat de la commande

$g_i(v_i(t))$, force de frottement du véhicule (i) à l'instant t

Les équations établies permettent de décrire en fonction du temps l'évolution de l'ensemble des paramètres cinématiques de la rame. Ce type de modélisation nous donne une simulation détaillée du mouvement et permet de connaître l'état du système à tout instant.

2.2.2.2. Algorithmes de régulation

L'objectif de la régulation consiste à minimiser l'écart entre les valeurs mesurées et celles déterminées en mode de fonctionnement nominal. Les actions de commandes permettant la régulation seront calculées à partir de la modification du jerk sachant que les paramètres cinématiques sont reliés entre eux par les équations suivantes :

$$a_i(t) = j_i(t) * t + a_i(t_0)$$
$$v_i(t) = \frac{1}{2} j_i(t) * t^2 + a_i(t_0) * t + v_i(t_0)$$
$$x_i(t) = \frac{1}{6} j_i(t) * t^3 + \frac{1}{2} a_i(t_0) * t^2 + v_i(t_0) * t + x_i(t_0)$$

avec : $j_i(t)$, jerk du véhicule (i) à l'instant t

$a_i(t)$, accélération du véhicule (i) à l'instant t

$a_i(t_0)$, accélération initiale du véhicule (i)

$v_i(t_0)$, vitesse initiale du véhicule (i)

$x_i(t_0)$, position initiale du véhicule (i)

Les sous objectifs de la régulation consistent à :

- Minimiser l'écart entre la position réelle $d_i(t)$ et la position théorique à l'instant t , soit :

$$\delta d_i(t) = d_i(t) - x_i(t)$$

- Minimiser l'écart entre la vitesse réelle et la vitesse théorique $v(t)$, soit :

$$\delta v_i(t) = v_i(t) - v(t)$$

- Minimiser la dérive de l'intervalle w_i entre deux rames successives soit :

$$\delta w_i(t) = x_i(t) - x_{i+1}(t)$$

Le choix d'un critère Q nous amène à pondérer les différents critères.

La solution au problème posé est obtenue après résolution de ces équations. La commande optimale à appliquer, le jerk en l'occurrence, est ainsi déterminée.

L'approche cinématique a l'avantage de pouvoir décrire l'évolution du système de façon très fine par des simulations appropriées et de déterminer des commandes théoriques efficaces.

Elle présente néanmoins certaines faiblesses pour la mise en oeuvre sur des cas réels, nous citerons principalement :

- Un coût élevé du système qui permet l'identification en temps réel de chaque rame.
- Les performances théoriques sont déterminées à partir d'estimations de grandeurs caractérisant l'interface rame - usagers à quai. Ces grandeurs sont difficilement maîtrisables.
- La modélisation de la force de traction en fonction de la vitesse est complexe : elle est généralement de forme non linéaire. Les hypothèses simplificatrices nécessaires à la modélisation provoquent un décalage par rapport au comportement réel du système.

La seconde famille de modèles se base sur les instants de départ de chacune des rames sur l'ensemble des quais.

2.2.3. Approche de séquentialisation temporelle

2.2.3.1. Modélisation de la régulation de trafic

Le principe de la séquentialisation temporelle consiste à établir des modèles mathématiques linéaires ou non linéaires caractérisant le transfert d'une rame entre deux quais successifs [VAN 89].

La modélisation proposée entraîne les hypothèses suivantes :

- Le temps de stationnement à quai varie linéairement avec le nombre de passagers à embarquer dans ce véhicule.
- Le nombre de passagers à embarquer est proportionnel à l'intervalle séparant deux départs successifs du même quai.
- Les caractéristiques de modélisation ne changent pas durant la période considérée.

On aboutit à deux équations. La première déterminant l'instant de départ du véhicule (i) du quai (k+1):

$$t_{k+1}^i = t_k^i + r_k^i + s_{k+1}^i$$

Avec t_k^i : Instant de départ du véhicule (i) du quai (k)

r_k^i : Temps de parcours du véhicule (i) entre les quais (k) et (k+1)

s_{k+1}^i : Temps d'arrêt à la station (k+1).

La deuxième équation est le temps de parcours entre les quais (k) et (k+1) :

$$r_k^i = R_k + u_k^i + W_{1k}^i$$

Avec R_k : Temps de parcours nominal entre les quais (k) et (k+1)

u_k^i : Le résultat de la commande appliquée au véhicule (i) entre les quais (k) et (k+1).

W_{1k}^i : Représente un terme de perturbation caractéristique de l'écart existant entre le temps de parcours nominal et le temps de parcours réel du véhicule (i) entre les quais (k) et (k+1).

On obtient grâce aux hypothèses le temps d'arrêt du véhicule (i) à la station (k+1).

$$s_{k+1}^i = C_{k+1}(t_{k+1}^i - t_k^i) + D_{k+1} + W_{2k}^i$$

Avec C_{k+1} : Facteur de proportionnalité liant le nombre de passagers à embarquer à l'intervalle de temps écoulé entre deux départs

D_{k+1} : Temps d'arrêt minimal au quai (k+1)

W_{2k}^i : Représente l'effet des perturbations et des imprécisions de modélisation.

2.2.3.2. Les équations caractéristiques

Ces modèles nous permettent de définir les équations qui serviront pour la mise en place de la régulation.

a) Equation aux différences

Nous pouvons définir l'écart horaire entre les instants de départ réel et théorique au quai (k) de la rame (i) :

$$x_k^i = t_k^i - T_k^i$$

Où T_k^i est l'horaire théorique de départ du véhicule (i) de la station (k), correspondant aux conditions idéales de fonctionnement sans aucune perturbation.

La combinaison des équations précédemment définies, permet d'établir l'équation aux différences qui prend la forme suivante :

$$f(x_{k+1}^i, x_{k+1}^{i-1}, x_k^i) = u_k^i + W_k^i$$

avec W_k^i fonction des deux termes d'erreur de modélisation W_{1k}^i et W_{2k}^i .

Elle nécessite de déterminer l'ensemble des conditions initiales du fonctionnement du système.

b) Equation aux intervalles

De la même manière, on peut exprimer l'intervalle de temps entre les instants de départ des véhicules (i) et (i-1) du quai (k) :

$$\delta t_k^i = t_k^i - t_k^{i-1}$$

On peut définir les écarts entre les résultats des commandes appliquées et entre les erreurs de modélisation sur le système :

$$\delta u_k^i = u_k^i - u_k^{i-1} \text{ et } \delta W_k^i = W_k^i - W_k^{i-1}$$

Ce qui permet d'établir, par substitutions successives, l'équation aux intervalles :

$$g(\delta t_{k+1}^i, \delta t_{k+1}^{i-1}, \delta t_k^i) = \delta u_k^i + \delta W_k^i$$

La généralisation de l'équation aux intervalles et de l'équation aux dérivations, permet de décrire l'évolution globale du système étudié.

Différents modèles peuvent être établis selon que :

- l'on se base sur l'équation aux intervalles pour la généralisation, ou au contraire sur l'équation aux différences,
- l'on désire décrire la dynamique quai par quai, rame par rame ou par une combinaison des deux aspects.

Les équations d'état décrivant l'évolution globale du système, indépendamment du modèle choisi, auront la forme générale suivante :

$$X(j+1) = A.X(j) + B.U(j) + C.W(j)$$

où les éléments des matrices A, B et C sont définis par l'ensemble des paramètres décrits précédemment, et où X(j) est fonction des écarts horaires, des intervalles entre les départs et écarts sur charge.

2.2.3.3. Algorithmes de régulation

Les algorithmes de régulation ont pour principaux objectifs :

- Une régulation sur horaires permettant de respecter au mieux l'horaire théorique.
- Une régulation sur intervalles visant à maintenir un intervalle constant entre les rames successives.
- Une régulation mixte sur horaire et intervalle.

La mise en place des algorithmes nécessite la détermination de critères de performances quadratiques à minimiser. Ces critères varient selon que :

- l'on connaît ou non les horaires théoriques de référence,
- l'on choisit l'un ou l'autre des modèles généraux décrivant la dynamique du système,
- l'on se place dans l'une ou l'autre des optiques de régulation.

Ce modèle est un modèle théorique basé sur un certain nombre d'hypothèses. Or en fonctionnement perturbé, ces hypothèses peuvent s'avérer fausses - d'où l'extrême fragilité de ce système de régulation.

2.2.4. Régulation dans le cas du canton fixe

Les limites attachées aux différents algorithmes présentés amènent à utiliser en pratique des algorithmes plus simples ; tout particulièrement sur le VAL de Lille utilisant le principe du canton fixe.

Une des finalités de la régulation de trafic est le respect, par l'ensemble des rames en ligne, d'un tableau d'horaires théoriques de passage en station [BAI 2-94].

C'est à dire qu'à chaque rame est associée une rame fictive qui progresse de façon idéale le long de la voie en suivant un profil de vitesse nominal. A chaque station, la rame réelle compare son heure d'arrivée réelle avec l'heure d'arrivée théorique.

L'heure théorique de départ à la station est calculée par :

- Les tables d'horaires de départs des terminus,
- Les durées d'arrêt en station,
- Les temps nominaux de parcours d'interstation.

Pour rattraper un retard, on peut :

- Choisir la consigne de vitesse à transmettre au train en fonction du profil de l'interstation sur laquelle il est envisagé qu'il récupère son retard.
- Moduler le temps de stationnement.
- Provoquer un décalage horaire (DDH) si le retard dépasse une certaine limite.

La régulation de trafic du système VAL utilise aussi la fonction d'injection/retrait des rames à partir des terminus. Elle utilise des places de parkings proches des terminus pour stocker des rames qui pourront être utilisées dans cette fonction.

Cette fonction gère la commande de cycle en terminus, le cycle pouvant être :

- un rebroussement,
- un retrait de la rame après arrêt au terminus arrivée,
- une injection d'une rame au terminus départ.

Elle n'utilise que la table des horaires de départ du terminus et le suivi des trains à l'arrivée.

Ces entrées lui permettent de gérer sans distinction :

- le nombre de véhicules en ligne en fonction de l'intervalle d'exploitation,
- la régulation de trafic en terminus.

Son principe de fonctionnement fait appel :

- à l'heure théorique du prochain départ,
- à la position des deux rames les plus proches de l'arrivée au terminus (une liste des prochaines rames arrivant au terminus est gérée sur les deux dernières interstations).

Une des particularités des futures lignes de métro réside dans la possibilité de définir un service partiel, comme c'est le cas dans la future ligne 2 du VAL de Lille.

2.3. Problématique du service partiel

Le service partiel consiste en l'imbrication ou l'inclusion d'un carrousel dans un autre (figure 1.3). Ceci permet à la ligne de desservir des portions de ligne à forte densité de population où la demande est forte en même temps que des portions à faible densité de population.

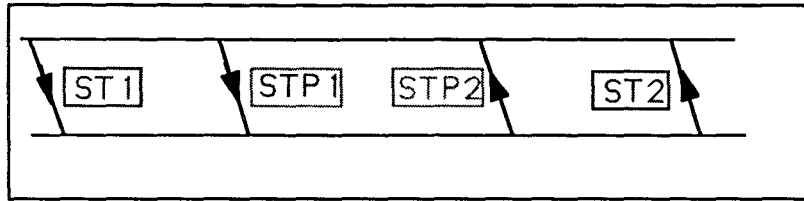


Figure 1.3. Schéma d'une ligne avec service partiel

ST1 station terminus 1

ST2 station terminus 2

STP1 station terminus partiel 1

STP2 station terminus partiel 2

Arrivées en terminus partiel, les rames dites "service partiel" rebroussement chemin, alors que les rames "service ligne" continuent jusqu'au prochain terminus.

Le principe du service partiel est particulièrement adapté pour les lignes où le trafic est concentré sur certaines zones de la ligne.

2.3.1. Services partiels envisageables

Le service partiel peut être géré différemment selon les hypothèses que l'on se fixe au départ. De ce fait, plusieurs types de services partiels peuvent être envisagés en fonction des besoins du trafic :

- Rebroussement d'une rame sur trois, conduisant à une augmentation de capacité de 50% sur la boucle service partiel par rapport au reste de la ligne. Toutefois, il est à noter que le rebroussement d'une rame sur trois établit sur l'extérieur de la boucle service partiel un carrousel à intervalle non régulier.

- Rebroussement d'une rame sur deux, permettant de doubler la capacité sur la boucle service partiel par rapport au reste de la ligne.

- Rebroussement de deux rames sur trois, permettant de tripler la capacité sur la boucle service partiel par rapport au reste de la ligne. Il faut néanmoins noter que cette solution triple l'intervalle à l'extérieur de la boucle service partiel ce qui, compte tenu de l'intervalle minimum sur la boucle service partiel (dû essentiellement au cantonnement de la ligne), risque de rendre la capacité de la ligne inférieure à la demande, surtout lorsqu'elle est forte en heure de pointe.

2.3.2. Service partiel ou lignes imbriquées

Après rebroussement d'une rame, deux cas sont envisageables :

- La rame reprend un service partiel

Ce cas se présente lorsque le temps de parcours de la boucle service partiel est un multiple pair de l'intervalle. Une rame service partiel ne prenant alors que des départs service partiel, on obtient bien la superposition d'un carrousel service partiel et d'un carrousel service ligne (figure 1.4).

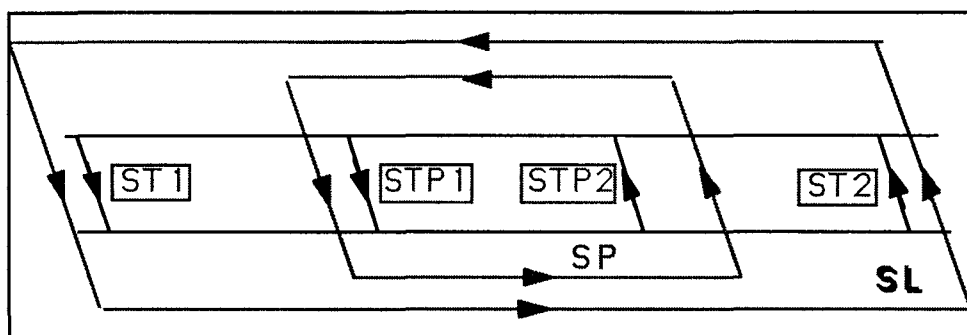


Figure 1.4. Boucle du service partiel dans la boucle du service ligne

- La rame prend un service ligne

Ce cas correspond à un temps de parcours de la boucle service partiel multiple impair de l'intervalle. La rame ayant rebroussé prenant un départ service ligne, la rame qui la suit au terminus partiel départ (rame venant de l'extérieur de la boucle service partiel), prendra un départ service partiel. On arrive à un système de deux lignes imbriquées avec un carrousel sur la portion ST1-> STP2 superposé à un carrousel sur la portion STP1 -> ST2 (figure 1.5).

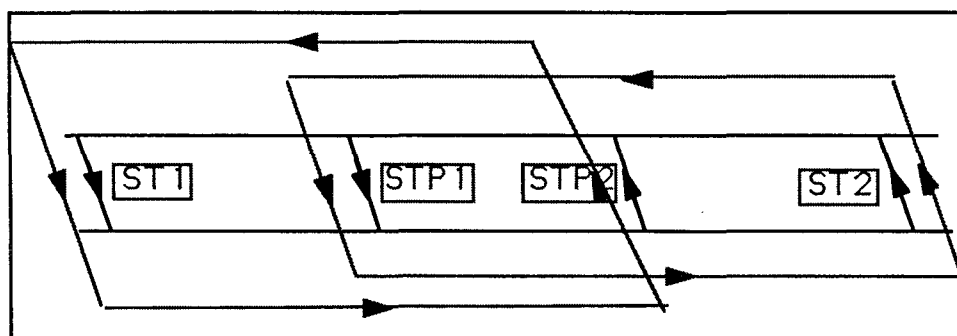


Figure 1.5. Boucles service partiel et service ligne imbriquées

Le deuxième cas offre une capacité de transport identique au premier sur le plan des intervalles. Il présente néanmoins l'inconvénient de ne pas relier toutes les stations entre elles. Ainsi un passager prenant une rame à ST1, sur la figure 1.5, devra changer de rame s'il veut se rendre à ST2. De ce fait, on ne retiendra pas ce deuxième cas comme solution.

2.3.3. Actions envisageables au niveau des terminus partiels

Le fonctionnement nominal du service partiel prévoit le rebroussement en terminus partiel des rames dites "service partiel" venant s'intégrer dans le carrousel service ligne de l'autre voie. Différentes actions particulières peuvent cependant être envisagées :

1. Rebroussement d'une rame "service ligne" devant normalement assurer un service ligne.

2. Passage en service ligne d'une rame "service partiel" devant normalement rebrousser.

3. Affectation de priorité

a) Priorité à une rame service ligne

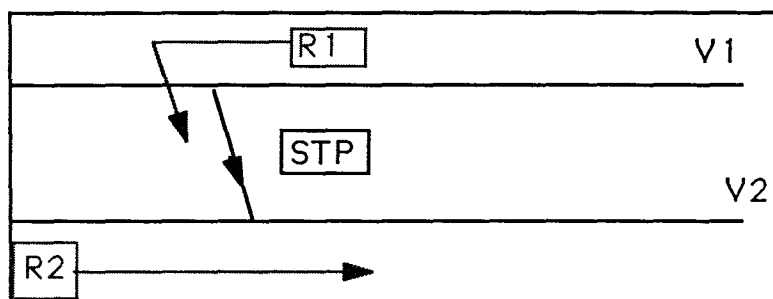


Figure 1.6. Priorité à une rame service ligne

b) Priorité à une rame service partiel

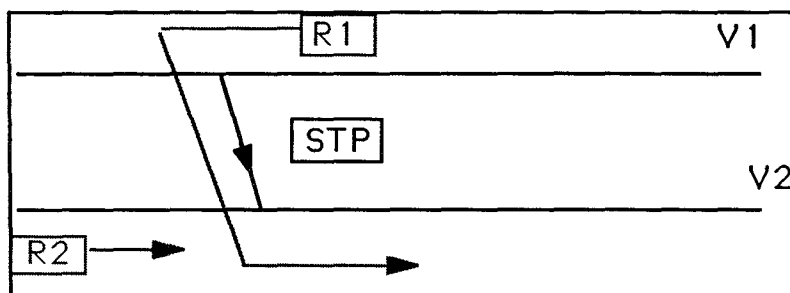


Figure 1.7. Priorité à une rame service partiel

Les solutions doivent prendre en compte la qualité de service offerte aux passagers. De ce fait, il faut éviter de faire descendre un passager de la rame pour lui faire prendre la suivante s'il veut arriver à destination. Ainsi, un passager prenant une rame à l'extérieur du service partiel ne changera pas de rame quelle que soit sa destination. De même, si un usager prend une rame dite service ligne à l'intérieur de la boucle service partiel, le rebroussement de la rame au terminus partiel sera interdit.

Cette seule exigence entraîne des conséquences sur les actions envisagées plus haut. Ainsi le rebroussement d'une rame "service ligne" devant normalement assurer un service ligne sera interdit. Le passage en service ligne d'une rame "service partiel" n'est pas gênant pour le passager, mais il sera interdit car il entraîne la propagation d'une lacune dans le carrousel service partiel.

Les affectations de priorité sont à proscrire parce qu'elles entraînent une inversion de l'ordre des rames service partiel/service ligne. Par conséquent, pour respecter les besoins de synchronisation à l'autre terminus partiel, il est nécessaire de faire passer en service ligne une rame service partiel ou de faire rebrousser une rame service ligne en la faisant passer en service partiel.

2.3.4. Conclusion

Au niveau d'un terminus partiel, les rames "service partiel" rebroussement chemin pour s'insérer dans le carrousel service ligne de l'autre voie. Cette condition impose une coordination obtenue par le biais d'une synchronisation des tables horaires de départ en terminus.

La satisfaction du client vis à vis du mode de transport doit justifier le choix de la régulation adoptée, quelle que soit la régulation employée. Cette satisfaction est définie par la qualité de service de ce système de transport.

3. Qualité de service d'un système de transport automatisé

La qualité de service est une notion présente à tous les niveaux de la vie industrielle. Elle l'est particulièrement dans le domaine tertiaire, en particulier les transports en commun. La définition par l'AFNOR (Association Française de NORmalisation) de la qualité de service adoptée en France de façon officielle adopte cette approche client de la qualité : "La qualité d'un produit ou d'un service est son aptitude à satisfaire les besoins d'un client". Pour certains chercheurs en marketing, le problème de la qualité de service, qui se définit comme une

différence entre ce que le client attend et ce qu'il reçoit, provient en grande partie de la différence entre ce que le client attend et ce que le prestataire suppose de l'attente du client.

Pour un exploitant de transport en commun, cette définition n'est mesurable qu'en soumettant les clients à un questionnaire complet sur leur satisfaction. Or ceci n'est pas facilement mesurable par des chiffres journaliers mais doit faire partie d'une grande étude de marketing sur l'impact du service sur les usagers. De ce fait, les exploitants ont essayé de définir l'attente du client vis à vis d'un système de transport en commun et ont mis en évidence les paramètres à ne pas négliger. Un système de transport en automatisation intégrale présente, face aux besoins des clients, un avantage qui est l'obtention d'une qualité de service minimisant les coûts d'exploitation et de maintenance. Ces deux contraintes ne sont généralement pas compatibles dans un système traditionnel mais peuvent le devenir grâce à l'automatisation intégrale. Les principaux avantages d'un système en automatisation intégrale pour les usagers sont :

- Un faible temps d'attente en station.
- Les intervalles attractifs en heures creuses.
- L'adaptation facile de l'offre et de la demande.
- La vitesse commerciale élevée.
- Le confort intérieur des véhicules.

3.1. Critères de qualité

3.1.1. Un faible temps d'attente en station

Le temps d'attente correspond à l'intervalle entre les trains sur la ligne. Cet intervalle est défini par la relation suivante :

$$\text{Intervalle(en minutes)} = 60 * \frac{\text{Capacité du train}}{\text{Débit horaire sur l' interstation la plus chargée}}$$

Pour écouler le trafic, on peut choisir de faire circuler N trains comprenant V voitures (avec un intervalle I) ou de faire circuler $k \times N$ trains comprenant $\left(\frac{V}{k}\right)$ voitures (soit à un

intervalle $\frac{I}{k}$). Le nombre de voitures au km est le même mais l'intervalle entre les trains est bien plus attractif dans la deuxième solution.

La qualité de service, pour un même débit, est donc améliorée en faisant circuler des rames de petites dimensions à intervalles faibles. Cependant, dans un système traditionnel avec conducteur (ou surveillant), cela entraîne une augmentation du nombre de conducteurs et par conséquent un accroissement important du coût d'exploitation.

Le fonctionnement avec des rames plus courtes permet également de diminuer la longueur des quais des stations, donc de gagner sur le coût des infrastructures. Il faut remarquer qu'un intervalle court (de l'ordre de la minute) ne peut de toute façon être tenu qu'avec un système automatique qui réagit plus rapidement que ne peut le faire un conducteur et surtout ne risque pas d'entraîner le "stress" que peut ressentir un conducteur devant assurer la conduite de trains rapprochés.

3.1.2. Intervalle attractif en heures creuses

L'absence de conducteur permet de maintenir un intervalle en heures creuses suffisamment faible (à peu près 5 minutes) pour rendre le système encore attractif durant cette période. Les systèmes traditionnels offrent en général des intervalles de l'ordre de 10 à 15 minutes durant ces périodes.

3.1.3. Adaptation facile de l'offre et de la demande

Les systèmes traditionnels ont des programmes d'exploitation très rigides basés sur des notions de "services", c'est-à-dire qu'à chaque horaire de départ des trains des garages ou des terminus est associé un nom de conducteur. Les couples trains/conducteurs ainsi définis pour la journée d'exploitation, ont pour mission d'effectuer un certain nombre de "tours" de ligne en cohérence avec les horaires et la rotation du personnel d'exploitation (tours de service). Cette gestion est généralement très lourde pour l'exploitant et difficilement adaptable à des circonstances particulières.

Pour un système en automatisme intégral, toutes les rames et les départs sont banalisés. La contrainte est de respecter pour chaque tranche horaire de la journée, l'intervalle permettant de satisfaire le débit. Les numéros de rames assurant les départs dans chacune des tranches horaires peuvent être quelconques et ne font pas l'objet d'une pré-affectation (aucun lien entre le départ et le numéro de rame). De même, l'ajout de trains supplémentaires pour des besoins particuliers liés à des événements occasionnels, est assuré sans contrainte de personnel.

3.1.4. Vitesse commerciale élevée

Une vitesse commerciale élevée a deux avantages :

- elle permet de diminuer le temps de trajet des usagers
- elle permet de diminuer le nombre de rames en ligne pour un intervalle donné d'où une diminution du coût d'investissement

$$\text{Nombre de trains en ligne} = \frac{\text{Temps de parcours}}{\text{Intervalle}}$$

Une vitesse commerciale élevée nécessite d'évoluer sur une voie en site propre. Les systèmes de transport de type bus ou tramway, même s'ils bénéficient d'aménagements privilégiés tels que couloirs réservés sont toujours tributaires des carrefours. Ceci est d'autant plus vrai aux heures de pointe (circulation automobile importante) lorsque ce système de transport devrait être le plus compétitif.

La qualité de service d'un système de transport doit être quantifiée par un certain nombre de paramètres. Ces paramètres peuvent être considérés comme des indicateurs définissant la qualité de service.

3.2. Indicateurs de la qualité de service

On peut, lors de la mise en oeuvre de la régulation, souhaiter :

- minimiser le retard d'un nombre maximum de passagers. Ce qui veut dire que l'intervalle entre les rames et la distribution des passagers le long des véhicules peuvent être non-uniformes. Lors d'une augmentation importante de passagers à un endroit de la ligne, si l'intervalle entre les rames est plus petit qu'ailleurs, le transport des passagers attendant sur le quai pourra se faire plus rapidement.

- minimiser l'écart du retard des trains par rapport à la table horaire.

Un compromis entre les deux objectifs doit être défini en privilégiant le respect de la table horaire pour maintenir un service cohérent le long de la ligne.

Taskin [TAS 85] propose de mesurer la qualité de service par des fonctions de coût. La régulation consiste alors en un problème d'optimisation. Cet algorithme est décrit à l'annexe 1. Ces indices de qualité de service sont très théoriques car ils prennent en compte la

connaissance totale du nombre de voyageurs non seulement sur les quais mais aussi dans les rames à l'instant de la perturbation. Or certains paramètres sont très difficiles à mesurer, comme la valeur exacte du retard subi par les voyageurs ou le temps d'attente au quai, qui est une notion ressentie de façon très subjective. "Dans un cas favorable, il est possible qu'un temps d'attente de 10 minutes soit ressenti comme un temps de 5 minutes et inversement dans un cas défavorable" [MOR 92]. La principale difficulté est d'évaluer cette notion.

Pour cela, les exploitants ont mis en place des indicateurs plus simples qui formalisent la qualité de service dans leur réseau. Les indicateurs utilisés par l'exploitant dépendent de sa ligne et de son mode de régulation. Nous allons détailler les indices de qualité de service sur le VAL de Lille et montrer qu'ils sont totalement dépendants de la régulation.

Dans le système VAL, certains indicateurs sont relevés chaque mois pour établir la notion de qualité de service. Ces indicateurs sont calculés chaque jour et constituent les statistiques d'exploitation. Au début de chaque journée, une série de tables horaires est définie en fonction de la date d'exploitation : jour de la semaine normal, période de vacances scolaires, samedi, dimanche, jour férié, manifestation ponctuelle, ... D'après la table horaire sélectionnée, les rames doivent assurer les "services" définis, à savoir les départs du terminus à l'heure prédéfinie dans cette table. Les heures de départ définissent l'intervalle entre les rames selon la période d'exploitation (heures de pointe ou heures creuses). Chaque rame sera affectée d'un numéro de service dépendant de son heure de départ au terminus départ. Ce numéro de service lui sera affecté jusqu'à son arrivée au terminus arrivée et un nouveau numéro de service lui sera affecté si la rame rebrousse chemin pour prendre un nouveau départ. De ce fait, les stations le long de la ligne attendent à une heure précise, non pas un numéro de rame, mais un numéro de service.

Ces indicateurs utilisent les données disponibles et sont utilisés en vue de la régulation du trafic pour donner une image de la qualité de service. Ces indicateurs décrits plus précisément à l'annexe 1 dépendent essentiellement des décalages horaires effectués par la régulation de trafic.

Bien évidemment, tous ces indices sont choisis en fonction de la régulation de trafic adoptée.

Dans la régulation de type VAL, le flux de voyageurs est intégré dans la définition des tables horaires. Ces tables horaires sont établies en fonction du jour, de l'heure et de la ligne concernée, à l'aide des estimations du nombre de personnes fréquentant le métro. Mais, de cette façon, le système ne prend pas en compte le nombre de voyageurs réel à l'instant présent. Si, pour une raison ou pour une autre, le nombre de voyageurs à certaines stations est plus

important qu'à l'habitude, la fréquence de passages des rames en station n'augmentera pas pour autant. Or, cet afflux supplémentaire peut provoquer des perturbations sur l'ensemble du réseau. En effet, les temps d'arrêts des rames aux stations, nécessaires pour faire entrer dans la rame un nombre de passagers supérieur à la normale, augmentent provoquant des retards qui s'accroissent graduellement sur ces trains. Du fait de ces retards, le nombre de passagers, attendant aux quais suivants, augmente et provoquera pour les mêmes raisons un retard supplémentaire sur les trains. Dans ce cas, l'attente du voyageur prenant habituellement le métro à cette heure, sera fortement augmentée et la qualité de service du réseau sera dégradée.

Comme nous l'avons vu précédemment dans la définition de la qualité de service, les desiderata d'un client vis à vis d'un système de transport en commun sont l'attente la moins longue possible en station et l'arrivée la plus rapide possible à destination. Pour cela, le critère le plus représentatif de la qualité de service est la durée moyenne d'un trajet d'un voyageur quelle que soit sa destination et quelle que soit l'heure à laquelle il emprunte le réseau. Or dans les indices de qualité de service du VAL de Lille, ce paramètre est défini comme une entrée du système et non pas comme une sortie puisqu'il est demandé par le système à l'opérateur. Il serait intéressant de pouvoir estimer ce temps de parcours moyen.

Malheureusement un certain nombre de paramètres utilisés dans la régulation de trafic ne sont pas précisément définis. Ils sont entachés d'une incertitude qui devra être prise en compte dans l'algorithme de régulation.

3.3. Paramètres introduisant incertitude et imprécision

Le système étant complexe, il n'est pas modélisé dans son intégralité. Aucune modélisation globale précise n'a été faite sur l'ensemble du système car certaines causes produisent des conséquences difficiles à prévoir. De ce fait, le système a été décomposé en sous-tâches dépendantes les unes des autres. De plus, un certain nombre de paramètres agissant sur le système ne sont que partiellement connus.

En premier lieu, nous citerons le retard des rames ou, plus précisément, les causes les produisant. Les pannes se déclarent d'une façon imprévisible et, si l'on peut modéliser la conséquence de chaque panne sur l'ensemble du réseau, on ne peut prévoir qu'une panne va se déclencher à l'instant T. Ceci est renforcé par la cause principale des pannes sur les lignes de métro : la malveillance des voyageurs qui échappe à toute loi de maintenance préventive.

Dans les systèmes en automatisme intégral, la position précise de la rame est méconnue entre deux stations. Au départ d'une station, une consigne de vitesse est donnée à la rame et

celle-ci suit un profil de vitesse pour aller à la prochaine station, indépendamment de l'environnement extérieur. Seules les consignes sécuritaires sont connues en dehors du sous-système rame. Par exemple, le franchissement par la rame du seuil d'un canton déjà occupé provoque l'arrêt immédiat de la rame. Au niveau du poste de commande, seul l'état du canton sur lequel se trouve la rame est connu. On peut donc dire que la rame se trouve entre deux positions : la position du début du canton et la position de fin.

Ignorant où se trouve la rame entre deux stations, nous ne connaissons que d'une manière imprécise la vitesse instantanée de la rame. Cependant, nous possédons la consigne envoyée à la rame à son départ de la station. Néanmoins, à chaque instant, nous ne savons ni sa vitesse, ni son accélération, ni son jerk.

Le calcul du paramètre "durée moyenne de trajet d'un voyageur" représentatif de la qualité de service, est rattaché à la notion du flux de passagers. Il faut connaître à chaque instant la répartition des voyageurs dans les rames et sur les quais des stations. Or, le nombre de passagers est mal connu en temps réel par les exploitants car, pour cela, il serait nécessaire de posséder un système de comptage de personnes. De plus, le nombre de passagers à l'intérieur d'une rame est lui aussi difficile à calculer. De nombreux travaux sont en cours pour dénombrer les passagers à l'aide de traitement d'image mais aucun de ces systèmes n'est en service actuellement sur un réseau de transport urbain. Nous pouvons donc considérer le nombre de voyageurs sur l'ensemble de la ligne comme un paramètre incertain et imprécis. Les exploitants ont une estimation de ce nombre grâce à des études statistiques journalières sur la fréquentation de leur ligne mais ne connaissent pas avec précision ce nombre à un instant donné. Néanmoins, l'avènement de la billetterie magnétique peut régler ce problème. Cette méthode a fait ces preuves dans le réseau de la ville de Rouen.

4. Conclusion

La régulation d'un réseau de transport en commun à haute densité, style le métro automatique, peut être réalisée par des approches rigoureusement établies issues des concepts généraux du contrôle et de la commande en automatique. Néanmoins, ces approches se basent sur des modèles qui ne reflètent pas les phénomènes réels. En effet, les équations sont fondées sur des hypothèses de départ souvent simplificatrices et ne représentant pas la situation réelle. La volonté de considérer des hypothèses plus réalistes nous donne des modèles d'une grande complexité pour lesquels il est difficile de déterminer des solutions optimales. En pratique des algorithmes plus simples, type régulation du VAL, sont utilisés, ignorant volontairement un certain nombre de paramètres, pour limiter la complexité de la régulation. De ce fait, ces

systèmes peuvent être satisfaisants si on considère les indicateurs de qualité actuels. Or nous l'avons vu, ces indicateurs de la qualité de service sont définis en fonction de la régulation de trafic employée et peuvent être beaucoup plus informatifs.

Les différents paramètres, non pris en compte par la régulation, nous amènent à nous intéresser à l'incertitude de leur évaluation. Celle-ci est due à un certain nombre de phénomènes physiques souvent ignorés car trop complexes à modéliser. Ceci nous amène à définir un nouveau modèle de régulation basé sur une approche prenant en compte l'imperfection des données telle que celle que peut fournir la théorie des sous ensembles flous. Mais, avant de définir plus précisément notre modèle, nous pouvons rappeler l'apport de la logique floue dans la commande de systèmes complexes.

CHAPITRE 2

APPORT DE LA LOGIQUE FLOUE DANS LA COMMANDE DE SYSTEMES COMPLEXES

1 . Introduction

La logique floue suscite actuellement un intérêt général de la part des chercheurs, des ingénieurs et des industriels, et plus généralement de la part de tous ceux qui éprouvent le besoin de formaliser des méthodes empiriques, de généraliser des modes de raisonnements naturels, d'automatiser la prise de décision, de construire des systèmes effectuant les tâches habituellement prises en charge par les hommes.

Les connaissances dont nous disposons sur une situation quelconque sont généralement imparfaites. Soit parce que nous avons un doute sur leur validité, elles sont alors incertaines. Soit parce que nous éprouvons une difficulté à les exprimer clairement, elles sont alors imprécises. Dans le fonctionnement de l'esprit humain, les imprécisions sont particulièrement courantes, par exemple dans les fonctions de reconnaissance et de raisonnement. La capacité d'établir des classes d'éléments de la nature ayant des propriétés analogues est très naturelle chez l'homme. Il lui est tout aussi naturel de traiter des données affectées d'incertitudes que d'utiliser des critères subjectifs, donc imprécis. Notre capacité à décrire précisément un système est en fonction inverse de sa complexité. Le besoin d'étudier ou de gérer des systèmes complexes conduit à la prise en compte de données vagues, imprécises, soumises à des erreurs, mal définies, dont la validité n'est pas absolue, soumises à une incertitude. Pourtant l'homme peut gérer ces données mal connues. Par exemple, la conduite d'une automobile oblige la prise en compte d'un certain nombre d'informations imparfaites (un enfant jouant sur le trottoir).

Les deux types d'imperfection dans les connaissances n'ont cependant pas eu la même importance dans les préoccupations des scientifiques. L'incertain a d'abord été abordé par la notion de probabilité par Pascal et Fermat. Cependant cette théorie ne permet pas de traiter des informations subjectives ni d'aborder les connaissances imprécises ou vagues [BOU 93].

Ces dernières n'ont été prises en considération que lorsque L. A. Zadeh [ZAD 65] a introduit la notion de sous ensembles flous, à partir de l'idée d'appartenance partielle à une classe dans une généralisation de la théorie classique des ensembles. Les développements de cette notion fournissent des moyens de représenter et de manipuler des connaissances imparfaitement décrites, vagues ou imprécises et ils établissent une interface entre des données écrites symboliquement et numériquement.

Dans la suite logique de la théorie des sous ensembles flous, la théorie des possibilités qui a été introduite en 1977, également par L. Zadeh [ZAD 77], constitue un cadre permettant de traiter des concepts d'incertitude de nature non probabiliste. Lorsqu'elle est considérée à partir de la notion d'ensembles flous, la théorie des possibilités constitue un cadre permettant d'exploiter, dans un même formalisme, imprécisions et incertitudes. La théorie des sous ensembles flous et la théorie des possibilités font partie de la famille d'approches de la théorie générale de l'information dont fait également partie la théorie de l'évidence.

2. La logique floue et ses applications à la commande floue et aux systèmes à base de connaissances

2.1. Principes de base de la logique floue

La logique floue peut être envisagée comme une extension de la logique booléenne. Ainsi, alors qu'en logique classique, l'appartenance d'un élément à un sous ensemble donné est une valeur binaire (0 ou 1), la théorie de la logique floue considère cette appartenance comme une valeur réelle dans l'intervalle $[0;1]$. Comme dans le cas de la logique classique, une valeur nulle implique la non appartenance de l'élément au sous ensemble, tandis que la valeur 1 indique que l'élément appartient complètement au sous ensemble considéré. De ce fait, la logique floue permet de manipuler des symboles et d'inférer des actions en utilisant des règles logiques à partir de prémisses imprécises ou incertaines. L'idée de Zadeh est de graduer l'appartenance. C'est ce qui distingue la logique floue de la logique ensembliste qui se contente de classer un objet dans une classe par le seul fait qu'il appartient ou n'appartient pas à la classe.

Par exemple, nous passons graduellement de la notion d'assez grand à la notion de grand. Par exemple un individu de 1,70 m peut être considéré comme assez grand ; plus la taille s'approchera de 1,80 m moins l'individu appartiendra à la classe assez grand pour s'approcher de la classe grand. Ceci se traduira par une diminution de son coefficient d'appartenance à la

classe assez grand et par une augmentation du coefficient d'appartenance à la classe grand. De ce fait un individu de 1,70 m aura un coefficient d'appartenance de 1 à la classe assez grand et de 0 à la classe grand. Celui de 1,75 m aura un coefficient de 0,86 à la classe assez grand et de 0,86 à la classe grand. De la même manière un individu de 1,80 m aura un coefficient d'appartenance de 0 à la classe assez grand et de 1 à la classe grand.

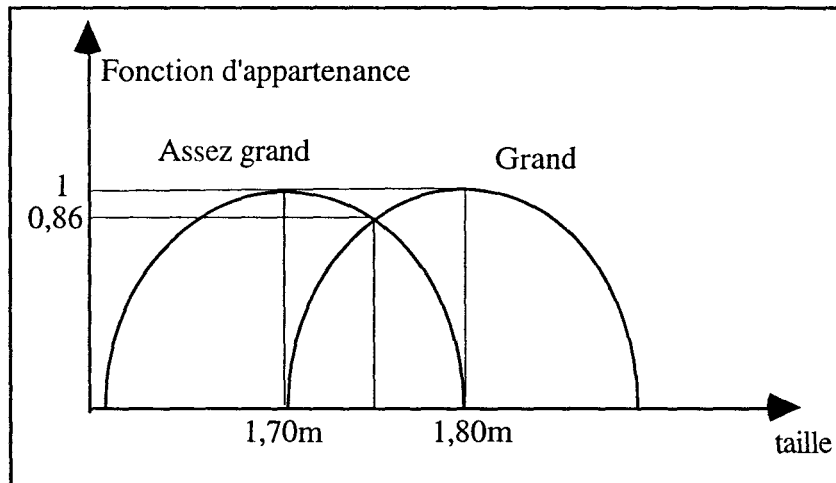


Figure 2.1. Fonction d'appartenance à la classe grande taille

Cette capacité à traiter des données avec des valeurs intermédiaires permet aux systèmes à base de logique floue de s'adapter avec plus de facilités aux environnements réels que ne le font les systèmes basés sur la logique binaire.

Un sous ensemble flou A de X est défini par une fonction d'appartenance qui associe à chaque élément x de X, le degré $\mu_A(x)$, compris entre 0 et 1, avec lequel x appartient à A :

$$\mu_A: X \rightarrow [0,1]$$

Le sous ensemble flou A devient un sous ensemble classique de X dans le cas particulier où $\mu_A(x)$ ne prend que des valeurs égales à 0 ou à 1. Un sous ensemble classique est donc un cas particulier de sous ensemble flou.

L'ensemble de tous les éléments appartenant de façon absolue (avec un degré 1) au sous ensemble flou A est appelé le noyau de A et noté $\text{noy}(A)$:

$$\text{noy}(A) = \{ \forall x \in U, \mu_U(x) = 1 \}$$

Les univers de discours représentent les domaines d'évolution des variables. Par exemple, si l'on considère une vanne dont l'ouverture est limitée à 80%, l'univers de discours

correspondant peut être défini par l'intervalle $[0;0,8]$ et on peut envisager, sur cet univers de discours, le sous ensemble flou "ouverture moyenne" caractérisé par la fonction d'appartenance :

$$\begin{aligned}\mu_{\text{Ouverture moyenne}} : [0;0,8] &\rightarrow [0;1] \\ x &\rightarrow \mu_{\text{Ouverture moyenne}}(x)\end{aligned}$$

Les principaux opérateurs sont détaillés en annexe 2. Ils permettent de définir des relations et opérations d'utilités diverses sur les ensembles flous.

2.2. La commande floue

La commande floue est le domaine dans lequel il existe le plus de réalisations connues, en particulier industrielles. Son but est de traiter des problèmes de commandes de processus, le plus souvent à partir des connaissances des experts ou d'opérateurs qualifiés travaillant sur le processus.

2.2.1. Typologie des contrôleurs flous

Lorsque l'on ne dispose pas d'un modèle défini au sens classique du terme, la commande floue peut être considérée comme une démarche adaptée. Dans ce cas, on synthétise la loi de commande en codant les connaissances des experts qui contrôlent le processus (opérateurs, pilotes...). L'utilisation d'un codage flou permet de représenter leur caractère graduel et nuancé. Ainsi, cette démarche revient à modéliser le comportement de l'opérateur de commande, et non le processus à commander.

Un contrôleur flou peut être vu comme un système expert simple fonctionnant à partir d'une représentation des connaissances basée sur les ensembles flous [FOU 94]. La base de connaissances contient les définitions des termes utilisés dans la commande et les règles caractérisant l'effet de la commande et décrivant la conduite de l'expert.

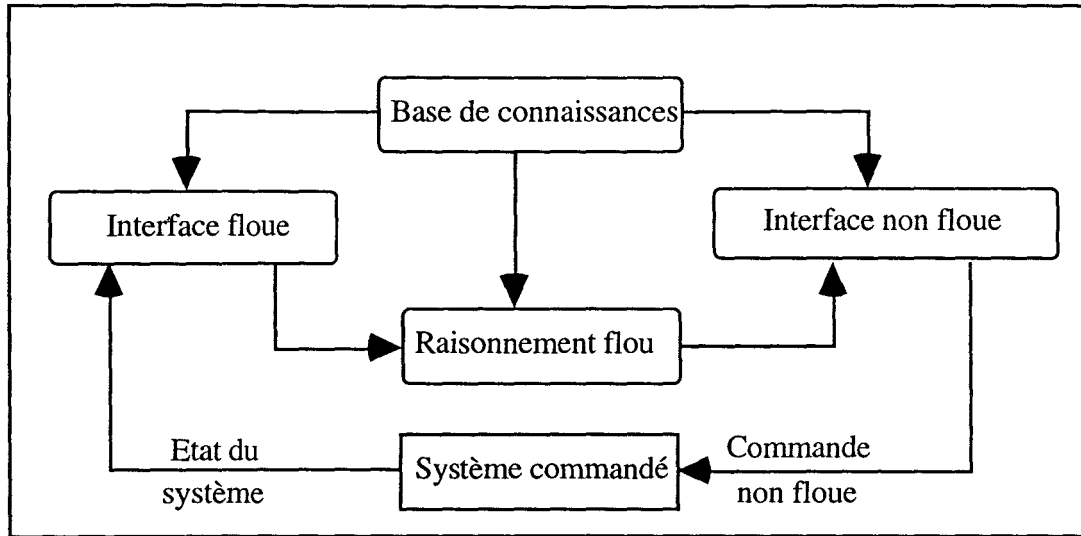


Figure 2.2. Configuration générale d'un contrôleur flou.

Un contrôleur flou est un système à base de connaissances utilisant un raisonnement dans une procédure de chaînage avant des règles. Les contrôleurs flous utilisent une expertise exprimée sous forme de règles, dont la forme est la suivante pour un contrôleur à deux entrées et une sortie :

Si $\{(X1 \text{ est } A1) \text{ et } (X2 \text{ est } A2)\}$ **alors** $(Y \text{ est } B)$

L'expression $\{(X1 \text{ est } A1) \text{ et } (X2 \text{ est } A2)\}$ est la prémisse de la règle, tandis que l'expression $(Y \text{ est } B)$ est la conclusion de la règle. Les variables $X1$, $X2$, Y représentent les variables physiques caractéristiques du processus à commander. $A1$, $A2$, B sont des valeurs linguistiques prises par les variables $X1$, $X2$ et Y . Les différentes valeurs linguistiques que pourront prendre les variables $X1$, $X2$ et Y sont représentées par des ensembles flous caractérisés, sur leurs univers de discours respectifs, par des fonctions d'appartenance.

Dans les contrôleurs flous les plus répandus, il n'y a pas d'enchaînement de règles ; toutes les règles activables sont activées. Une règle est activable dès que sa prémisse est caractérisée par un degré d'appartenance non nul.

On peut distinguer plusieurs étapes dans le traitement des règles : la fuzzification, l'inférence floue qui repose sur l'utilisation d'un opérateur d'implication, l'agrégation des règles et la défuzzification.

En automatique classique, la commande est calculée à partir d'un modèle mathématique du processus. La commande floue permet de décrire le comportement du processus sous forme de règles dont le schéma bloc peut être représenté par la figure suivante :

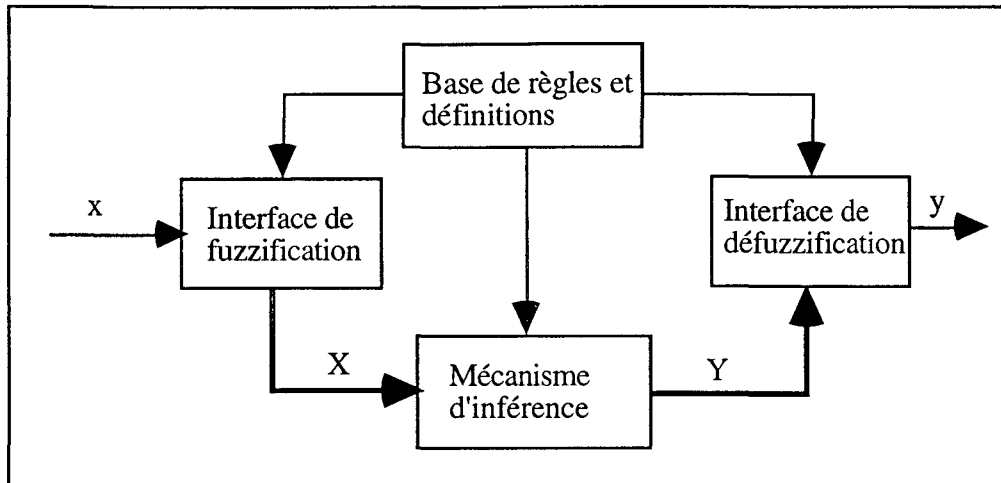


Figure 2.3. Schéma bloc d'un contrôleur flou

2.2.2. Bases de règles et définitions

On regroupe dans ce bloc l'ensemble des définitions utilisées dans la commande floue (univers de discours, partition floue, choix des opérateurs, ...), ainsi que la base de règles, transcription sous forme de règles floues de la stratégie de commande fournie par l'expert.

2.2.3. Interface de fuzzification

Pour classer chaque donnée, il faut convertir les entrées en données floues. La "fuzzification" est une conversion floue qui permet le passage du numérique au symbolique.

La fuzzification consiste à établir la relation entre l'ensemble des valeurs que peut prendre la variable et l'ensemble des grandeurs linguistiques qui la définissent. Cette évaluation est sujette à l'appréciation et au jugement de l'expert humain pour le problème posé. La forme et les valeurs numériques représentatives des fonctions caractéristiques relatives à un qualificatif sont dépendantes de la perception de l'individu.

2.2.4. Inférence floue

A partir de la base de règles et du sous ensemble flou X_0 correspondant à la fuzzification de la variable mesurée x_0 , le mécanisme d'inférence calcule le sous ensemble flou Y relatif à la commande à appliquer au système.

On rappelle qu'en logique classique le modus ponens permet, à partir de la règle "Si X est A alors Y est B " et du fait " X est A " de conclure le fait " Y est B ", qui sera ajouté à la base

des faits, dans un raisonnement par chaînage avant. Zadeh a étendu ce principe au cas flou, principe que l'on appelle alors *modus ponens généralisé*.

Fait	$x \text{ est } X_0$	$x \text{ est } X_0'$
Règle	Si $x \text{ est } X_0$ alors $y \text{ est } Y$	Si $x \text{ est } X_0'$ alors $y \text{ est } Y'$
Déduction	$y \text{ est } Y$	$y \text{ est } Y'$

A partir de la règle "Si $x \text{ est } X_0$ alors $y \text{ est } Y$ " et du fait X_0' , on en déduit un nouveau fait Y' qui est caractérisé par un ensemble flou dont la fonction d'appartenance est :

$$\mu_{Y'}(y) = \sup_{x \in X} ((\mu_{X_0'}(x) \wedge \mu_R(x, y)))$$

Les fonctions d'appartenance $\mu_{X_0'}$ et μ_R caractérisent respectivement le fait X_0' et la règle. Si X_0' correspond à un singleton $\{x_0\}$, $(\mu_{X_0'}(x) = 1 \text{ si } x = x_0, 0 \text{ ailleurs})$. L'opérateur $\wedge$ est une T-norme (voir annexe 2).

2.2.5. Défuzzification

Le sous ensemble Y de l'univers de discours de la commande ayant été calculé par le mécanisme d'inférence, l'interface de défuzzification a pour objectif de le transformer en une valeur réelle y_0 permettant ainsi la commande effective du système. Il existe plusieurs algorithmes de défuzzification comme par exemple la méthode du centre de gravité, la méthode du maximum, la méthode des hauteurs.

Un grand nombre d'applications a vu le jour utilisant la commande floue. On peut citer par exemple la commande de machines outils, de groupes d'ascenseurs, d'appareils électroménagers, de caméras, ... [SAN 89] et [TON 85]. De grands projets ont été réalisés comme la commande vocale d'un hélicoptère ou la réalisation d'une commande floue pour le métro de Sendai au Japon [OSH 88] et [YAS 85]. Le contrôle flou des rames de métro consiste en une commande de moteurs. Les critères retenus pour la conduite de la rame sont l'amélioration de la sécurité, l'économie d'énergie, le confort, le suivi de vitesse, la précision de l'arrêt et le temps de circulation. Les sorties du contrôleur sont une commande de traction ou de freinage à appliquer sur la rame. Les résultats de ce système de commande ont été comparés à un système de commande réalisé avec un PID classique donnant : un gain d'énergie de 10%, un temps de parcours raccourci entre deux stations (4%), une meilleure précision de l'arrêt et une amélioration du confort des passagers par un changement de commande sur l'accélération et la décélération intervenant deux fois moins souvent. Le problème posé est cependant un peu différent du nôtre, car c'est un contrôle flou réalisé sur la commande des moteurs en vue

d'améliorer la marche des trains entre deux stations. Ce contrôle permet d'optimiser la marche des trains au niveau local. Il n'est aucunement question de réguler l'ensemble des rames circulant sur la ligne pour optimiser son fonctionnement, ni de permettre d'agir en fonctionnement dégradé.

2.3. Représentation de l'incertain

Dans certaines applications il peut être utile d'appliquer les connaissances disponibles sur le système à commander. Dans ce cas, on peut recourir à une base de connaissances sur les événements passés pour prédire les événements futurs ou plus simplement pour commander le système en utilisant l'expérience du passé. Un certain nombre de méthodes nous permettent d'utiliser le passé du système, entre autres les statistiques, la théorie des probabilités et la théorie des possibilités.

2.3.1. Représentation statistique

Dans le langage courant, le mot "statistique" désigne des collections de nombres présentées parfois sous forme de tableaux ou de graphiques, qui regroupent toutes les observations faites sur des faits nombreux et relatifs à un même phénomène. Le statisticien devra se pencher sur les nombres qui lui sont soumis de manière à obtenir des rapports numériques sensiblement indépendants des anomalies du hasard et qui dénotent l'existence de causes régulières dont l'action s'est combinée avec celle des causes fortuites.

L'observation d'un phénomène nous permet de définir des classes dans lequel il se déroule. A chaque classe on attribue un effectif, à savoir le nombre d'événements qui se sont déroulés dans cette classe. Nous pouvons représenter ces données sous différentes formes graphiques comme les histogrammes, les graphiques circulaires... qui permettent de mettre en évidence la notion de fréquence d'apparition des événements. D'après ces données, nous pouvons établir certaines relations mathématiques permettant de représenter les événements. Il est également possible de définir des caractéristiques de dispersion qui ont pour but d'apprécier la dispersion des valeurs observées autour de la valeur centrale qui est en général la moyenne arithmétique. Cette représentation des données permet de matérialiser les données observées lors du fonctionnement d'un système et leur répartition. Par contre elle ne permet pas de prédire un événement futur d'après les informations disponibles sur le système.

2.3.2. Représentation probabiliste

2.3.2.1. Probabilité élémentaire

La théorie des probabilités fait appel au rapport du nombre d'occurrences d'un événement et du nombre total d'événements. Elle est basée sur le fait qu'on possède un nombre relativement important d'informations sur le passé. De ce fait, lorsque nous pouvons observer un nombre important d'événements, nous pouvons définir la probabilité de A (élément de E). A partir d'événements qui se sont déroulés durant le fonctionnement du système, nous établissons, pour chaque événement, sa probabilité d'apparition. Nous déduisons de ce fait que l'événement A a $x\%$ de chance de se produire. Pour chaque événement, nous calculons sa densité de probabilité ; c'est à dire la répartition de la conséquence de l'événement dans le temps. Par exemple, lors du fonctionnement d'une ligne de métro, certaines pannes se produisent. La probabilité d'avoir une panne notée A est égale au nombre de pannes de type A constatées rapporté au nombre de pannes totales. Pour les pannes de type A, la répartition de leur durée nous donne une densité de probabilité. On pourra savoir que la panne de type A a 10 % de chance de se produire et que sa durée moyenne est de 5 minutes.

2.3.3. Représentation possibiliste

La théorie des possibilités permet de raisonner sur des connaissances imprécises ou vagues, en introduisant un moyen de prendre en compte des incertitudes sur ces connaissances. Si des informations de nature fréquentielle étaient disponibles, indiquant dans quelle proportion des cas une règle est valide ou avec quelle probabilité un fait est vrai, on pourrait effectuer un raisonnement de nature probabiliste. On suppose ici que l'on ne dispose pas de telles informations et que l'on est seulement capable de dire dans quelle mesure il est possible et certain que les règles sont valides et les faits vrais.

2.3.3.1. Mesure de possibilités

Etant donné un ensemble fini X, on attribue à chaque sous ensemble de X, noté A_i , un coefficient compris entre 0 et 1 qui évalue à quel point cet événement est possible. On définit pour cela une mesure de possibilités Π définie sur chaque sous partie de X et telle que :

$$\begin{aligned}\Pi(\text{vide}) &= 0 \\ \Pi(X) &= 1 \\ \forall A_i \subset X, \Pi\left(\bigcup_{i=1}^{i=n} A_i\right) &= \sup_{i=1, \dots, n} (\Pi(A_i))\end{aligned}$$

2.3.3.2. Distribution de possibilités

Une distribution de possibilités est définie à partir d'une mesure de possibilités si l'on peut attribuer à chaque partie élémentaire x de X un degré de possibilité. Elle est notée π et satisfait les propriétés suivantes :

$$\begin{aligned}\pi(x) &\in [0,1] \\ \sup_{x \in X} (\pi(x)) &= 1 \\ \forall A \subset X, \Pi(A) &= \sup_{x \in A} \pi(x)\end{aligned}$$

Une distribution de possibilités est représentable par un ensemble flou.

2.3.3.4. Mesure de nécessité

Une mesure de possibilités fournit une information sur l'occurrence d'un événement A relatif à un ensemble de référence X , mais elle ne suffit pas pour décrire l'incertitude existant sur cet événement. Par exemple, si $\Pi(A) = 1$, il est tout à fait possible que A soit réalisé, mais on peut avoir en même temps $\Pi(A^c) = 1$, avec A^c l'événement contraire de A , ce qui exprime une indétermination complète sur la réalisation de A . Pour compléter l'information sur A , on indique le degré avec lequel la réalisation de A est certaine, par l'intermédiaire d'une mesure de nécessité, grandeur duale de la mesure de possibilité. Une mesure de nécessité N d'un événement a une valeur d'autant plus grande que l'événement contraire est impossible. N est définie sur chaque sous partie de X et telle que :

$$\begin{aligned}N(\text{vide}) &= 0 \\ N(X) &= 1 \\ \forall A_i \subset X, N\left(\bigcap_{i=1}^{i=n} A_i\right) &= \inf_{i=1 \dots n} (N(A_i))\end{aligned}$$

Les propriétés imposées aux mesures de possibilités et de nécessité suggèrent une association. Plus précisément, sur un ensemble de référence X , une mesure de nécessité N peut être obtenue à partir de la donnée d'une mesure de possibilités Π , par l'intermédiaire du complémentaire A^c de toute partie A de X :

$$\forall A \in X, N(A) = 1 - \Pi(A^c)$$

Plus un événement A est affecté d'une grande nécessité, moins l'événement complémentaire A^c est possible, donc plus on est certain de la réalisation de A . $\Pi(A)$ mesure le

degré avec lequel l'événement A est susceptible de se réaliser, $N(A)$ indique le degré de certitude que l'on peut attribuer à cette réalisation.

3. Décomposition hiérarchique

3.1. Décomposition des systèmes complexes

Les systèmes de transport sont complexes et sujets à des perturbations aléatoires. Dans beaucoup de cas, il est difficile de savoir si la complexité apparente du système est due à une connaissance insuffisante sur le comportement du système ou aux sources de bruits extérieurs. Chaque composant du système peut être parfaitement connu et modélisable mais, pris dans son ensemble, le système ainsi formé peut devenir quasiment impossible à gérer automatiquement. Ce caractère est encore accentué par la possible émergence d'un comportement de groupe qui ressort des modèles fractals ou chaotiques. Ceci est notamment dû à la présence de l'homme en tant qu'utilisateur du système.

De ce fait, pour commander l'ensemble du système il est nécessaire de le hiérarchiser afin d'appliquer une commande pertinente sur chaque niveau mis en évidence dans la hiérarchisation. Le terme système hiérarchique fait référence à 3 phénomènes différents [BRE 84] :

- Il peut être utilisé pour la compréhension du système. Par exemple une structure hiérarchique peut réduire la complexité du système et le rendre compréhensible.
- Il peut faire référence à une organisation particulière du système.
- Il peut faire référence à un principe particulier pour le contrôle du système. Ce cas correspond à notre cas concernant le système de gestion d'une ligne de métro.

Un système peut être analysé grâce à une hiérarchie sans pour autant être réellement hiérarchique. Il peut aussi avoir une structure hiérarchique sans aucune structure de contrôle hiérarchique réelle.

Le système de gestion doit être décomposé en plusieurs niveaux interdépendants les uns des autres [ERS 84] où chaque commande est adaptée au niveau sur lequel elle agit.

Dans la représentation du système par une structure hiérarchique, plus le niveau est bas plus on se rapproche des composants du système. Plus on monte dans la hiérarchie plus on

s'éloigne des composants du système pour prendre en compte des fonctions de plus en plus complexes.

Chaque niveau doit prendre en compte les informations sur les niveaux précédents ainsi que les commandes que l'on a appliquées aux niveaux voisins. De plus chaque niveau doit posséder un ensemble de contraintes qui lui est propre. Ainsi, par exemple le niveau régulation d'un système de transport devra prendre en compte les contraintes locales propres à ce niveau comme le respect de la table horaire ainsi que les commandes appliquées au niveau précédent.

Ainsi, chaque niveau doit posséder une autonomie propre permettant de définir facilement des contraintes sur ce niveau. Mais cette autonomie est relative car pour un bon fonctionnement de l'ensemble du système, il est impératif de prendre en compte les commandes imposées par le niveau inférieur [ERS 76].

Cette approche multiniveaux entraîne un certain nombre de conséquences :

- Adaptation des modèles selon le niveau. Le caractère des problèmes traités aux différents niveaux devient de plus en plus global au fur et à mesure que l'on considère des niveaux plus élevés de la structure. C'est pourquoi, il faut adapter les modèles utilisés et leur degré de finesse à la nature du problème de décision propre à chaque niveau. Ceci implique donc des notions d'agrégation de données et d'informations pour définir de tels modèles.

- Adaptation de l'horizon selon le niveau. La prise de décision à un certain niveau, si on la considère comme un processus dynamique, suppose connus deux intervalles de temps : un horizon de décision, qui correspond à l'intervalle de temps sur lequel la décision retenue est exécutoire ; un horizon de prévision, qui correspond à l'intervalle de temps qu'il faut considérer pour prendre convenablement la décision. Dans l'approche multiniveaux, l'horizon de prédiction et l'horizon de prévision décroissent lorsque l'on descend dans la structure.

- Autonomie des niveaux. Dans une telle structure, une décision prise à un niveau définit les objectifs du niveau inférieur qui constituent pour lui des contraintes. Mais pour qu'un niveau de décision ait un rôle réel, il doit disposer d'une certaine autonomie dans l'exécution des décisions du niveau supérieur. Ceci lui permet en particulier de s'adapter aux perturbations qu'il doit prendre en compte sans remettre en cause la décision du niveau supérieur. Pour avoir une telle autonomie, il est nécessaire que la décision prise à un certain niveau laisse au niveau inférieur un choix de solutions possibles pour ses propres décisions. Ceci laisse donc à chaque niveau une certaine indépendance. Cette autonomie est en fait essentielle pour un fonctionnement harmonieux de l'ensemble de la structure, car elle permet à chaque niveau de

traiter les perturbations qui sont de sa compétence sans remettre en cause, en général, les décisions des niveaux supérieurs.

- Flexibilité des décisions à un niveau. Pour assurer sa fonction dans le cadre de la structure, chaque niveau réalise les objectifs du niveau supérieur par une suite de décisions élémentaires qui s'enchaînent. Dans de tels processus séquentiels de décision, un point important semble être la mesure de flexibilité que chaque décision laisse pour la prise des décisions ultérieures. En effet, c'est par cette flexibilité laissée par la décision qu'il est possible d'adapter les décisions d'un niveau aux perturbations produites sans remettre en cause la décision d'un niveau supérieur.

Cette approche multiniveau a été développée pour les systèmes complexes classiques. Nous pouvons étendre ce concept à la commande floue des systèmes complexes.

3.2. Extension aux systèmes flous

Il existe deux types de situations : la commande et l'aide à la conduite [FOU 94].

- La commande monovariante concerne les processus de type mono-entrée, mono-sortie. Les contrôleurs flous sont des extensions des contrôleurs classiques (PI, PD, PID), car les contrôleurs flous ont de larges plages de fonctionnement non linéaires, ce qui leur confère une robustesse vis-à-vis des perturbations tant internes qu'externes. Dans le cas où un processus multivariable peut être découplé, sa commande s'effectue à l'aide de plusieurs contrôleurs monovariants, et l'on est ramené au cas précédent.

- Les systèmes d'aide à la conduite concernent des processus qui possèdent déjà des boucles classiques de contrôle ; les techniques floues peuvent être utilisées pour différentes tâches de supervision du processus telles que le réglage automatique, le diagnostic, l'aide à la décision, la détection de défauts, etc. Mais rien n'empêche que les boucles classiques soient réalisées par des contrôleurs flous.

Si l'on considère seulement la commande des processus, chaque niveau à commander correspond à une commande réalisée par un contrôleur flou. Les contraintes de chaque niveau sont comprises dans les règles de production du contrôleur flou de ce niveau ainsi que dans la définition de l'univers du discours.

On peut concevoir de deux façons différentes l'interaction entre les différents niveaux. Les contrôleurs flous peuvent être reliés entre eux par un coordinateur qui agence les entrées et

les sorties de chaque contrôleur pour donner au système une commande adaptée (figure 2.4) [JAM 95].

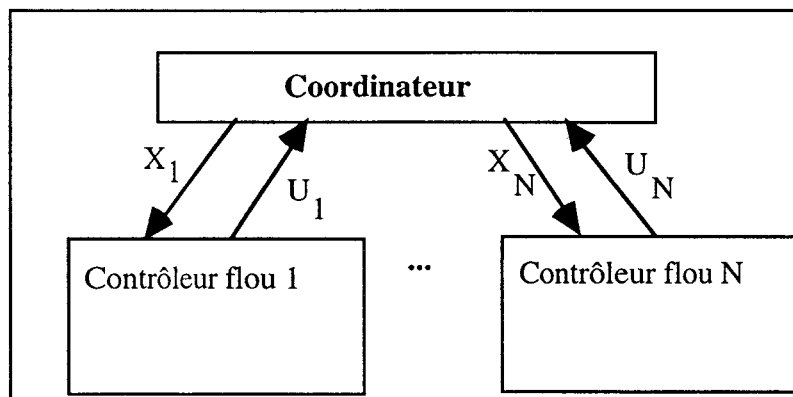


Figure 2.4. Système coordonné

L'interaction entre les niveaux peut être prise en compte par la sortie du niveau inférieur qui sera considérée comme une entrée d'un niveau. Dans ce cas on aura un système hiérarchisé (figure 2.5) comme le système décrit dans le paragraphe précédent.

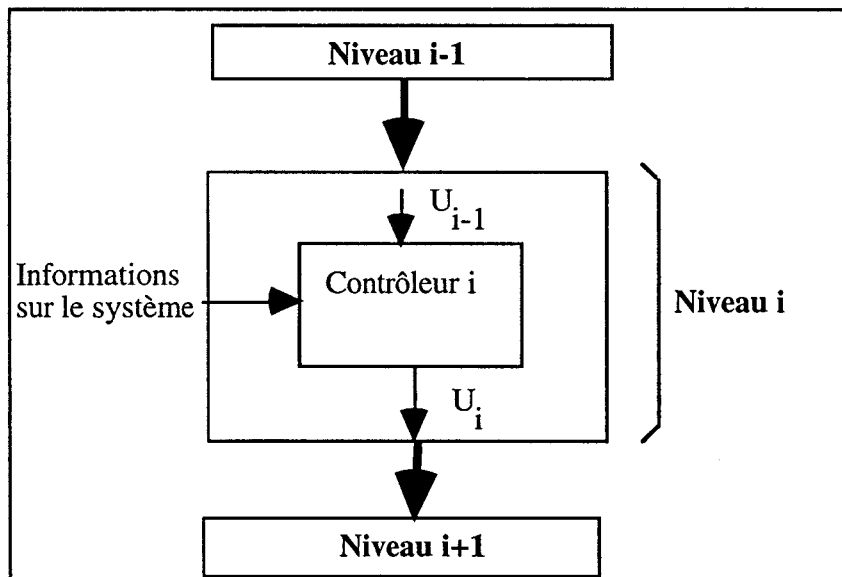


Figure 2.5. Structure hiérarchique appliquée à la commande floue

La sortie du dernier niveau N sera la commande réellement appliquée au système. Cette façon de procéder permet d'établir des contrôleurs flous interdépendants les uns des autres. Ils sont néanmoins simples à réaliser car chaque contrôleur est défini d'après des contraintes locales. Ainsi le contrôleur du niveau régulation d'un système de transport sera réalisé d'après les contraintes dérivées d'une table horaire de référence ; ce qui simplifie sa mise en oeuvre. A

ce niveau, on ne prendra en compte que les données disponibles sur le système nécessaires au respect des contraintes définies.

De ce fait, en vue de la réalisation d'un contrôle par logique floue d'une ligne de métro automatisée, nous allons montrer pourquoi cette ligne de métro et son système de gestion peuvent être considérés comme un système hiérarchique.

3.3. Décomposition hiérarchique de système de gestion d'une ligne de métro

Dans un premier temps, toutes les informations disponibles sur l'état du système sont centralisées et traitées au Poste de Commande Centralisé (PCC). Ce PCC constitue le niveau le plus élevé dans la hiérarchie.

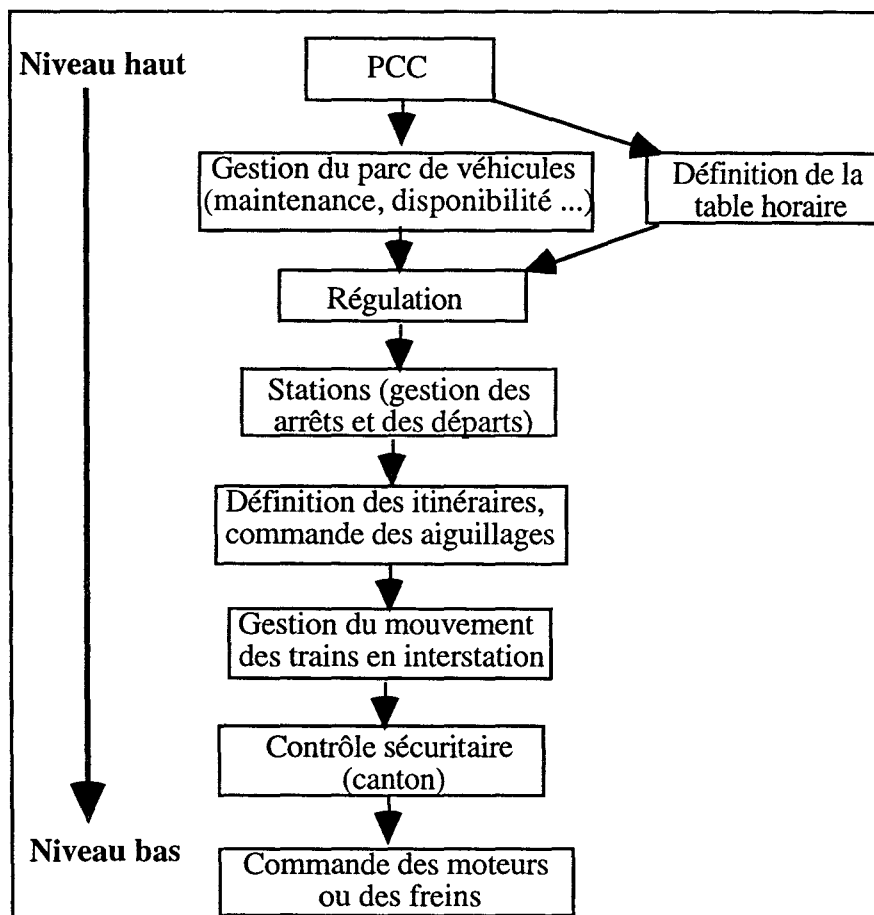


Figure 2.6. Décomposition hiérarchique d'une ligne de métro automatisée

Cette représentation se veut simpliste et ne décompose pas l'ensemble de toutes les fonctionnalités du système. Par exemple, la définition de la table horaire peut être décomposée en plusieurs niveaux [ADA 93]. Tout d'abord, les horaires sont créés hors ligne, en fonction du nombre prévu de voyageurs présents sur la ligne de métro durant les différents jours de la semaine. Le niveau inférieur prendra en compte les événements imprévisibles à long terme (par exemple une manifestation) pour redéfinir une table horaire différente de celle prévue. Au début de chaque journée, une table horaire doit être sélectionnée. La table horaire de chaque journée est, elle-même, décomposée en une succession d'intervalles différents selon l'heure de la journée (heures creuses, heures de pointe). La régulation en ligne prendra en compte les horaires théoriques de référence définis par la table horaire pour cette tranche d'heure.

Nous avons décidé d'utiliser la logique floue au niveau régulation et au niveau des stations, plus particulièrement, dans l'autorisation ou l'interdiction des départs des rames des stations.

Dans notre régulation, nous utilisons deux niveaux. Le premier niveau prend en compte le retard des rames pour permettre de réguler les rames selon l'horaire théorique de référence. Pour chaque rame à commander, nous régulons d'après trois informations : le retard de la rame et les retards des rames amont et aval. Les autres informations susceptibles d'être disponibles sur la (les) cause(s) des retards constatés ne sont pas prises en compte à ce niveau. Le deuxième niveau permet de prendre en compte les événements qui se sont déroulés sur la ligne pour interdire les départs des rames des stations où elles se trouvent. Dans le niveau supervision représenté sur la figure 2.7, nous observons si une rame est perturbée sur la ligne et nous estimons son temps de perturbation. D'après ces informations, nous agissons sur les rames pour éviter les situations de conflits et permettre un fonctionnement le moins dégradé possible de l'ensemble de la ligne.

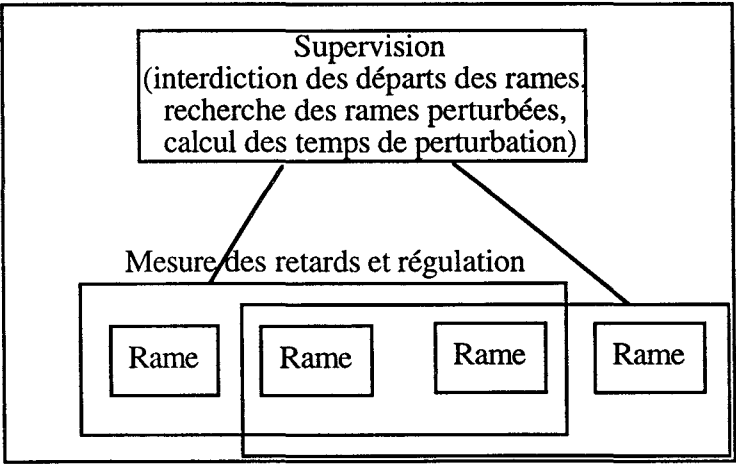


Figure 2.7. Structure hiérarchique de la régulation

CHAPITRE 3

APPLICATION DE LA LOGIQUE FLOUE A LA REGULATION DU VAL DE LILLE

1 . Introduction

L'amélioration de la qualité de service des systèmes de transport automatisés guidés passe d'abord par une bonne exploitation du trafic. Pour cela, il faut faire évoluer les algorithmes actuels de régulation à l'aide de nouveaux outils permettant de prendre en compte les fonctionnements pouvant être soumis à de l'incertain. Les outils développés à partir de la logique floue répondent à cette attente.

Le but est d'utiliser les sous ensembles flous pour permettre de réguler le trafic d'une manière qui se rapproche le plus possible d'un contrôle humain. Le principal avantage d'un contrôleur humain est que celui-ci a une expérience des différentes perturbations de la ligne. Il régulera le trafic en fonction des problèmes déjà rencontrés.

Nous considérons pour cela deux niveaux de régulation. Le niveau le plus bas régulera les rames en fonction des retards constatés lors de l'exploitation. Il s'agit de réguler au mieux en fonction d'objectifs définis préalablement. Le deuxième niveau permet d'améliorer la qualité de service du réseau en agissant lors de perturbations plus importantes qui modifient la marche normale des véhicules.

Dans un premier temps, nous présentons le premier niveau de régulation réalisé à l'aide d'un contrôleur flou.

2 . Commande floue

2.1. Définition du modèle

Le but d'une régulation est de permettre au système de transport d'avoir un comportement optimal pour transporter le maximum de personnes. Pour cela, des horaires

théoriques de référence ont été élaborés en fonction de la capacité nominale de la ligne et du nombre de personnes à véhiculer, dans des tranches horaires prédéfinies. Le but de la régulation est de respecter ces horaires théoriques de passages en stations afin de permettre à la ligne d'avoir un fonctionnement optimal (la régulation doit respecter les horaires théoriques de référence). Mais, vis à vis du passager attendant sur le quai, il est indispensable de lui assurer une attente limitée (la régulation doit assurer un intervalle régulier entre les rames).

Pour cela, les deux objectifs indissociables de la régulation sont :

- respecter le mieux possible les horaires théoriques de référence,
- conserver le plus possible des intervalles réguliers entre les rames.

La première étape dans la définition du modèle est d'obtenir une image du retard de chaque rame. A chaque période d'échantillonnage, nous observons la position de la rame. Comme nous connaissons le temps mis par cette rame pour arriver à cette position, nous pouvons en déduire son retard à cet instant d'échantillonnage (figure 3.1).

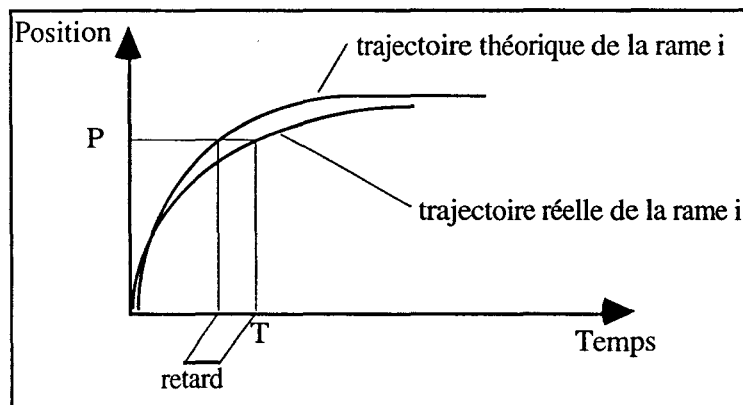


Figure 3.1. Déviation par rapport à l'horaire théorique

Pour la construction du modèle, il est nécessaire de définir le système de commande.

Les entrées X seront les retards des rames par rapport à l'horaire théorique de référence. X_i correspond donc au retard de la rame i .

Les sorties Y seront les temps à rattraper ou à perdre pour les rames sur la portion de voie suivante. Y_i est donc le temps que doit rattraper la rame i .

Un contrôleur flou prend en compte le retard de chaque rame de la ligne pour commander la rame i (figure 3.2). Ceci permet d'avoir une vision globale de la ligne et de gérer le trafic de façon optimale.

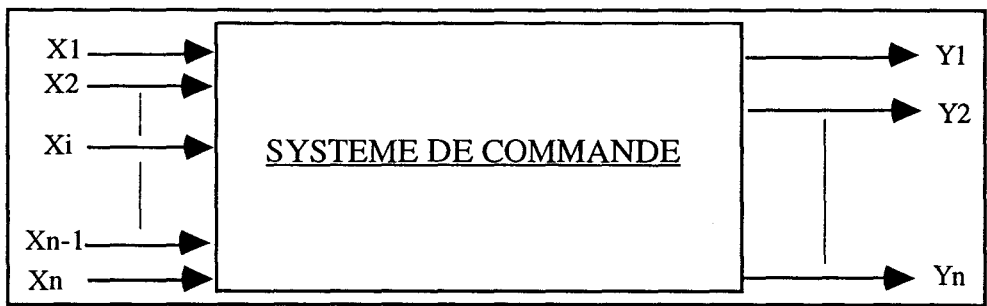


Figure 3.2. Architecture globale du système

Dans une telle architecture, le contrôleur mesure, à chaque pas d'échantillonnage, le retard de chaque rame de la ligne et donne une commande à chaque rame. Pour commander la rame i , le contrôleur utilisera le retard de toutes les rames de la rame 1 à la rame n . De toute évidence, cette façon de procéder nous donnerait un nombre considérable de règles utilisées par le contrôleur flou. Ce contrôleur serait très difficile à réaliser à cause du nombre important de règles et de sa dépendance au nombre de rames en ligne.

De ce fait, nous avons simplifié le modèle : pour commander une rame i seule l'image des retards de cette rame X_i , de la rame amont X_{i-1} et de la rame aval X_{i+1} est nécessaire. Lorsque la période d'échantillonnage est faible, la répercussion des retards sur les rames amont et aval de la ligne lors d'une perturbation sera relativement rapide ; ainsi l'influence de la perturbation est bien prise en compte sur l'ensemble du système. Cette simplification ne nuira donc pas à la régulation mise en oeuvre. La nouvelle architecture peut se représenter de la façon suivante :

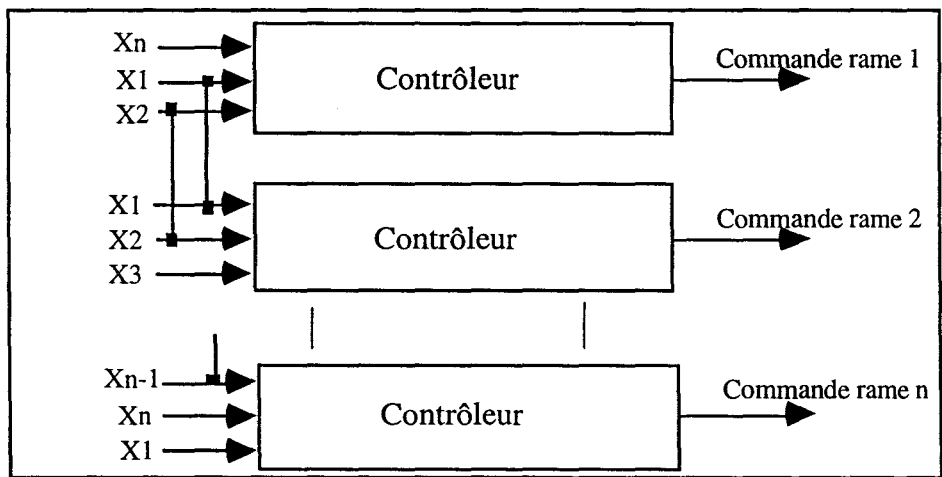


Figure 3.3. Architecture simplifiée

Le système de commande donne en sortie des consignes qui devront être interprétées par chacune des rames. A chaque rame, on associera un système de commande propre qui conditionnera l'évolution de l'état de chacune des rames. Une observation des rames nous montre clairement que la stratégie de contrôle dépend du contexte dans lequel évolue la rame. Ainsi quatre modes opératoires différents se distinguent :

- **mode 1** : La rame est à l'arrêt et l'objectif est de la maintenir arrêtée pendant le temps nominal d'arrêt.
- **mode 2** : La rame est en phase d'accélération avec objectif prioritaire d'atteindre une vitesse de consigne.
- **mode 3** : La rame a atteint la vitesse de consigne et l'objectif est de la maintenir à cette vitesse.
- **mode 4** : La rame est en phase de décélération et l'objectif est l'arrêt de la rame en un point précis.

Ces quatre modes sont visibles ci-dessous sur le diagramme de vitesse simplifié de la rame entre deux stations consécutives :

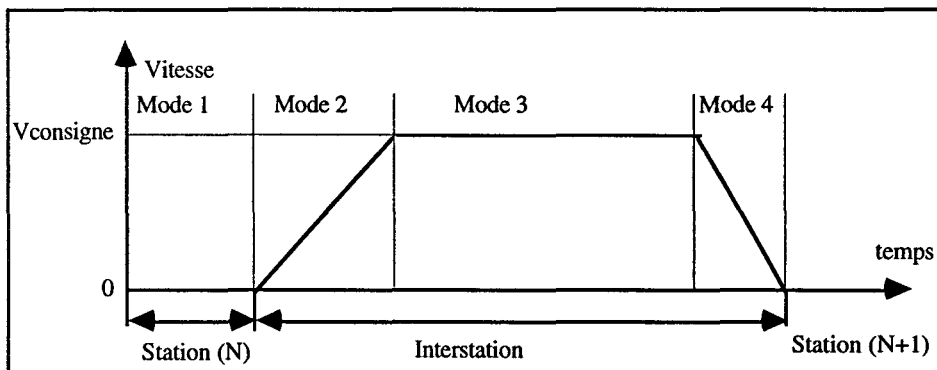


Figure 3.4. Allure du profil théorique de vitesse

Pour réguler le trafic, nous pouvons procéder de la façon suivante :

Les entrées du contrôleur flou seront le triplet des retards (ou avances) X_i , X_{i+1} et X_{i-1} . La sortie du contrôleur sera le temps à rattraper, sur la prochaine portion de voie, pour la rame i , Y_i . Un module en aval du contrôleur permettra de convertir ce temps en une consigne d'accélération (figure 3.5). Cette consigne dépendra de la position de la rame sur le profil de vitesse.

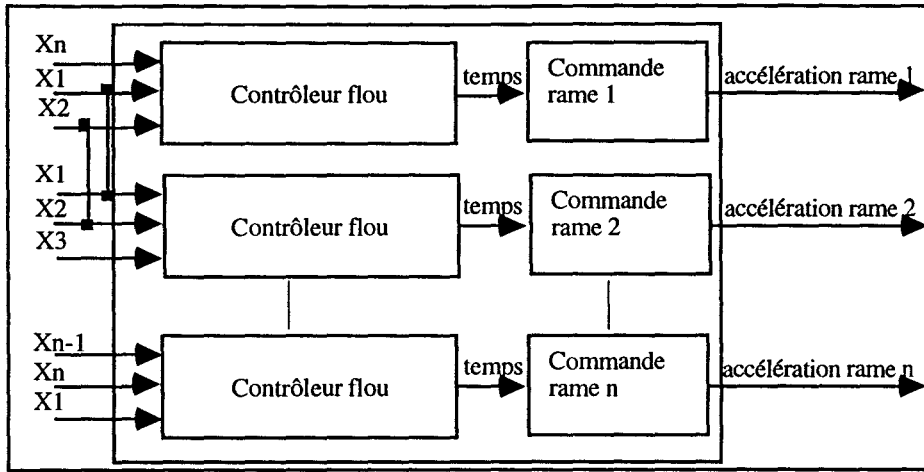


Figure 3.5. Architecture globale du modèle

Les architectures de tous les contrôleurs flous sont identiques. De ce fait, un seul jeu de règles sera écrit et de la même manière tous les modules de commande seront identiques.

2.1.1. Structure du module de commande

La structure de ce module dépend de l'emplacement de la rame sur le profil de vitesse.

Si la rame est en mode 1, elle est arrêtée en station, alors nous modifierons le temps d'arrêt en station en fonction de son retard (ou de son avance). Par exemple, si le contrôleur flou calcule un temps à rattraper de 4 secondes et si la rame est arrêtée en station, la commande sur la rame sera : diminuer le temps d'attente en station de 4 secondes.

Lorsque la rame est en mode 2 ou 3, un algorithme calcule une nouvelle accélération (ou décélération) permettant à la rame de gagner, par rapport à l'horaire théorique de référence, le temps fourni par le contrôleur flou. Cette nouvelle accélération tient évidemment compte des consignes de sécurité : ne pas dépasser l'accélération maximale, ne pas dépasser la vitesse limite autorisée sur la portion de voie correspondante, etc.

Lorsque la rame est en mode 4, nous ne pouvons pas changer la décélération durant le freinage de la rame. La fin du mode 3 et le début du mode 4 sont définis par un algorithme cinématique calculant le début du freinage de la rame pour arriver en station avec une vitesse nulle. Ce moment précis de changement de zone est calculé sachant que la rame a une décélération constante D_{nom} . Si l'on modifie cette décélération, la rame risque de "rater la station" ; c'est-à-dire qu'elle peut s'arrêter avant la station ou s'arrêter après. Dans les deux cas, une panne "station ratée" serait générée puisque les voyageurs ne peuvent pas descendre. Un

simple test de la position de la rame sur le profil de vitesse pourra être effectué avant d'appeler le contrôleur flou.

2.1.2. Les fonctions d'appartenance

L'efficacité du contrôleur flou est fortement influencée par le choix des fonctions d'appartenance des entrées et de la sortie.

2.1.2.1. Les entrées

Dans un premier temps, pour diminuer la complexité et le nombre des règles, nous avons limité, pour chaque entrée, la couverture de l'univers de discours par trois termes linguistiques. De ce fait, le nombre maximum de règles que nous obtiendrons sera égale à $3^3 = 27$.

Les fonctions d'appartenance des entrées X_{i-1} et X_{i+1} sont identiques à la fonction d'appartenance de l'entrée X_i .

Le choix des bornes des fonctions d'appartenance doit être fait après une étude qualitative sur l'importance des pannes sur une ligne de métro. Seule l'étude d'une ligne peut nous dire si une panne de deux minutes, par exemple, est une perturbation très pénalisante pour la ligne. Cette étude doit être faite sur une ligne de métro automatisée ayant un certain nombre d'années d'exploitation ainsi les perturbations rencontrées seront significatives. Nous avons utilisé la ligne 1 du VAL de Lille qui répond à cette contrainte [COU 88].

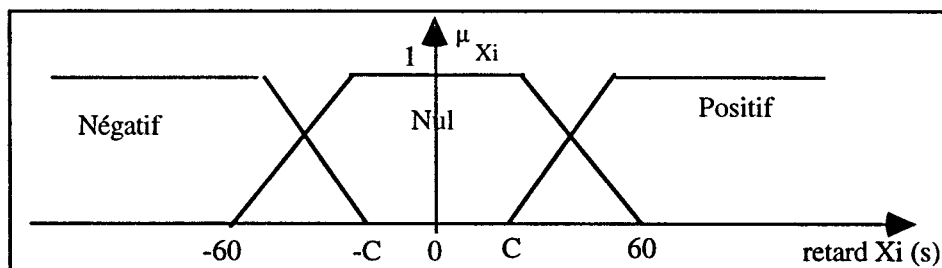


Figure 3.6. Fonctions d'appartenance pour l'entrée X_i

La borne C doit être choisie en fonction du retard que l'on considère comme nul, nous prendrons C égal à 1 seconde.

2.1.2.2. La sortie

La sortie est le temps qui doit être rattrapé par la rame dans la prochaine portion de voie. Celle-ci est plus difficile à classer par partitionnement de l'univers de discours. Nous considérons que le contrôleur donne un temps à rattraper au système et que la rame ne peut pas toujours l'exécuter. En effet, si nous demandons à la rame de rattraper 60 secondes étant donné que la période d'échantillonnage est relativement faible (d'environ une seconde), il est évident que la rame ne peut rattraper tout ce temps en un instant d'échantillonnage. Mais le fait de fournir un temps important permettra d'accélérer fortement la rame, et, de ce fait, elle reviendra plus vite à son fonctionnement nominal.

Nous classerons la sortie en quatre zones de variations : négative, nulle, positive et grande positive. Une commande grande négative n'a pas vraiment de sens car, par définition, le comportement naturel de la ligne engendre des retards et non de l'avance ; de ce fait nous aurons plus de grands retards à rattraper que de grandes avances.

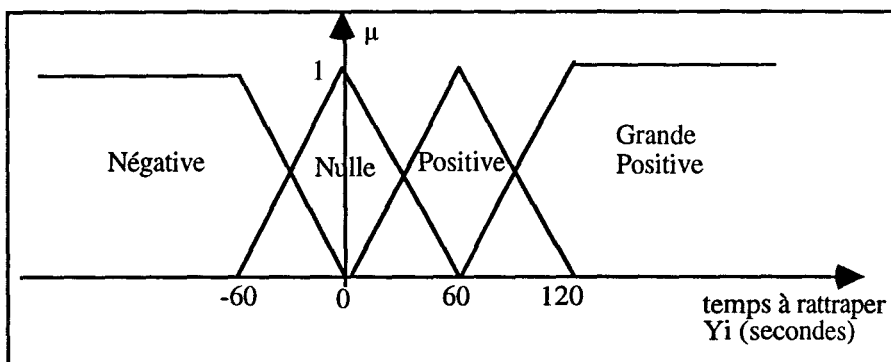


Figure 3.7. Fonction d'appartenance pour la sortie Y_i

Le contrôleur flou utilisé sera du type Mamdani [MAM 77].

Remarque : Une partie de l'opération d'inférence permettant l'agrégation des règles est réalisée à l'aide de l'opérateur max et la défuzzification est effectuée par la méthode du centre de gravité. Ce qui implique que les bornes de la fonction d'appartenance de la sortie Y_i doivent être choisies pour que la valeur numérique de Y_i issue de la défuzzification soit cohérente avec la ligne de métro à réguler.

Nous constatons, lors des premiers essais, une absence de stabilité pour les retards proches de zéro. En resserrant les bornes des fonctions d'appartenance, nous augmentons la stabilité ; par contre, nous diminuons l'importance de la commande pour les retards élevés. Il faut donc trouver un compromis entre une commande importante pour que la rame récupère son

retard rapidement et une commande pas trop élevée pour les petits retards afin de stabiliser système.

La première solution est de fragmenter plus finement l'univers de discours pour mieux contrôler la stabilité pour les faibles valeurs et conserver une commande importante ailleurs. Mais le fait d'augmenter le nombre d'éléments constituant l'univers de discours accroît le nombre de règles d'une façon combinatoire, ce qui peut poser des problèmes de temps de calcul.

L'autre solution consiste à mettre en oeuvre deux contrôleurs flous. Le premier traite les petits retards, avec des règles privilégiant la commande Y_i nulle (pour assurer la stabilité). Le deuxième traite les retards plus importants, avec des règles qui favorisent les commandes positives et minimisent le temps mis par la rame pour revenir à son comportement nominal.

Les fonctions d'appartenance pour les deux contrôleurs flous ont les mêmes formes triangulaires, seuls, les supports changent. Les bornes seront plus élevées pour le contrôleur 2 s'occupant des grands retards et moins élevées pour le contrôleur 1 s'occupant des petits retards.

Pour le contrôleur correspondant au cas où le retard est supérieur à 20 secondes, on privilégie les commandes négatives pour répartir l'intervalle entre les rames : lorsque le retard est classé positif, il est forcément supérieur à 20 secondes. Par contre, pour le contrôleur 1, nous privilégierons les commandes nulles puisque le retard sera forcément inférieur ou égal à 20 secondes. Ceci évitera de provoquer un retard sur les autres rames. Cette façon de procéder nous amène à écrire des règles spécifiques pour les contrôleurs 1 et 2.

2.1.3. Les règles

Les règles du contrôleur seront de la forme :

Si X_{i-1} est Nul et si X_i est Nul et si X_{i+1} est Nul
alors Y_i est Nul.

Pour répartir l'intervalle entre les rames nous utiliserons les règles du type :

Si X_{i-1} est Positif et si X_i est Nul et si X_{i+1} est Nul
alors Y_i est Négatif.

Pour respecter l'horaire théorique des rames nous utiliserons les règles du type :

Si X_{i-1} est Nul et si X_i est Positif et si X_{i+1} est Nul
alors Y_i est Positif.

Nous allons maintenant présenter les bases de règles des deux contrôleurs flous.

2.1.3.1. Premier Contrôleur

Par définition ce premier contrôleur s'occupe des petits retards, donc nous privilégierons le retour des rames à leur horaire théorique de référence sans répartir le retard sur les rames amont et aval.

x _R		X _{i-1}								
		NG			Z			PO		
		X _{i+1}								
		NG	Z	PO	NG	Z	PO	NG	Z	PO
X _i	NG	NG	NG	NG	NG	NG	NG	NG	NG	NG
	Z	Z	Z	Z	Z	Z	Z	Z	Z	Z
	PO	PO	PO	PO	PO	PO	PO	PO	PO	PO

Avec NG : négatif, Z : Nul, PO : Positif, GP : Grand positif.

2.1.3.2. Deuxième Contrôleur

Ce contrôleur flou s'occupe des grands retards. Il permet la répartition des retards entre les rames tout en appliquant des commandes élevées aux rames ayant un retard positif.

x _R		Xi-1								
		NG			Z			PO		
		Xi+1								
		NG	Z	PO	NG	Z	PO	NG	Z	PO
Xi	NG	NG	NG	NG	NG	NG	NG	NG	NG	NG
	Z	PO	PO	Z	PO	Z	NG	Z	NG	NG
	PO	PO	PO	PO	PO	PO	PO	PO	PO	GP

2.2. Simulation sur la ligne 2 du VAL de Lille

Dans la régulation type VAL, la mesure du retard des rames par rapport à leur horaire théorique de référence, ainsi que les commandes découlant de ces retards à appliquer sur les rames, ne sont effectuées qu'en station par l'intermédiaire des EAS (Électronique d'Arrêt en

Station). Une fois la rame envoyée, il ne sera plus possible d'intervenir tant qu'elle ne sera pas arrivée à la prochaine station. Les conséquences d'une action appliquée sur une rame en station ne seront visibles que sur le retard de la rame à la station suivante. Nous avons ajouté à ces EAS dans la simulation le contrôleur flou permettant de commander la rame arrêtée en station en fonction de son retard et des retards des rames amont et aval. C'est à dire qu'à chaque arrêt d'une rame i en station :

- son retard par rapport à son horaire théorique de référence est calculé,
- les retards de la rame i , de la rame amont et aval sont envoyés au contrôleur flou,
- le résultat de l'inférence floue est obtenu sous la forme d'un temps à rattraper (ou à perdre), pour la rame i , sur la prochaine inter-station.

Lors de l'arrêt d'une rame en station, seul son retard est connu. Les retards des rames amont et aval ne sont pas connus si ces rames ne se trouvent pas en station puisque aucune mesure de position n'existe en inter-station. De ce fait nous devons calculer ces retards par un calcul d'approximation [BAI 3-95].

2.2.1. Mesure des retards des rames amont et aval

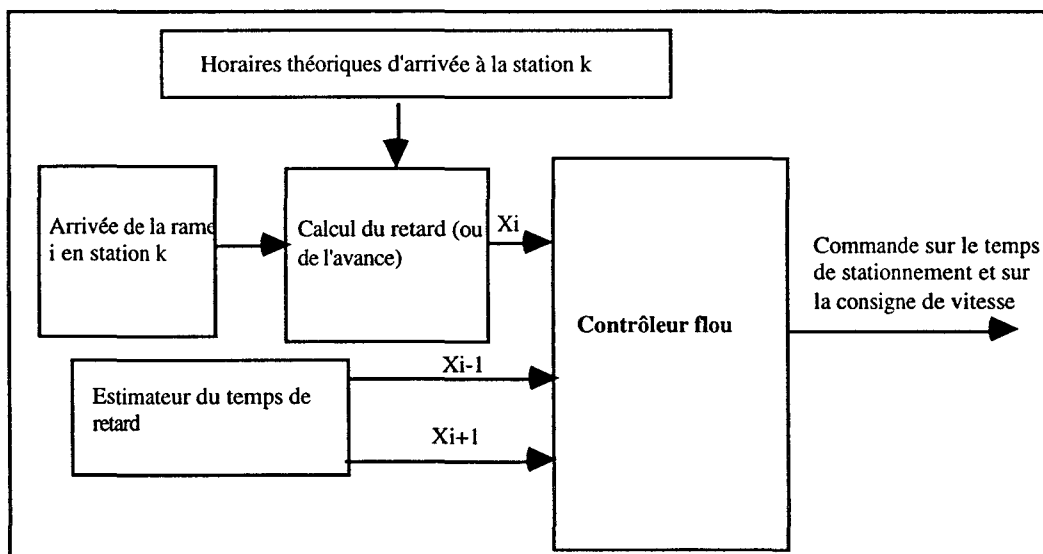


Figure 3.8. Mesure du retard et application de l'inférence floue

Contrairement à un premier modèle développé qui supposait connu le retard de toutes les rames à l'instant d'échantillonnage T [BAI 1-94] et [BAI 2-95], seul le retard de la rame i (X_i) est connu en considérant l'instant T comme l'instant où la rame i arrive à la station. Les

retards des rames amont et aval ont été calculés à l'instant où ces rames sont arrivées à une station. Ce qui veut dire qu'à l'instant T, les retards des rames $i-1$ et $i+1$ ne sont pas connus.

Il est donc nécessaire d'estimer le retard des rames à chaque instant d'échantillonnage pour que celui-ci corresponde au retard réel des rames à l'instant T.

L'estimation de ce retard doit se faire d'après :

- Le retard de la rame mesuré à la dernière station rencontrée.
- La commande sur la consigne de vitesse appliquée à la rame à cette station.
- Le temps de stationnement réel de cette rame.

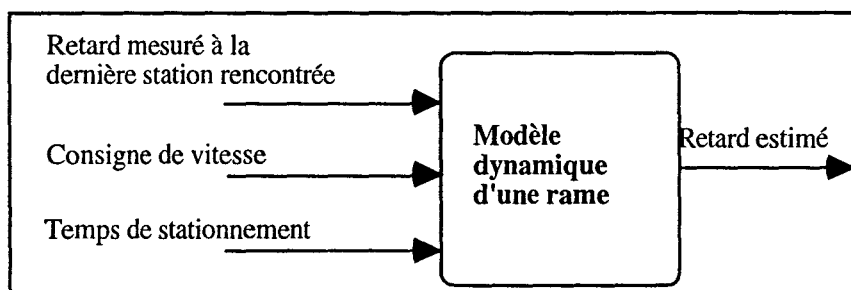


Figure 3.9. Calcul du retard estimé.

De ce fait la commande à appliquer à la rame i sera calculée en fonction des retards X_i , X_{i-1} et X_{i+1} à l'instant T.

2.2.2. Calcul du retard estimé

2.2.2.1. La rame est en station

Ce retard estimé est égal au retard réel lorsque la rame arrive en station puisqu'il vient d'être calculé grâce à la comparaison de l'heure d'arrivée en station et de l'horaire théorique de référence d'arrivée en station.

$$\hat{X}_i^t = X_i$$

Avec $\hat{X}_i^t$: retard estimé de la rame i à l'instant t et X_i retard de la rame à l'arrêt en station.

2.2.2.2. La rame repart de la station

Lorsque la rame repart de la station, le retard estimé est égal au retard mesuré en station, duquel on soustrait le temps gagné - ou perdu - en station.

$$\hat{X}_i' = X_i - (\text{temps d'arrêt nominal} - \text{temps d'arrêt réel})$$

2.2.2.3. La rame est en inter-station

Lorsque la rame est en inter-station, le retard estimé est égal au retard mesuré en station, moins le temps gagné - ou perdu - en station, moins le temps gagné par une consigne de vitesse supérieure à la consigne de vitesse nominale.

$$\hat{X}_i' = X_i - (\text{temps d'arrêt nominal} - \text{temps d'arrêt réel}) - (\text{temps rattrapé sur l'interstation suivante})$$

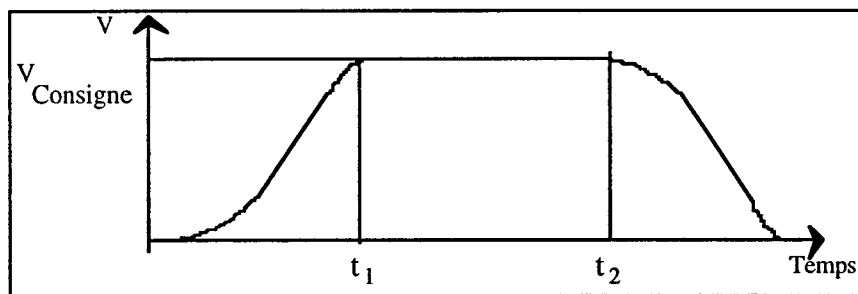


Figure 3.10. Profil de vitesse entre deux stations

Le temps gagné depuis l'instant de départ de la rame en station dépend de la position sur le profil de vitesse.

Soit T temps écoulé depuis le départ de la rame en station.

a) Si $T < t_1$

La rame est dans sa phase d'accélération, elle n'a pas encore atteint sa vitesse nominale.

$$t_r = \left(\frac{\text{accélération appliquée} - \text{accélération nominale}}{\text{accélération nominale}} \right) \times T$$

Avec t_r temps gagné (ou perdu) depuis la dernière station.

b) Si $T \geq t_1$ et $T \leq t_2$

$$t_r = \left(\frac{V_{\text{consigne appliquée}} - V_{\text{consigne nominale}}}{V_{\text{consigne appliquée}}} \right) \times (T - t_1)$$

c) Si $T > t_2$

La rame est en phase de décélération.

$$t_r = \left(\frac{\text{accélération appliquée} - \text{accélération nominale}}{\text{accélération nominale}} \right) \times (T - t_2)$$

Or, l'accélération est toujours égale à l'accélération nominale pour permettre à la rame d'arriver le plus vite possible à la vitesse de consigne et de s'arrêter précisément en station ; seule la vitesse de consigne peut donc être modifiée. De ce fait, nous ne gagnons du temps que lorsque la rame a atteint sa vitesse de consigne. C'est à dire qu'entre deux stations, le temps gagné par rapport à l'horaire théorique de référence dépend de la consigne de vitesse que l'on donne à la rame. Il est égal à :

$$t_r = \left(\frac{V_{\text{consigne appliquée}} - V_{\text{consigne nominale}}}{V_{\text{consigne appliquée}}} \right) \times (T - t_1)$$

Cette façon de procéder nous permet d'agir sur les rames s'arrêtant en station de la même façon que l'on agirait sur les rames à chaque pas d'échantillonnage.

2.2.3. Résultats

Nous avons testé l'algorithme sur la ligne 2 du VAL de Lille lorsqu'une rame est perturbée. Pour mesurer la validité de l'estimateur, nous calculons l'erreur moyenne commise (1) et l'erreur quadratique moyenne (2) à chaque arrêt en station.

$$\epsilon_1 = \frac{1}{n} \sum_{i=1}^n \frac{|\hat{X}_i^t - X_i|}{X_i} \quad (1)$$

$$\epsilon_2 = \sqrt{\frac{\sum_{i=1}^n (\hat{X}_i^t - X_i)^2}{n}} \quad (2)$$

Où t est l'instant où la rame i arrive en station, i le numéro de la rame et n le nombre de passages en station.

Temps de perturbation	ϵ_1	ϵ_2
0	0	0
30 s	0,0117	0,17
100 s	0,0127	0,1959
150 s	0,0157	0,24
200 s	0,0484	0,41

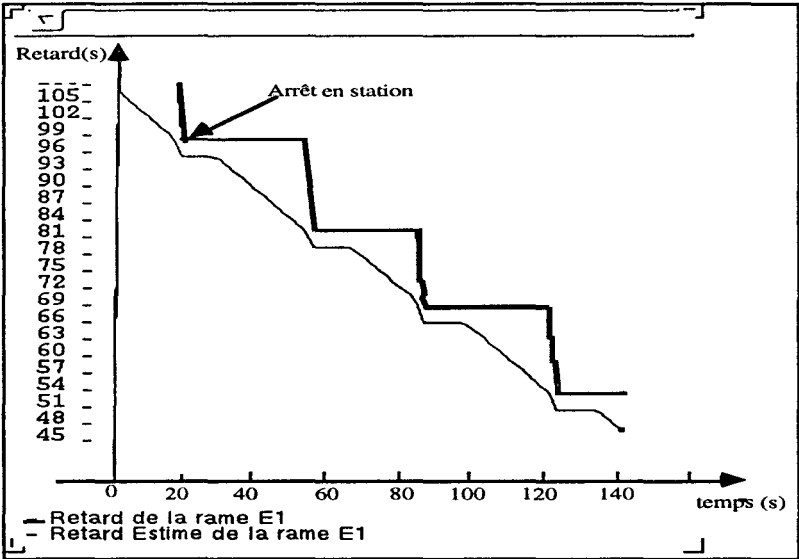


Figure 3.11. Evolution du retard et du retard estimé de la rame E1 en fonction du temps

Nous pouvons constater qu'à chaque arrêt en station, le retard estimé coïncide presque avec le retard mesuré en station. Nous constatons également que le retard estimé représente bien le retard de la rame à chaque pas d'échantillonnage.

2.2.4. Répercussion du résultat de l'inférence floue sur les rames

Comme il a été dit au paragraphe 2.1.1, une commande sera envoyée à la rame en fonction de la consigne. Cette commande sera appliquée au temps de stationnement et à la consigne de vitesse sur la prochaine inter-station. Un module en aval du contrôleur permettra de convertir le résultat de l'inférence floue en une consigne de vitesse et / ou une diminution (ou augmentation) du temps d'arrêt en station. Par exemple, si le contrôleur flou donne un temps à rattraper $t=4$ secondes et que la rame est arrêtée en station ; si, d'autre part, le temps d'arrêt en station ne peut être diminué au maximum que de 3 secondes ; alors la commande sera de

diminuer le temps d'attente en station de 3 secondes et d'augmenter la vitesse de consigne de la rame sur la prochaine inter-station pour que la rame rattrape 1 seconde par rapport à son parcours nominal. La consigne de vitesse peut être augmentée ou diminuée à condition qu'elle soit supérieure à une vitesse minimum et inférieure à la vitesse limite admissible sur le tronçon.

De plus, un module d'injection - retrait a été implémenté. Ce module réalise une approximation des heures de départ et d'arrivée des rames au terminus ; cette approximation se fait par un calcul des temps de parcours des rames proches du terminus [BAI 3-95]. Connaissant les heures futures de départ et d'arrivée des prochaines rames, nous décidons si une rame doit être injectée ou retirée du carrousel.

2.3. Conclusion

Dans un premier temps une commande floue a été réalisée permettant de réguler l'ensemble des rames par rapport à deux critères : le respect des intervalles entre les rames et le respect des horaires théoriques de référence. Dans cet objectif, un module de régulation a été créé pour permettre de réguler les rames à chaque pas d'échantillonnage. Dans un souci de faire correspondre notre modèle aux modèles de régulation existants, en vue d'une comparaison entre les deux types de régulation, nous avons adapté notre modèle pour que l'algorithme fonctionne avec une mesure des retards des rames à chacun de leurs arrêts en station.

Dans ce module de régulation, les perturbations ne sont pas prises en compte. Au début d'une panne, les rames amont et aval ne sont pas trop perturbées par le système de régulation. Cependant, la non perturbation volontaire de ces rames peut entraîner des conséquences importantes sur l'ensemble du réseau, comme le blocage complet de la ligne si la perturbation se prolonge. De ce fait, il serait intéressant de contrôler l'ensemble des rames circulant sur la ligne pour éviter les situations de blocage de l'ensemble du réseau. C'est dans ce but qu'un superviseur de l'ensemble des rames a été réalisé.

3. Supervision floue

3.1. Introduction

Jusqu'à présent nous ne régulons les rames qu'après un retard constaté. Or, dans certains cas, il serait préférable de savoir si une rame est perturbée avant d'arriver à des situations où la perturbation d'une rame entraîne sur d'autres rames.

Nous pouvons considérer que la régulation comporte plusieurs niveaux. Le premier niveau hiérarchique est sécuritaire et interdit tous les événements contraires à la sécurité. Le deuxième niveau est local aux rames. Ce niveau permettra de gérer le temps d'attente en station, de contrôler les organes de traction et de freinage pour permettre à la rame de suivre le profil de vitesse entre deux stations. Le troisième niveau est le contrôle flou des rames permettant de les réguler selon la table horaire et selon l'intervalle entre elles. Le quatrième niveau est une supervision de l'ensemble des mouvements des rames. Ce niveau agira en cas de fortes perturbations sur les rames pour éviter un blocage complet de la ligne. Les niveaux sont interdépendants les uns des autres et permettent d'assurer un fonctionnement optimal de la ligne sans menacer la sécurité (figure 3.12).

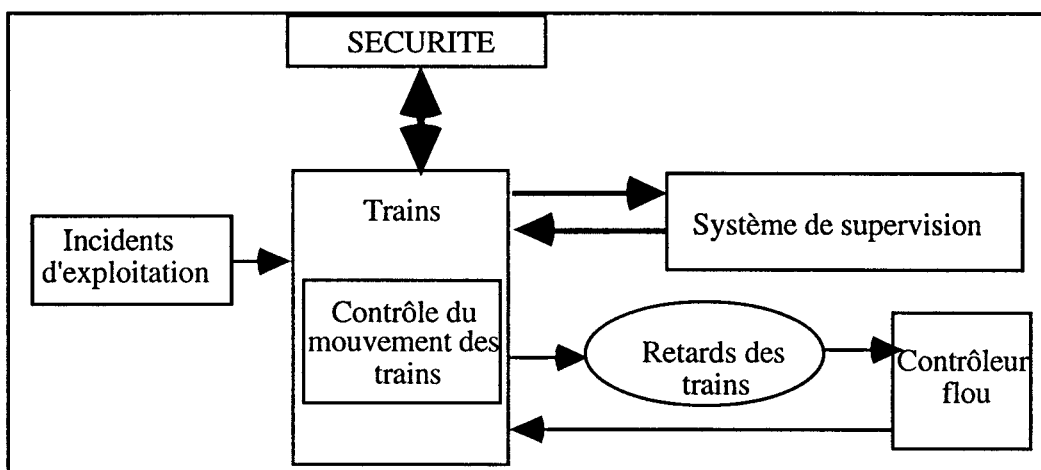


Figure 3.12. Synoptique de la régulation

La deuxième phase de l'étude nous a amené à nous demander quels étaient les cas où nous pouvions agir sur les rames lors des perturbations. Le but est d'éviter de faire partir une rame d'une station si dans l'inter-station suivante une rame est arrêtée suite à une perturbation. Cette façon de procéder permet d'éviter un entassement des rames en anticollision derrière une rame perturbée. Il est préférable pour l'usager d'être dans une rame bloquée en station d'où il peut descendre, plutôt qu'entre deux stations.

Nous étudierons les différents scénarii de pannes. Nous en dégagerons certaines situations afin d'éviter les arrêts en anticollision.

3.2. Scénarii de pannes

Pour expliciter l'algorithme, nous présentons les quatre cas de figures qui résument l'ensemble des scénarii, lors de perturbations.

3.2.1. La rame perturbée est arrêtée en station

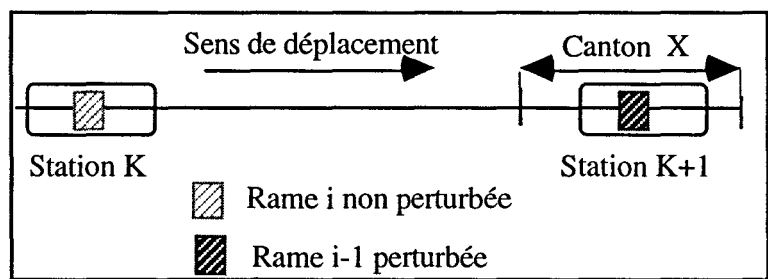


Figure 3.13. La perturbation se produit en station, la rame aval est en station

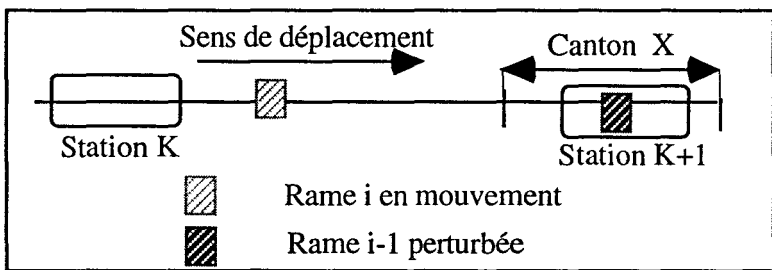


Figure 3.14. La perturbation se produit en station, la rame aval est en inter-station

3.2.2. La rame perturbée est arrêtée en inter-station

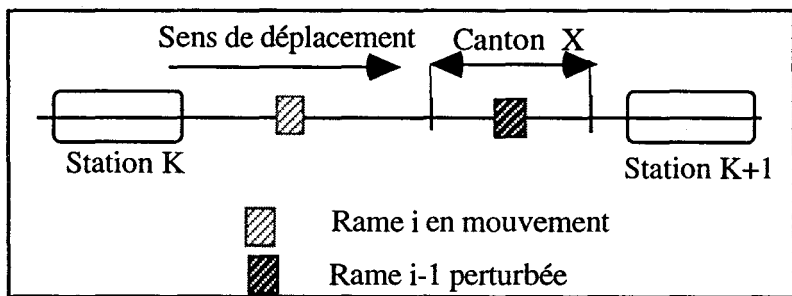


Figure 3.15. La perturbation se produit en inter-station, la rame aval est en inter-station

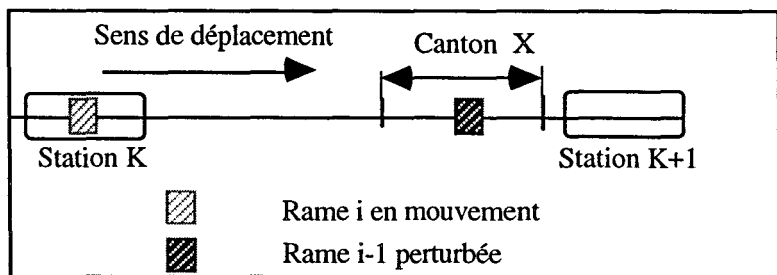


Figure 3.16. La perturbation se produit en inter-station, la rame aval est en station

Dans les cas des figures 3.14 et 3.15, la rame i est déjà partie de la station. Nous ne pouvons agir sur celle-ci si la perturbation se prolonge. Dans ces deux cas, il y aura arrêt de la rame i en anticollision à la limite du canton X.

Pour les cas des figures 3.13 et 3.16, une réflexion sur la stratégie nous conduit à nous demander si la rame i, se trouvant à la station K, doit être envoyée ou non. En effet, il est préférable, pour le confort des usagers, de bloquer une rame en station plutôt qu'en inter-station. Ces deux cas se ramènent donc à un seul.

3.3. Introduction d'une notion de risque

La décision de faire quitter la station K à la rame i doit être prise d'après les informations suivantes :

- Le retard de la rame i par rapport à son horaire théorique en station.
- Le temps mis par la rame i pour parcourir la distance entre la station K et le début du canton X. C'est le temps s'écoulant entre le départ de la rame i de la station K et son arrêt éventuel sur anticollision si la rame (i-1) est encore perturbée.
- La durée de la perturbation de la rame (i-1) ou plus précisément la durée séparant le départ de la rame i de la station K de la libération du canton X par la rame (i-1). Ceci pour éviter la mise en fonctionnement du dispositif d'anticollision.

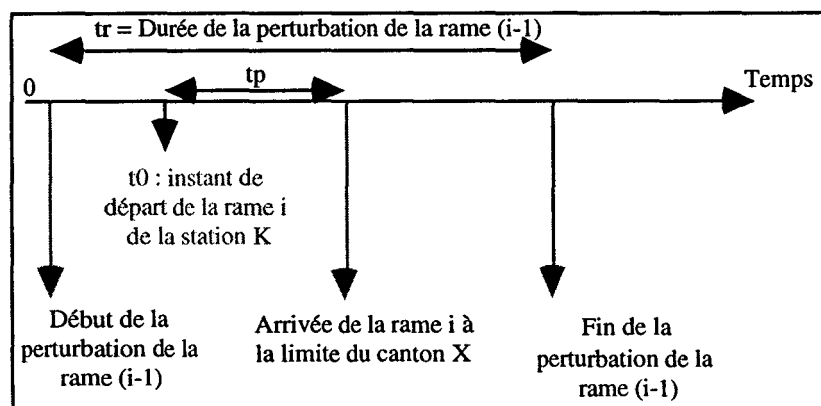


Figure 3.17. Comparaison entre les différents instants caractéristiques

La rame i ne s'arrêtera pas sur anticollision si :

$$t_r - t_0 + t_L < t_p \quad (1)$$

$t_r - t_0$, temps que va encore durer la perturbation de la rame (i-1) au moment du départ de la rame i.

t_L , temps mis par la rame (i-1) pour sortir du canton X.

t_p , temps de parcours de la rame i entre le départ de la station K et son arrivée à la limite du canton X.

Les données historiques des pannes constatées lors de l'exploitation nous permettent d'avoir une définition probabiliste de chaque panne survenue en fonction des situations dans lesquelles elle s'est produite : la rame est en station ou non, la période correspond à une heure de pointe ou à une heure creuse, la durée de la panne, le type de panne [CRO 94]. A tout moment, nous connaissons le temps qui s'est écoulé depuis le début de la panne sur la rame (i-1). Ce temps nous permet de réaliser une représentation probabiliste de la perturbation, à chaque instant, en fonction du temps écoulé depuis le début de la panne, ce qui permet d'obtenir une estimation d'une durée de la perturbation t_r en train de se produire sur la rame (i-1). Grâce aux transformations développées par Dubois - Prade [DUB 93] et Sandra Sandri [SAN 91], il est possible d'obtenir une fonction de possibilités représentant cette même estimation. Nous supposons donc que l'estimateur nous fournira un intervalle de temps avec une incertitude sous forme d'une distribution de possibilités. Une étude menée montre l'intérêt d'utiliser une distribution de possibilités par rapport à une distribution de probabilité [BAI 1-95]. La pertinence de la décision d'envoyer la rame i dépendra de la qualité de l'estimation du retard de la rame (i-1).

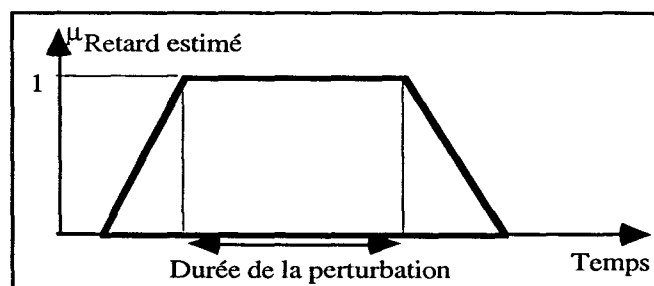


Figure 3.18. Fonction d'appartenance du retard estimé t_r

Le temps de parcours t_p de la rame i de la station K à la limite du canton X peut être considéré comme un intervalle flou. On se ramène donc à un problème de comparaison de nombres flous ; cette comparaison nous permettra de prendre une décision concernant l'envoi de la rame.

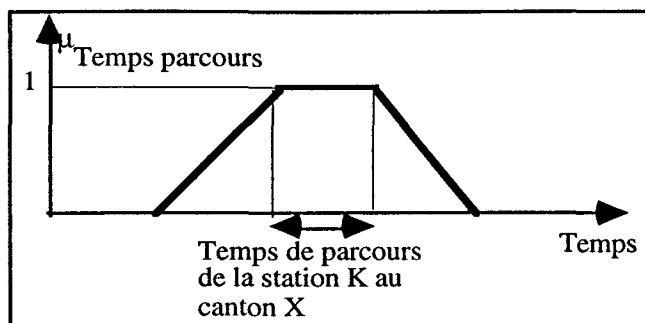


Figure 3.19. Fonction d'appartenance du temps de parcours t_p

Ce temps de parcours de la rame i a une valeur minimum (la distance parcourue à la vitesse maximum admissible) et une valeur maximum (la distance parcourue à la vitesse minimum admissible).

L'envoi de la rame i se décide à la suite d'une comparaison entre la fonction d'appartenance du retard estimé et celle du temps de parcours. Cette comparaison nous conduit à définir deux critères : le risque pris d'envoyer la rame et le non-risque. Plus le risque est grand, plus la rame i a des chances de s'arrêter en anticollision devant l'entrée du canton X. A l'inverse, plus le risque se rapproche de 0, moins la rame i a des chances de s'arrêter en anticollision.

Nous définissons un indice non-risque comme étant la fonction inverse du risque. Lorsque le risque est nul alors le non-risque sera maximum ; et lorsque le risque est maximum alors le non-risque est nul.

La prise de décision d'envoyer la rame sera effectuée par une comparaison de ces deux indices :

- Si $\text{risque} < \text{non-risque}$ alors nous laissons la rame i partir de la station K,
- si $\text{risque} \geq \text{non-risque}$ nous bloquons la rame i à la station K car il sera fortement possible que la rame i s'arrête en anticollision.

3.4. Calcul du "risque" et du "non-risque"

Nous ne connaissons pas l'instant précis où se déclenche la perturbation. La mesure des retards se fait en station, donc, nous savons qu'une rame est perturbée lorsque la station ne la voit pas arriver. Un estimateur du temps de retard de la rame a été intégré [BAI 2-96]. Le

déclenchement de cet estimateur se fera en fonction du temps de parcours des rames dans les cantons qu'elles occupent. Si une rame entre dans un canton de x mètres à l'instant t_1 et qu'elle n'en est pas ressortie à l'instant t_2 dépendant de la vitesse prévue sur ce canton et de la longueur de celui-ci, alors nous pouvons dire que la rame est certainement perturbée. De ce fait, nous déclencherons l'estimateur de pannes qui indiquera si la rame a une chance d'être perturbée de y secondes à cause de la panne numéro z .

3.4.1. Retard de la rame

On peut représenter le retard de la rame sous la forme d'un sous ensemble flou (figure-3.20).

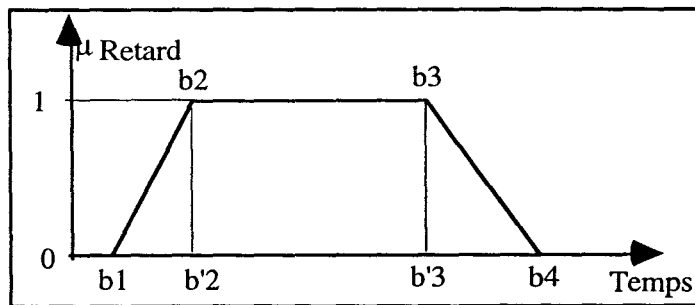


Figure 3.20. Fonction d'appartenance du retard estimé

$b'2$ = temps de perturbation de la rame (i-1) à l'instant t_0 , instant de départ de la rame i de la station K

$b'3$ = $b'2$ + temps maximum mis par la rame (i-1) pour libérer le canton X

$$b'3 = b'2 + \frac{\text{Position rame} - \text{position de la limite du canton X} + \text{longueur de la rame}}{\text{vitesse minimum de la rame (i - 1) sur le canton X}}$$

3.4.2. Temps de parcours de la rame

De la même façon le temps de parcours de la rame peut être schématisé selon la figure 3.21.

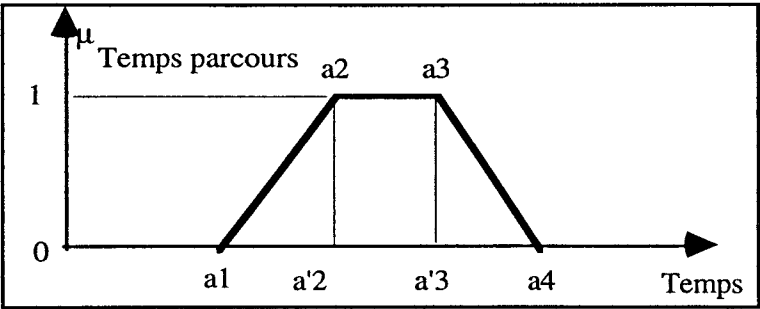


Figure 3.21. Fonction d'appartenance du temps de parcours

$a'2$ = Temps de parcours de la rame i de la station K jusqu'au canton X avec la vitesse maximale.

$a'3$ = Temps de parcours de la rame i de la station K jusqu'à la limite du canton de la station (K+1) avec la vitesse minimale.

Le calcul du risque et du non-risque est effectué par une comparaison des deux intervalles flous définis ci-dessus.

3.5. Différents indices de risque

Plusieurs indices de risque - et de non-risque - peuvent être calculés en comparant deux ensembles flous.

Le premier indice de risque sera calculé en comparant l'ensemble flou représentant le temps mis par la rame i pour aller de la station K au canton X, et l'ensemble flou représentant le temps mis par la rame $(i-1)$ pour sortir du canton X.

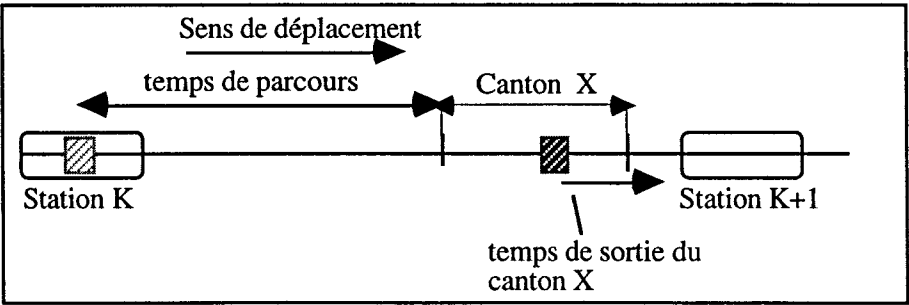


Figure 3.22. Eléments de calcul du premier indice de risque

Le deuxième indice de risque sera calculé en comparant l'ensemble flou représentant le temps mis par la rame i pour aller de la station K à la limite du canton de la station $K+1$, et l'ensemble flou représentant le temps mis par la rame $(i-1)$ pour sortir du canton de la station $K+1$.

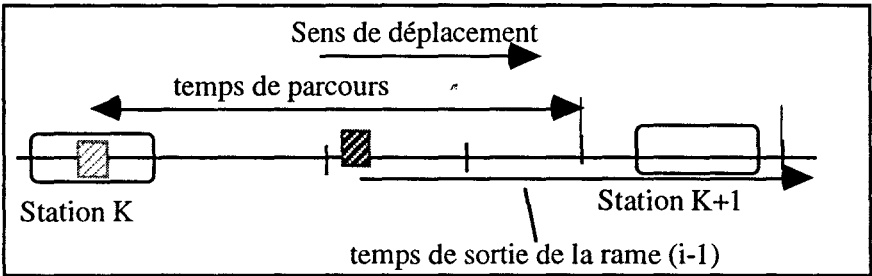


Figure 3.23. Eléments de calcul du deuxième indice de risque

Le temps de parcours de la rame i serait le temps mis par la rame pour parcourir la distance de la station K à la limite du canton de station. Ce temps est un intervalle dont la limite supérieure est cette distance parcourue avec la vitesse minimum admissible sur le tronçon et dont la limite inférieure est cette distance parcourue avec la vitesse maximum admissible.

Le temps de sortie de la rame $(i-1)$ serait la somme du temps de sortie du canton où elle se trouve du temps pour arriver à la station, du temps d'arrêt de la rame $(i-1)$ à la station $K+1$ et du temps mis par la rame pour sortir du canton de la station $K+1$.

Ce deuxième indice est plus pénalisant pour la rame i mais lorsque celle-ci repartira de la station K il est probable que la rame $(i-1)$ sera déjà partie de la station $K+1$.

Par contre si nous ne prenons en compte que le premier indice, la rame i ne s'arrêtera pas en anticollision devant le canton X , mais elle pourrait être bloquée à la limite du canton de la station $K+1$ encore occupé par la rame $(i-1)$.

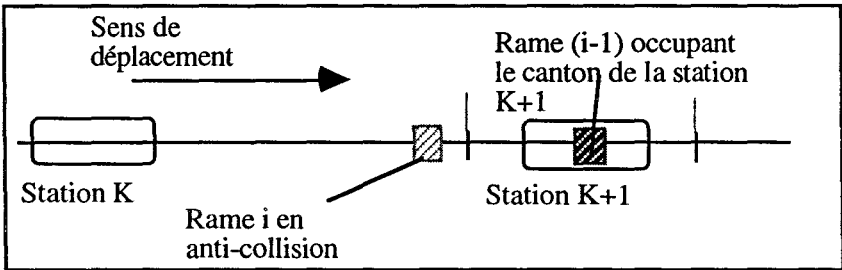


Figure 3.24. Arrêt en anticollision selon le calcul du premier indice de risque

3.6. Extension de la notion de calcul du risque

Le calcul précédent n'intervient que si une rame se trouve perturbée en station ou en inter-station. Nous pouvons le généraliser en faisant intervenir ce calcul à chaque fois qu'une rame i va partir d'une station. Plusieurs cas se présentent :

- La rame (i-1) est perturbée.
- La rame (i-1) est en anticollision suite à une perturbation de la rame (i-2).
- La rame (i-1) sur-stationne à la station K+1 suite à une décision de non-départ (prise après un calcul de risque effectué précédemment).
- La rame (i-1) est en marche normale.

3.6.1. Si la rame (i-1) est perturbée

Nous calculons le temps de sortie du canton X et le temps de parcours de la rame i de la station K à la limite du canton X. Pour calculer le risque, nous comparons les deux sous ensembles flous suivants :

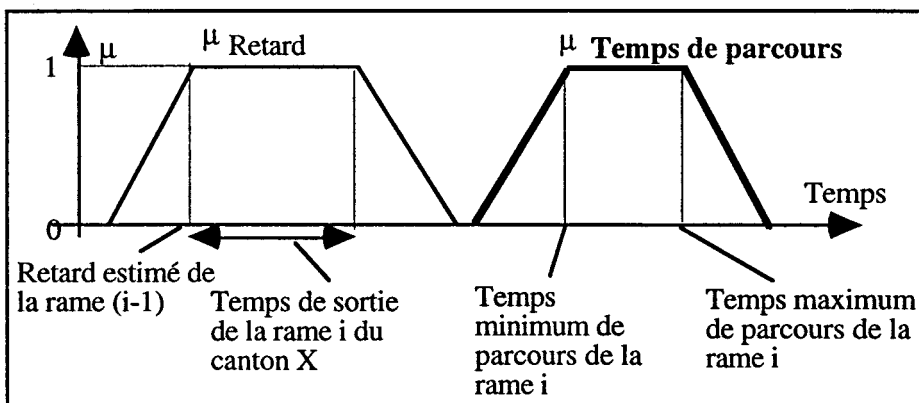


Figure 3.25. Définition des deux sous ensembles flous pour le calcul du risque

3.6.2. Si la rame (i-1) est en marche normale

Nous calculons le temps de sortie du canton X et le temps de parcours de la rame i de la station K à la limite du canton X. Nous comparons les deux sous ensembles flous suivants :

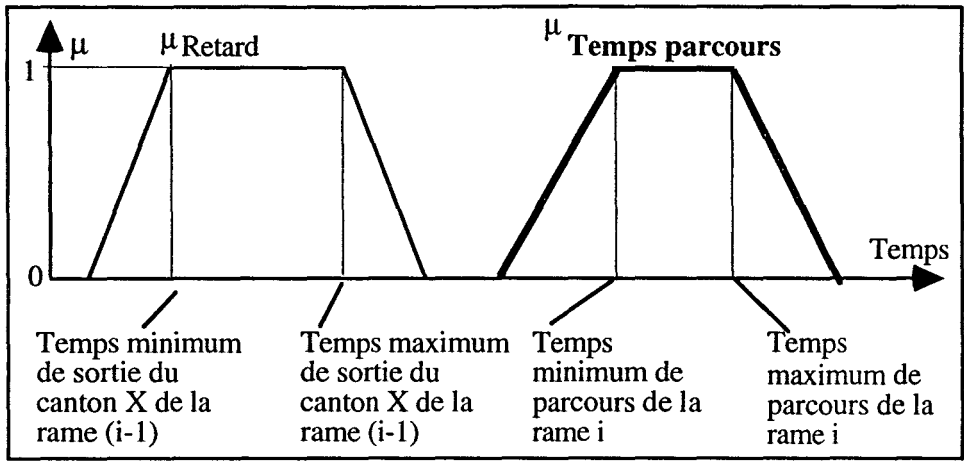


Figure 3.26. Définition des deux sous ensembles flous pour le calcul du risque

3.6.2. Si la rame (i-1) est en anticollision ou en sur-stationnement

Le calcul revient à considérer la rame (i-1) comme perturbée. Le temps de perturbation sera le temps pendant lequel la rame (i-1) est arrêtée suite à une anticollision ou suite à un sur-stationnement.

Remarque : Le calcul du deuxième indice de risque se fait sur la même base. Le temps de parcours sera le temps mis par la rame de la station K à la limite du canton de la station (K+1) et le temps de sortie de la rame (i-1) sera le temps de parcours de sa position actuelle jusqu'à sa sortie du canton de la station (K+1).

3.7. Comparaison entre les deux mesures de distance

Il existe une grande quantité d'algorithmes servant à classer les nombres flous. Dans notre cas, il suffit de comparer deux nombres flous pour dire si le nombre flou représentant le retard estimé est plus grand ou plus petit que le nombre flou représentant le temps de parcours. Il ne suffit pas d'utiliser un algorithme calculant les distances : nous devons savoir comment un nombre flou est placé par rapport à l'autre. Nous avons, d'abord, utilisé un algorithme qui calcule les surfaces de recouvrement (ou de non recouvrement) d'un nombre flou par rapport à l'autre. Ce calcul de surface nous permet de calculer deux indices : le risque et le non-risque. La comparaison de ces deux indices nous permet de conclure sur la grandeur d'un nombre par rapport à l'autre ; à savoir si le risque est supérieur au non-risque, alors le nombre flou représentant le retard estimé est supérieur au nombre flou représentant le temps de parcours.

Un certain nombre de méthodes de comparaison sont décrites et comparées par Bortolan [BOR 85] ainsi que Fortemps [FOR 95], nous avons implémenté les méthodes pouvant conclure l'importance d'un nombre par rapport à un autre, dans tous les cas proposés. Nous avons retenu deux méthodes car certaines des méthodes proposées ne solutionnent pas tous les cas : la méthode de Dubois et Prade [DUB 83] et la méthode proposée par Jain [JAI 76] et [JAI 77]. $\tilde{r}_e$ est le nombre flou représentant le retard estimé et $\tilde{t}_p$ le nombre flou représentant le temps de parcours.

3.7.1. La méthode de Dubois et Prade

Ils proposent la construction de quatre indices capables de décrire la localisation de deux nombres flous $\tilde{r}_e$ et $\tilde{t}_p$. En particulier ils définissent :

1. Possibilité de surclassement (de $\tilde{r}_e$ sur $\tilde{t}_p$)

$$\begin{aligned} PD(\tilde{r}_e) &= \prod_{\tilde{r}_e} ([\tilde{t}_p, +\infty)) \triangleq \text{Poss}(\tilde{r}_e \geq \tilde{t}_p) \\ &= \sup_{z_{r_e}} \min[\mu_{\tilde{r}_e}(z_{r_e}), \sup_{z_{t_p} \leq z_{r_e}} \mu_{\tilde{t}_p}(z_{t_p})] \\ &= \sup_{\substack{z_{r_e}, z_{t_p} \\ z_{r_e} \geq z_{t_p}}} \min[\mu_{\tilde{r}_e}(z_{r_e}), \mu_{\tilde{t}_p}(z_{t_p})] \end{aligned} \quad (2)$$

2. Possibilité stricte de surclassement

$$\begin{aligned} PSD(\tilde{r}_e) &= \prod_{\tilde{r}_e} ([\tilde{t}_p, +\infty)) \triangleq \text{Poss}(\tilde{r}_e > \tilde{t}_p) \\ &= \sup_{z_{r_e}} \inf_{z_{t_p}, z_{t_p} \geq z_{r_e}} \min[\mu_{\tilde{r}_e}(z_{r_e}), 1 - \mu_{\tilde{t}_p}(z_{t_p})] \end{aligned} \quad (3)$$

3. Nécessité de surclassement

$$\begin{aligned} ND(\tilde{r}_e) &= \prod_{\tilde{r}_e} ([\tilde{t}_p, +\infty)) \triangleq \text{Nec}(\tilde{r}_e \geq \tilde{t}_p) \\ &= \inf_{z_{r_e}} \sup_{z_{t_p}, z_{t_p} \leq z_{r_e}} \max[1 - \mu_{\tilde{r}_e}(z_{r_e}), \mu_{\tilde{t}_p}(z_{t_p})] \end{aligned} \quad (4)$$

4. Nécessité stricte de surclassement

$$\begin{aligned} NSD(\tilde{r}_e) &= \prod_{\tilde{r}_e} ([\tilde{t}_p, +\infty)) \triangleq \text{Nec}(\tilde{r}_e > \tilde{t}_p) \\ &= 1 - \sup_{z_{r_e} \leq z_{t_p}} \min[\mu_{\tilde{r}_e}(z_{r_e}), \mu_{\tilde{t}_p}(z_{t_p})] \\ &= 1 - PD(\tilde{t}_p) \end{aligned} \quad (5)$$

La valeur de ces indices permet de classer $\tilde{r}_e$ par rapport à $\tilde{t}_p$. Nous illustrons cet algorithme par quelques exemples illustrés figure 3.27.

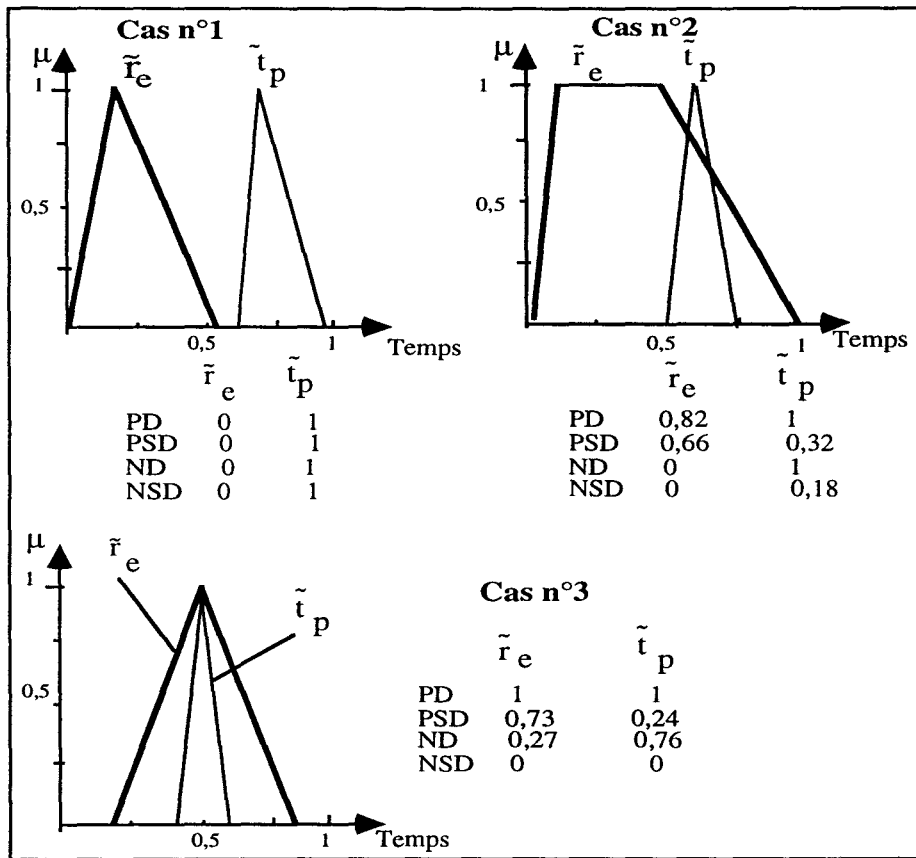


Figure 3.27. Comparaison de deux nombres flous par la méthode de Dubois et Prade

Dans la plupart des cas, les indices donnent la position de $\tilde{r}_e$ par rapport à $\tilde{t}_p$. Mais dans d'autres cas, l'algorithme ne nous permet pas de trancher. En référence aux cas exposés par Bortolan [BOR 85] nous avons choisi comme indice le degré de surclassement strict. Dans le cas n°3 de la figure 3.27, $\tilde{r}_e$ est supérieur à $\tilde{t}_p$ et dans ce cas, seul cet indice permet de trancher. De plus, la relation désirée est la relation strictement supérieure : il faut que le nombre flou représentant le temps de parcours soit strictement supérieur au nombre flou représentant le retard estimé pour que la conclusion ait un sens dans notre application.

3.7.2. La méthode de Jain

Dans cette méthode, nous calculons deux indices $\mu_o(1)$ et $\mu_o(2)$. Si $\mu_o(2)$ est supérieur à $\mu_o(1)$ alors $\tilde{t}_p$ est supérieur à $\tilde{r}_e$. Ces deux indices sont calculés de la façon suivante :

On définit l'ensemble $\tilde{z}$ tel que sa fonction d'appartenance $\mu_{\max}$ soit définie par (6).

$$\mu_{\max} = \left[\frac{z}{z_{\max}} \right]^k, (k > 0) \quad (6)$$

On obtient $\mu_o(1)$ maximum de l'intersection de $\tilde{r}_e$ avec $\tilde{z}$ et $\mu_o(2)$ maximum de l'intersection de $\tilde{t}_p$ avec $\tilde{z}$, c'est à dire $\mu_o(i) = \text{hgt}(\tilde{u}_i \cap \tilde{u}_{\max})$. Nous illustrons la méthode par l'exemple de la figure 3.28.

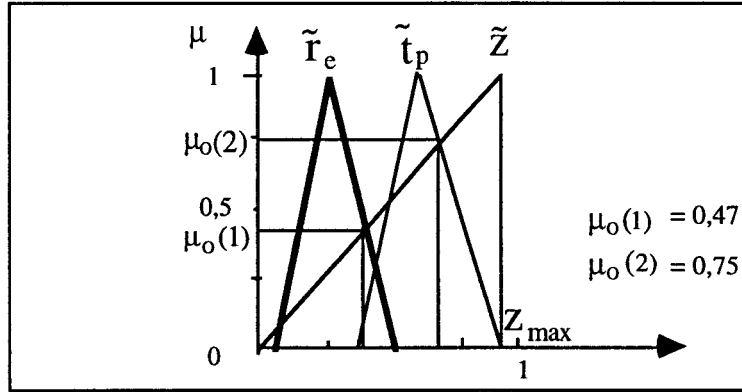


Figure 3.28. Comparaison de deux nombres flous par la méthode de Jain

Dans notre cas, le risque sera $\mu_o(1)$ et le non-risque sera $\mu_o(2)$.

3.7.3. La méthode des surfaces

D'après le calcul des surfaces générées par les ensembles flous nous classons un nombre flou $\tilde{r}_e$ par rapport à un autre $\tilde{t}_p$ en calculant la valeur du risque et du non-risque [BAI 1-96].

Risque = aire entre la limite haute de $\tilde{r}_e$, c'est à dire la droite b4-b3, et la limite basse de $\tilde{t}_p$, la droite a1- a2.

Non-risque = aire entre la limite basse de $\tilde{r}_e$, la droite b1 - b2, et la limite haute de $\tilde{t}_p$, la droite a3 - a4.

Nous obtenons donc si non-risque > risque , $\tilde{t}_p > \tilde{r}_e$

Les figures suivantes représentent ce calcul dans tous les cas possibles.

1. Si $b4 \leq a1$

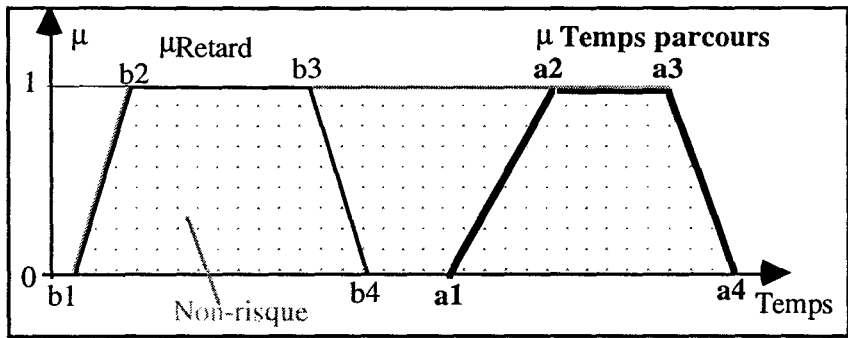


Figure 3.29. Le temps de parcours est supérieur au retard estimé

Risque = 0
Non-risque = Aire de $(b1, b2, a3, a4)$

2. Si $a1 < b4$ et $a2 \geq b3$

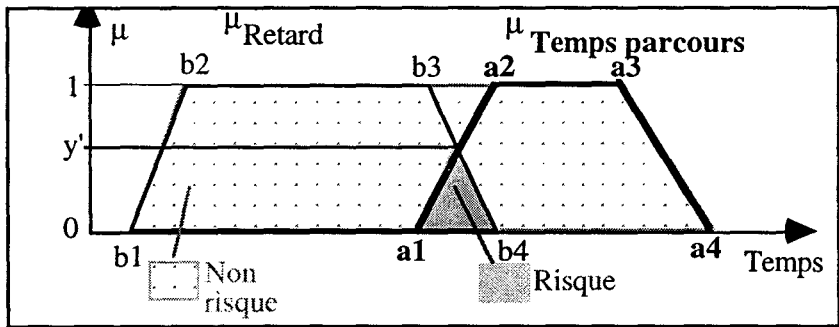


Figure 3.30. Chevauchement des deux sous ensembles flous

Risque = Aire du triangle $(a1, y', b4)$
Non-risque = Aire de $(b1, b2, a3, a4)$

3. Si $a1 \leq b4$ et $a2 \leq b3$ et $a3 \geq b2$

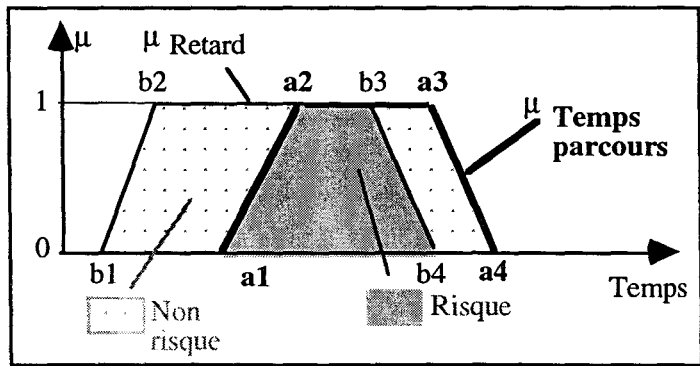


Figure 3.31. Chevauchement des deux sous ensembles flous

Risque = Aire de (a1, a2, b3, a4)

Non-risque = Aire de (b1, b2, a3, a4)

4. Si $a1 < b4$ et $a2 < b3$ et $b1 \leq a4$ et $a3 < b2$

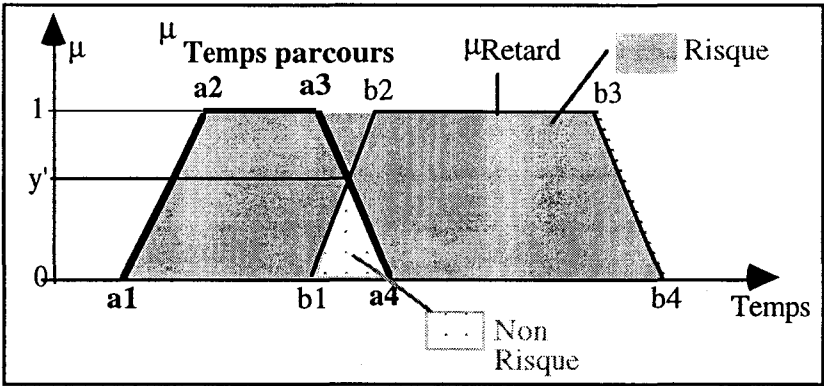


Figure 3.32. Chevauchement des deux sous ensembles flous

Risque = Aire de (a1, a2, b3, a4)

Non-risque = Aire du triangle (a4, y', b1)

5. Si $a4 < b1$

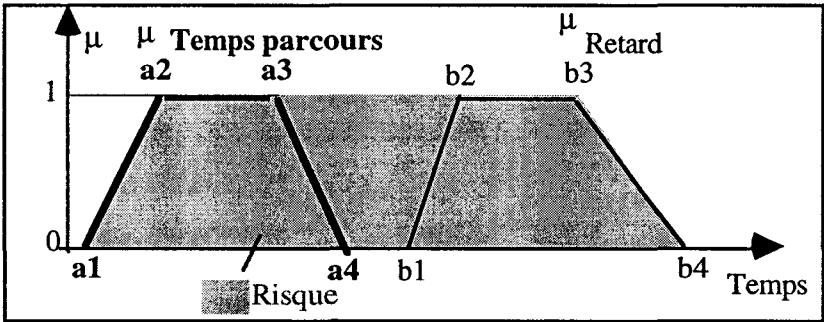


Figure 3.33. Le temps de parcours est inférieur au retard estimé

Risque = Aire de (a1, a2, b3, b4)

Non-risque = 0

Pour déterminer la méthode à employer, nous simulons une panne sur la ligne de métro. Nous supposons le temps de perturbation connu ; de ce fait, les deux sous ensembles flous sont parfaitement définis [BAI 3-96]. Pour utiliser la méthode la plus appropriée, nous mesurons les heures de passage des rames en station proche de la perturbation simulée. Nous calculons, avant chaque départ d'une rame en station, un indice de risque et de non-risque pour décider de laisser ou non partir la rame. Ce calcul de risque se fait par les trois méthodes : le calcul des surfaces, l'algorithme de Dubois et Prade et l'algorithme de Jain avec $k=1$. La

méthode qui respectera le plus le critère de régularité, sans provoquer d'arrêt en anticollision, sera considérée comme la plus efficace dans notre cas.

Nous représentons dans le tableau 3.1 la somme de la différence entre l'intervalle réel des passages des rames aux stations et l'intervalle nominal (150 secondes) pour toutes les stations rencontrées par toutes les rames pour différentes perturbations. Les temps de calcul sont donnés à titre indicatif et dépendent de la machine utilisée.

	Perturbation de 150 s			Perturbation de 200 s			Perturbation de 250 s		
Méthode	Calcul des surfaces	Dubois et Prade	Jain avec k=1	Calcul des surfaces	Dubois et Prade	Jain avec k=1	Calcul des surfaces	Dubois et Prade	Jain avec k=1
Somme des différences entre les intervalles	1253 s	1268 s	1264 s	1642 s	1666 s	1673 s	2523 s	2572 s	2566 s
Somme des temps de calcul (ms)	99	155685	6403	174	336393	12725	324	614704	20883

Tableau 3.1. Sommes des différences d'intervalles de passages des trains en station avec l'intervalle nominal (150 s) pour les trois algorithmes.

Nous pouvons voir sur ce tableau que, pour toutes les perturbations simulées, l'algorithme de calcul des surfaces donne une différence d'intervalle de passages en station qui se rapproche le plus du fonctionnement nominal. De plus, bien que les algorithmes de Dubois et Prade et de Jain ne soient pas optimisés, la différence est suffisamment significative pour dire que l'algorithme de calcul des surfaces est bien plus rapide que les deux autres algorithmes. De ce fait, pour comparer les deux nombres flous, nous utiliserons cet algorithme de calcul de surface pour calculer le risque et le non-risque.

3.8. Application de l'algorithme au service partiel.

Le traitement du service partiel dépend des hypothèses fixées lors de sa définition.

Comme nous l'avons vu au chapitre 1, le service partiel le plus adapté sera une boucle du service partiel imbriquée dans une boucle du service ligne. Ainsi, un passager, prenant une

rame à l'extérieur de la boucle service partiel, doit pouvoir terminer son trajet sans changement de rame quelle que soit sa destination. Ce qui veut dire qu'une rame assure le trajet sur l'ensemble de la ligne sans faire demi-tour au terminus partiel. Le problème se résume donc à un problème de synchronisation : assurer, sur la boucle du service partiel, une série de rames présentant toujours le même ordre entre rame "service partiel" et rame "service ligne". Ainsi, au terminus service partiel les rames "service partiel" vont rebrousser chemin et les rames "service ligne" vont continuer (figure 3.34).

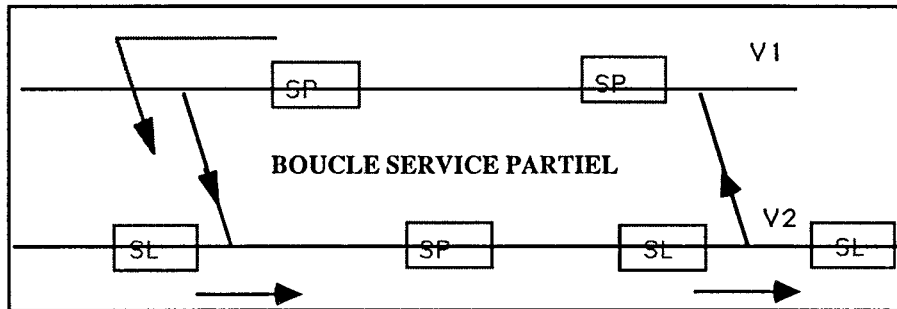


Figure 3.34. Série de rames à l'intérieur de la boucle service partiel

Le problème de synchronisation se situe à l'entrée de la boucle du service partiel. Ceci revient à considérer le système comme une fourche où, systématiquement, le passage d'une rame service ligne est suivie d'une rame service partiel avec une fréquence à l'intérieur de la boucle service partiel deux fois supérieure à celle de l'ensemble de la ligne.

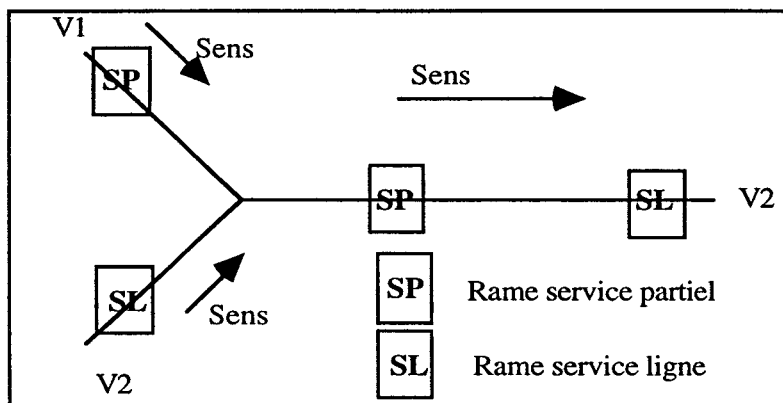


Figure 3.35. Autre représentation de l'entrée dans la boucle service partiel

Dans le cas de la figure 3.35, la rame service ligne doit passer avant la rame service partiel. Il est donc impératif d'empêcher la rame service partiel de rebrousser chemin car ceci entraînerait un déséquilibre au niveau de la succession des rames. Nous devons immobiliser cette rame service partiel jusqu'à ce que la rame service ligne SL passe dans la boucle service

partiel. Pour ce faire, nous pouvons utiliser l'algorithme de calcul de risque décrit précédemment. Si on considère une rame fictive située devant la rame service partiel SP1, sur le premier canton après l'aiguillage et si on considère que cette rame fictive est perturbée d'un temps t_0 , cette rame fictive ne sera visible que pour la rame SP1. La rame aval de la rame service ligne SL1 sera la rame réelle sur la boucle service partiel, notée SP2 sur la figure ci-dessous.

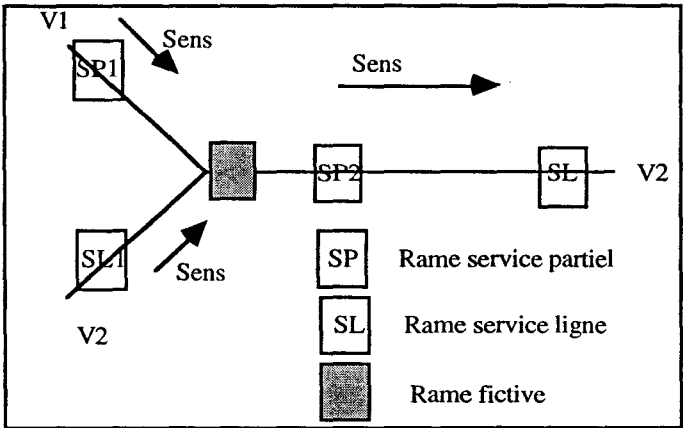


Figure 3.36. Introduction d'une rame fictive pour la rame service partiel

On considérera le temps t_0 de perturbation de la rame fictive comme le temps nécessaire à la rame service ligne SL1 pour arriver dans la boucle service partiel ou plus précisément, le temps nécessaire pour aller de sa position actuelle jusqu'à la fin du canton X2. Si la rame service ligne est perturbée, nous ajouterons au temps t_0 le temps de perturbation estimé. De cette façon, on empêchera, par un calcul de risque et de non-risque, la rame service partiel SP1 de rebrousser chemin et de prendre la place de la rame service ligne (figure 3.37).

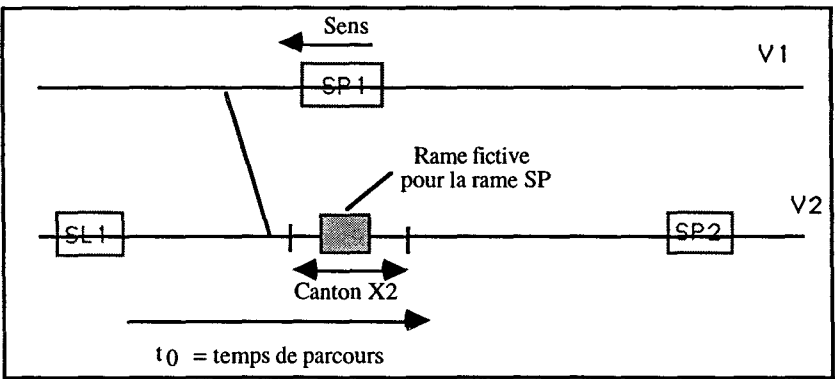


Figure 3.37. Définition du temps de perturbation de la rame fictive

Si la rame service ligne SL1 est arrêtée à la station K, lors de son départ, on calculera un indice de risque, pour éviter que cette rame ne s'arrête en anticollision si la rame précédente

est perturbée. Ce calcul se fera en comparant le temps de parcours de cette rame jusqu'au canton occupé par la rame aval, rame notée SP2, et le temps de perturbation éventuel de la rame SP2. Avant de faire partir la rame service partiel SP1 du terminus service partiel, un autre indice de risque sera calculé en comparant le temps de parcours de cette rame jusqu'à la limite du canton occupé par la rame fictive et le temps de perturbation t_0 de la rame fictive.

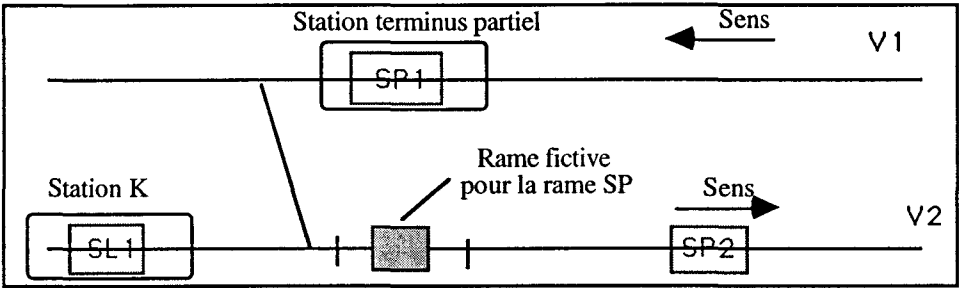


Figure 3.38. Calculs de risque pour les rames SL et SP

Aussitôt que la rame service ligne SL1 a libéré le canton X2, nous faisons disparaître cette rame fictive. De ce fait, la rame SP1 attendant au terminus partiel partira de la station où elle sur-stationnait.

Dès la disparition de la rame fictive qui empêchait la rame service partiel de partir, une autre rame fictive sera placée au même endroit que précédemment. De ce fait, la prochaine rame service ligne SL2 aura pour rame aval la rame fictive qui ne sera pas visible par la rame service partiel SP1.

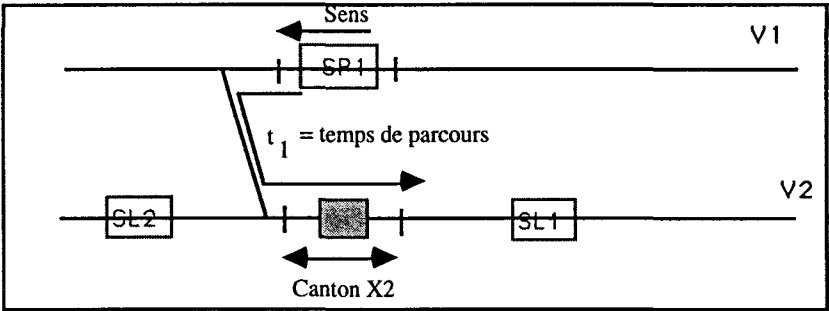


Figure 3.39. Blocage de la rame service ligne jusqu'au passage de la rame service partiel

La rame fictive sera considérée comme perturbée durant un temps t_1 , temps de parcours de la rame service partiel SP1 de sa position jusqu'à la fin du canton X2. De la même façon, un calcul de risque et de non-risque sera réalisé sur la rame service ligne SL2 et la rame fictive pour éviter que la rame service ligne s'intercale dans la boucle avant la rame service partiel normalement attendue.

De cette façon, on peut assurer une synchronisation entre les rames du service ligne et du service partiel même lorsque les rames sont perturbées. Cette opération s'effectue d'une manière identique sans service partiel, lors de fortes perturbations, quand nous bloquons, successivement, en station, les rames service ligne et service partiel le long de la ligne en évitant les arrêts en anticollision.

Néanmoins, lorsqu'une rame est perturbée sur la boucle service ligne durant un temps très important, nous serons amenés à supprimer la rame fictive qui empêche la rame service partiel de rebrousser chemin. Ceci permet à une partie de la ligne de fonctionner lorsqu'un incident important bloque l'autre partie. Bien entendu, la succession des rames sur la boucle service partiel sera interrompue mais les passagers se trouvant dans les zones non concernées par la perturbation pourront continuer à circuler normalement.

4. Conclusion

Dans un premier temps, l'algorithme basé sur la commande floue nous a permis de réguler l'ensemble des rames par rapport à un horaire théorique de référence. Les retards sont mesurés d'abord à chaque pas d'échantillonnage. Cependant, dans un souci de faire correspondre notre modèle aux régulations en service sur les différents métro automatiques, nous avons modifié notre modèle en tenant compte des retards mesurés à chaque arrêt d'une rame en station. Ce modèle de régulation répartit l'intervalle entre les rames et permet aux rames retardées de revenir à leur fonctionnement nominal le plus rapidement possible. Ce modèle est très efficace pour les perturbations de faible ampleur mais, parce qu'il ne prend pas en compte les événements susceptibles de se produire lors du fonctionnement, il n'interdit pas les arrêts prolongés provoqués par les algorithmes de sécurité. De ce fait, lors d'une perturbation importante, la rame située en amont de la perturbation est, certes retardée mais rien ne l'empêche de repartir de la station où elle se trouve. Si la perturbation se prolonge, et si elle est située sur l'inter-station de la rame venant de partir de la station, il y aura arrêt de cette rame non perturbée en anticollision. Cet arrêt se prolongera aussi longtemps que la rame située en aval restera bloquée. Il est préférable, pour la qualité de service, de bloquer cette rame avant son départ de la station pendant un temps dépendant du temps estimé de la durée de la panne et de la dynamique de la rame non perturbée. De ce fait, lors de grosses perturbations, très peu de rames se trouveront entre deux stations : cela permet d'arrêter si nécessaire, d'une façon plus souple pour la clientèle, l'ensemble du trafic sur la ligne. De plus, ce calcul de risque et de non-risque permet de synchroniser d'une façon efficace l'ensemble des rames lors du fonctionnement de la ligne en service partiel. Il permet de supprimer les conflits dus à la nature même du fonctionnement d'une ligne de métro avec un service partiel.

CHAPITRE 4

PLATEFORME DE SIMULATION ET RESULTATS

1 . Introduction

Pour tester le bien fondé de l'application, il est nécessaire de posséder un outil de validation de l'algorithme. Le meilleur outil disponible serait le site réel du métro de Lille. Il faudrait avoir à notre disposition la Ligne 1 du VAL de Lille comprenant le matériel roulant, les voies pour faire circuler les rames et tout le matériel informatique au poste de commande centralisé. Le coût d'une telle validation serait très élevé. La seule solution nous permettant de valider efficacement notre modèle est de posséder une simulation d'une ligne de métro automatisée. La simulation doit être la plus réaliste possible afin de comparer les différents systèmes de régulation. De ce fait, la simulation permettra de simuler les différentes lignes de métro automatisées et ne sera pas réservée à une seule ligne de métro. De plus, la simulation nous permet de mettre au point notre algorithme - chose difficilement concevable si nous utilisons le site réel. Ainsi, l'INRETS s'est doté de moyens suffisants en vue d'une simulation la plus réaliste possible [ROD 1-94].

Les applications de la simulation dans les transports guidés sont nombreuses. Nous retiendrons les suivantes :

- Modélisation de chaînes cinématiques,
- Conception de nouveaux principes d'espacement entre véhicules (cantons fixes, cantons mobiles déformables),
- Mise en oeuvre de plans de circulation,
- Aide à la décision pour le choix de la structure d'un réseau (terminus à une, deux ou trois voies, lignes à voies uniques),
- Aide à l'exploitation (extrapolation d'une situation présente et étude de la répercussion de décisions de régulation),

- Formation d'opérateurs ("simulateur de formation", reconstitution de cas difficiles, rares ou dangereux),
- Dimensionnement des installations fixes de traction électrique.

La simulation consiste à modéliser le système à étudier de telle façon que le comportement du modèle reproduise le plus fidèlement possible celui du système réel. Pour un même système, il existe plusieurs modèles possibles suivant les paramètres que l'on souhaite étudier. Lorsque l'on modélise un système, on est donc toujours amené à faire des approximations qui dépendent des objectifs recherchés dans la simulation. Le but de cette plateforme de simulation réalisée en langage orienté objet est, entre autres, de tester les nouveaux modèles de régulation et de les comparer à ceux existant actuellement. Cette plateforme, ainsi que la notion de programmation par objet, est décrite dans l'annexe 3.

2. Simulation de la ligne 1 du VAL de Lille

Nous nous servons de la simulation décrite en annexe 3 pour évaluer notre nouvelle régulation de trafic et la comparer à celle existant. Pour cela il faut que la ligne simulée soit une représentation réaliste de la ligne.

Un éditeur graphique réalisé à l'INRETS-ESTAS [GOG 94] nous permet de réaliser et de représenter les caractéristiques de la ligne réelle (dimension de la ligne, dimension et position des cantons, dimension des rames, dimension et position des stations, vitesses limites admissibles sur chaque portion de voie ...).

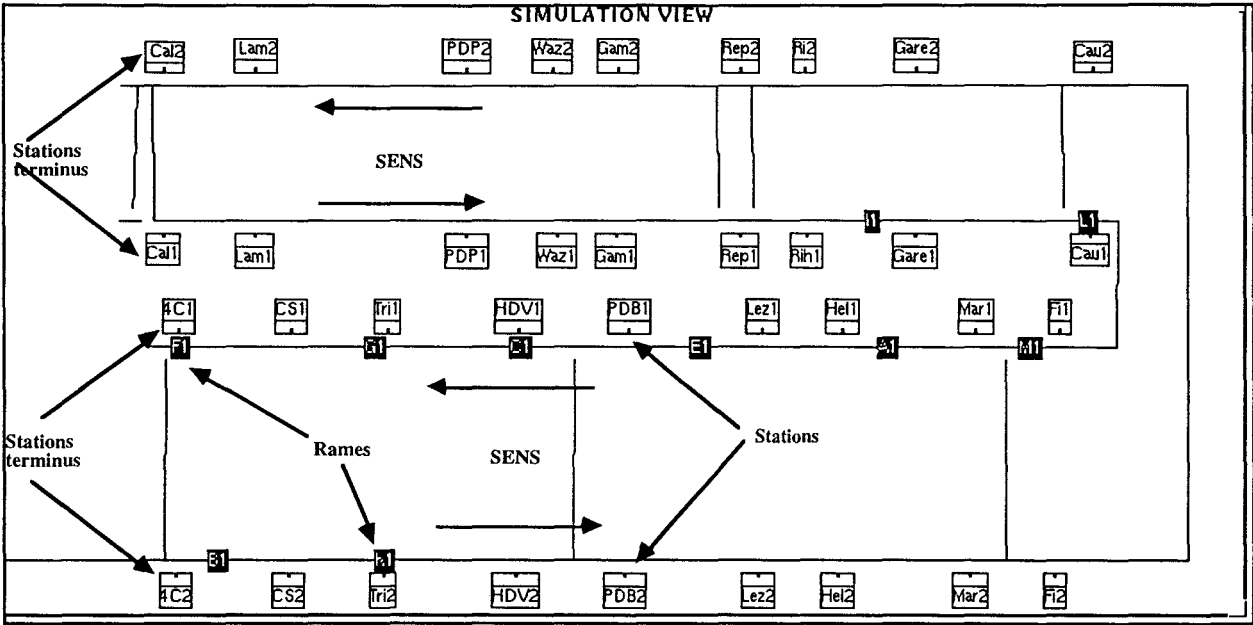


Figure 4.1. Simulation de la ligne 1 du VAL de Lille

Dans un premier temps, nous avons implémenté l'algorithme de simulation décrit dans le chapitre 3, pour la simulation de la ligne 1 du VAL de Lille [BAI 2-95].

2.1. Simulation d'une perturbation

Nous avons simulé une perturbation de 72 secondes sur une rame située au milieu de la ligne, la rame A1 sur la figure 4.1. Les 72 secondes correspondent au temps maximum de perturbation possible sur la rame pour que cette perturbation n'en provoque pas d'autres sur les rames aval, suite à des arrêts en anticollision. Nous pouvons voir sur la figure ci-dessous le retard de la rame perturbée ainsi que le retard de la rame la précédant (E1). La figure 4.3 nous montre le retard des rames M1 et L1 situées derrière la rame perturbée.

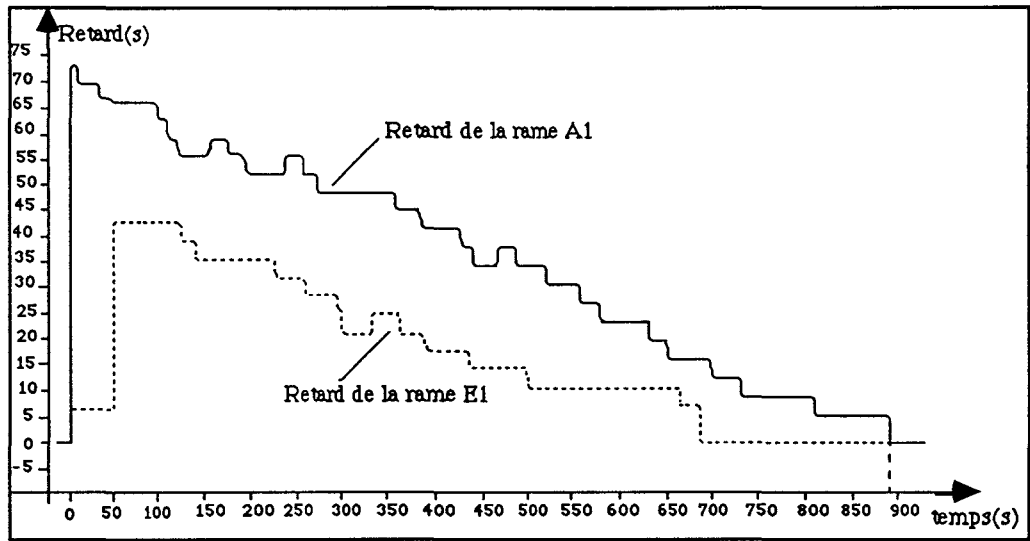


Figure 4.2. Retard des rames E1 et A1 en fonction du temps

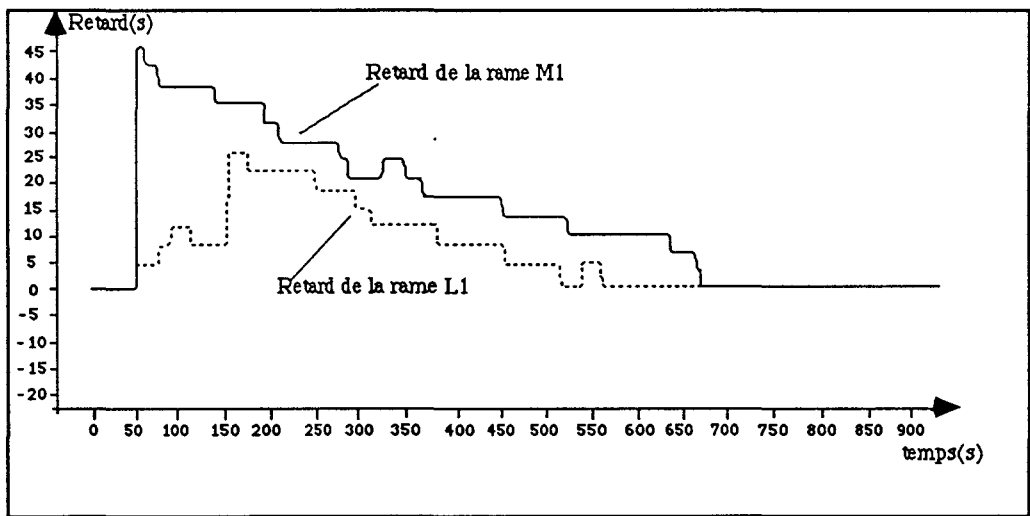


Figure 4.3. Retard des rames M1 et L1 en fonction du temps

On constate que la rame A1 rattrape son retard en 891 secondes. On notera une répartition plus grande de l'intervalle entre les rames, puisqu'on provoque un retard de 42 secondes sur la rame E1. De ce fait, l'intervalle entre la rame E1 et A1 devient égal à 150 secondes (intervalle nominal $-42 + 72$) ce qui se rapproche le plus fortement de l'intervalle nominal, ici 120 secondes.

2.2. Comparaison avec l'algorithme VAL

Nous avons perturbé la rame A1 de 72 secondes au même endroit que dans les essais précédents. Les autres rames ne prenant aucun retard du fait de l'algorithme de régulation du VAL, nous ne visualisons que le retard de la rame A1.

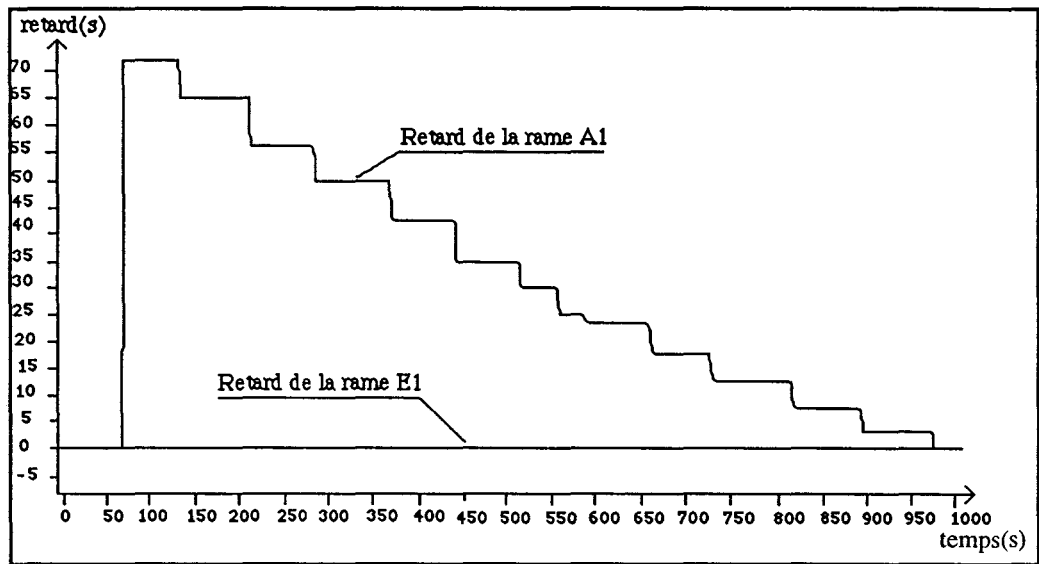


Figure 4.4. Evolution du retard des rames A1 et E1 en fonction du temps pour la régulation type VAL.

Nous pouvons voir que la rame A1 met 913 secondes pour revenir à son parcours nominal. Nous pouvons dire que la rame A1 met le même temps pour récupérer son retard dans la régulation floue que dans la régulation type VAL. Cependant, le fait de provoquer un retard sur les rames amont et aval dans l'algorithme flou nous permet de répartir beaucoup mieux l'intervalle entre les rames successives. Lors des essais, la rame M1 située en aval de A1 s'arrête à la station Marbrerie 1 pour répartir l'intervalle. Si l'intervalle nominal entre les rames était de 60 secondes (et non pas de 120 secondes comme dans les essais de simulation), la rame M1 s'arrêterait en interstation, dans le cas de l'algorithme du type VAL. Alors que dans l'algorithme flou, la rame ne s'arrêterait pas en anticollision car la rame M1 aurait attendu à la station Marbrerie 1. La rame M1 s'est arrêtée en station au lieu de repartir, provoquant un retard sur sa marche normale, l'intervalle séparant la rame A1 et M1 devient égal à 40 secondes, ce qui permet à la rame de circuler normalement. Tandis que dans l'algorithme VAL, la rame M1 s'arrête obligatoirement en anticollision puisqu'elle n'a pas pris de retard à cette station et que la rame A1 a un retard de 72 secondes alors que l'intervalle nominal entre les rames est de 60 secondes.

Pour illustrer la répartition de l'intervalle entre les rames, il suffit de mesurer l'intervalle de temps entre deux passages pour chaque station.

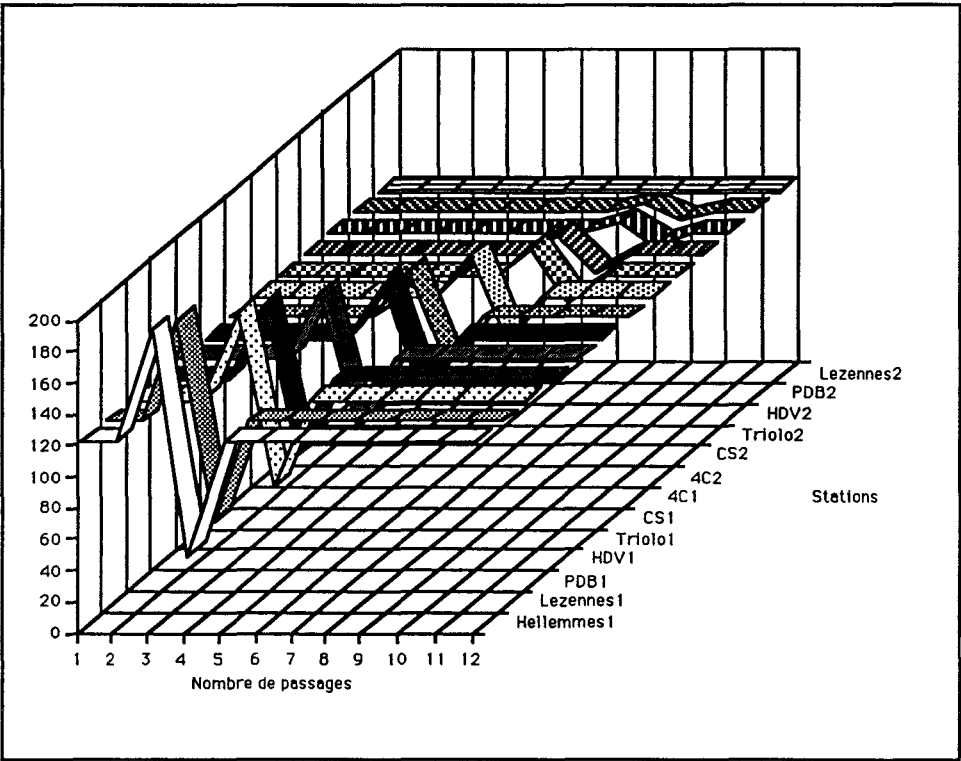


Figure 4.5. Fréquence de passage des rames aux stations successives pour l'algorithme de régulation type VAL

Sur ces graphiques, plus les pics et les vallées sont importants, moins la régulation par intervalle est réalisée. On peut constater, pour l'algorithme type VAL, que les rames perturbées sont moins nombreuses mais, par contre, les pics et les vallées sont plus importants.

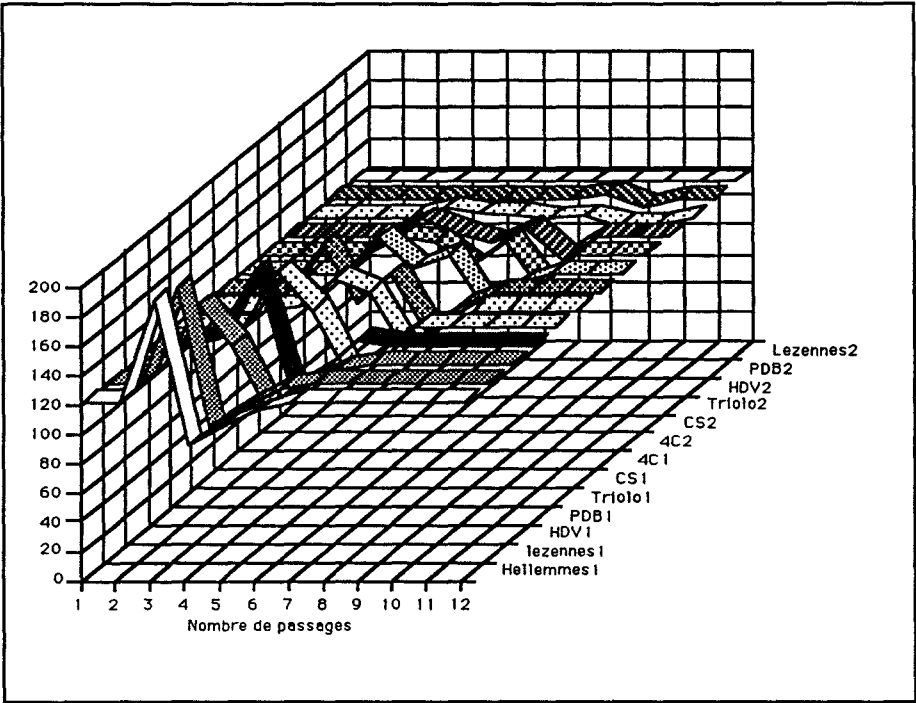


Figure 4.6. Fréquence de passages des rames aux stations successives pour l'algorithme de régulation floue

3. Simulation de la ligne 2 du VAL de Lille

Comme nous l'avons signalé au chapitre 3, dans un souci de faire correspondre la nouvelle régulation avec la régulation du VAL de Lille, nous avons décidé de ne mesurer les retards des rames qu'en station. De ce fait, nous utilisons la simulation de la ligne 2 du VAL de Lille actuellement en construction [ROD 2-94]. Notre régulation a été implémentée sur cette simulation comme la régulation actuelle du VAL de Lille.

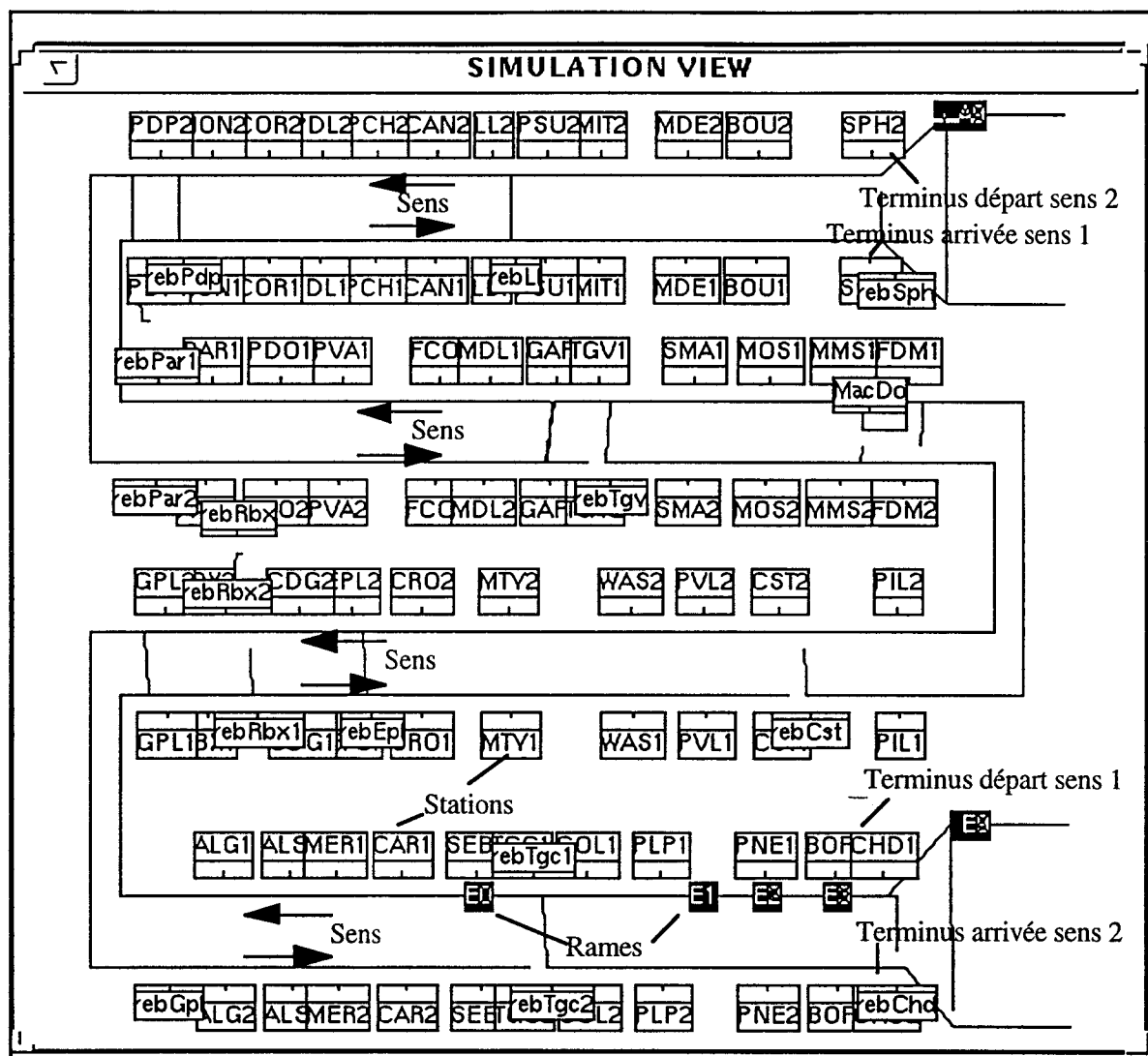


Figure 4.7. Schéma de la simulation de la ligne 2 du VAL

3.1. Simulation d'une perturbation

Une perturbation a été provoquée entre deux stations sur la rame E1 de 100 secondes (figure 4.7). Nous visualisons l'évolution des retards des rames dans le voisinage de la perturbation en fonction du temps.

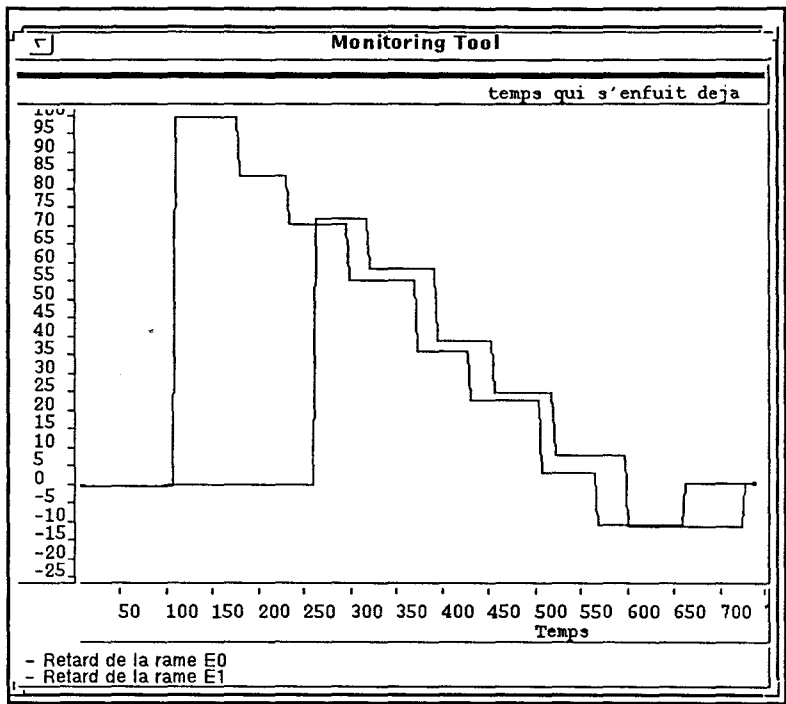


Figure 4.8. Evolution du retard des rames E1 et E0 en fonction du temps

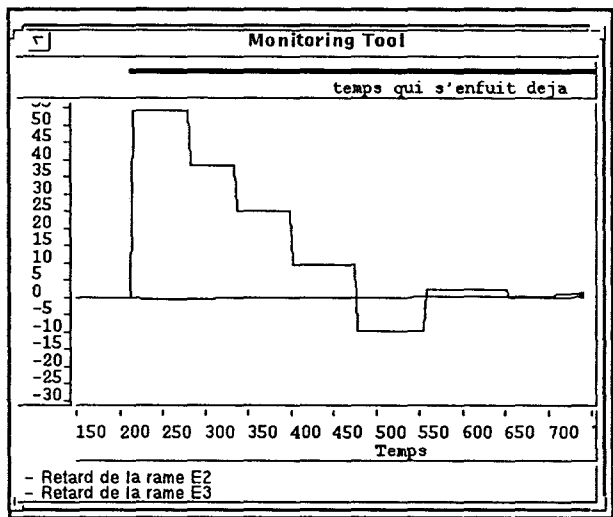


Figure 4.9. Evolution du retard des rames E2 et E3 en fonction du temps

Comparaison avec l'algorithme VAL

Pour faire une comparaison objective de notre nouvelle façon de réguler le trafic, il serait nécessaire de comparer ces résultats avec la régulation de type VAL puisque ces résultats ont été obtenus avec une simulation de la ligne 2 du VAL de Lille. Le retard est mesuré exclusivement lorsque la rame arrive en station ; si la rame est en retard, nous diminuons le temps d'arrêt au maximum de 3 secondes et nous donnons une consigne de vitesse supérieure.

Nous avons perturbé la rame E1 de 100 secondes au même endroit que pour les essais précédents. Les autres rames ne prenant aucun retard, nous ne visualisons que le retard de la rame E1.

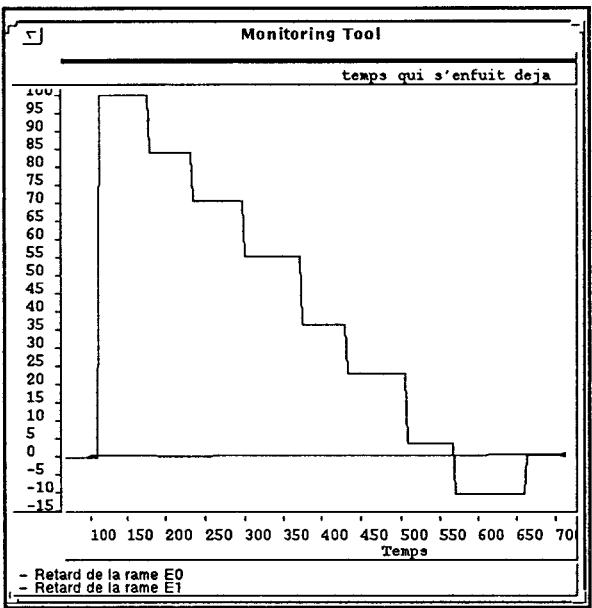


Figure 4.10. Evolution du retard de la rame A1 et E1 en fonction du temps pour la régulation type VAL

L'allure du retard de la rame E1 par rapport au temps est la même pour la régulation floue que pour la régulation type VAL. Nous pouvons dire que la rame E1 met le même temps pour récupérer son retard dans la régulation floue que dans la régulation type VAL (550 secondes pour revenir à son parcours nominal). Cependant, le fait de provoquer un retard sur les rames amont et aval dans l'algorithme flou nous permet de répartir bien mieux l'intervalle entre les rames successives.

3.2. Simulation de deux perturbations

Nous simulons deux perturbations de 100 secondes survenant sur la ligne. Ces perturbations sont provoquées sur deux rames E1 et E5 au même endroit que précédemment (entre les stations PNE1 et PLP1). Nous pouvons voir sur la figure 4.11 l'évolution des retards des rames en fonction du temps.

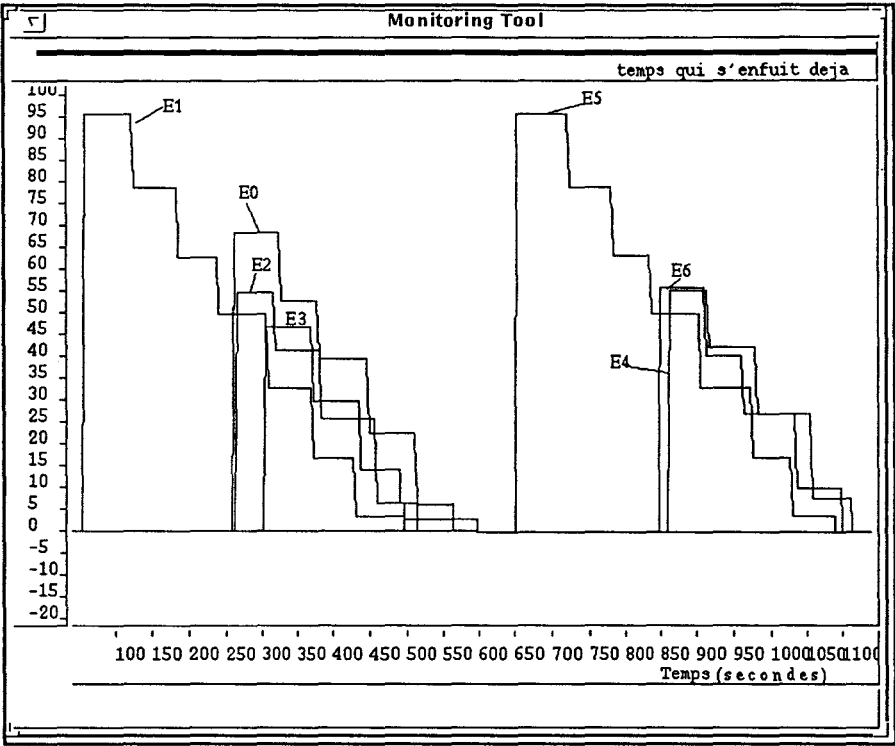


Figure 4.11. Evolution des retards en fonction du temps lorsque deux rames sont perturbées

Comparaison avec l'algorithme VAL

Nous perturbons, de la même façon, les deux rames E1 et E5 d'une durée de 100 secondes et nous visualisons les retards des rames en fonction du temps sur la figure ci-dessous :

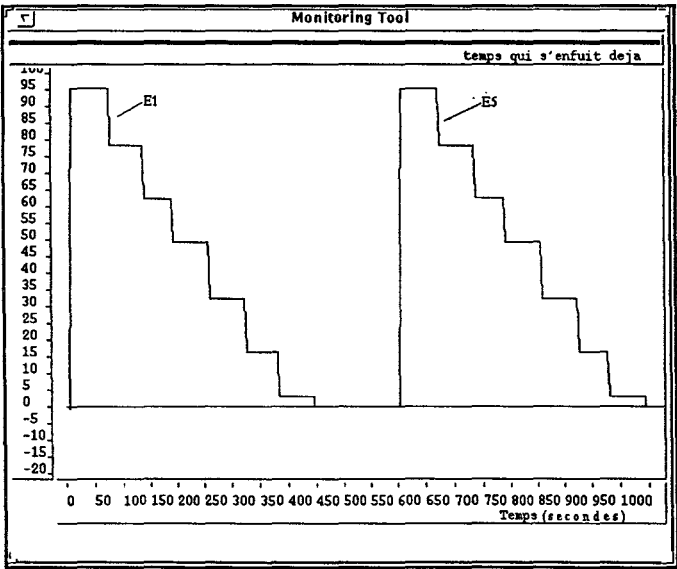


Figure 4.12. Evolution des retards en fonction du temps lorsque deux rames sont perturbées

De la même façon que pour une perturbation, nous pouvons voir que les rames perturbées mettent le même temps pour revenir à leur parcours nominal. Mais, parce que nous perturbons les rames amont et aval, l'intervalle entre les rames est mieux respecté.

Dans les deux lignes de métro simulées, nous pouvons constater une meilleure régularité du carrousel lorsque nous utilisons l'algorithme basé sur la commande floue. Le choix d'une stratégie de mesure des retards soit en station soit à chaque pas d'échantillonnage est un problème délicat. En effet, l'installation d'une méthode de mesure sur la ligne permettant de dialoguer avec les rames, à chaque pas d'échantillonnage, peut poser des problèmes matériels importants. Par contre, mesurer le retard à chaque station et estimer le retard à chaque pas d'échantillonnage permet de diminuer la contrainte matérielle. D'un autre côté, le fait de mesurer le retard à chaque pas d'échantillonnage permet de savoir pratiquement à tout instant si une rame est perturbée et d'appliquer aussitôt les procédures pour éviter l'empilement derrière la rame perturbée afin d'assurer un dépannage rapide de cette rame. Si un système de transport futur permet un tel dialogue avec les rames en circulation, il sera alors possible de se servir de cette fonction pour permettre une mesure des retards et l'envoi des consignes sur les rames à chaque instant.

Nous nous proposons maintenant de tester l'algorithme de supervision flou décrit au chapitre 3.

4. Simulation de la régulation avec indice de risque

Pour mesurer l'efficacité de notre algorithme, nous perturbons une rame sur la ligne suffisamment longtemps pour qu'elle modifie la marche normale des rames situées derrière elle. Pour ce faire, nous perturbons la rame de 200 secondes sur la ligne 2 du VAL de Lille. Afin que les résultats de simulation de l'algorithme ne dépendent pas exclusivement de l'estimation du temps de retard, nous prendrons le temps exact de perturbation simulé et non le temps de retard estimé. Pour être cohérent avec le système actuel, ce temps de perturbation ne sera connu par le système que lorsqu'une rame ne sortira pas du canton à l'instant où elle devrait en sortir.

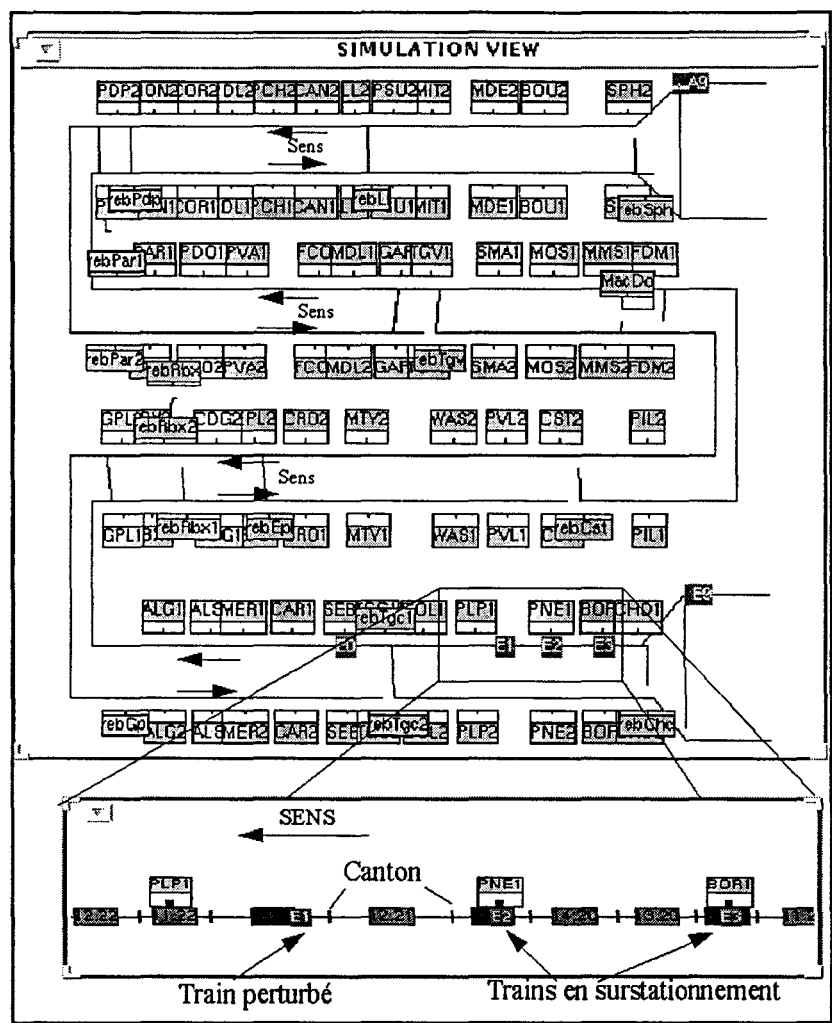


Figure 4.13. Simulation de la ligne 2 du VAL de Lille avec arrêt des rames E2 et E3 en station suite à une perturbation sur la rame E1

Trois algorithmes de régulation sont implémentés :

- La régulation type VAL,
- Une régulation floue sans indice de risque,
- La régulation floue avec calcul de risque à chaque départ de station d'une rame. Les différents indices de risque décrits au chapitre 3 ont été calculés, nous ne retiendrons que l'indice le plus pénalisant : celui qui maintient le plus longtemps la rame arrêtée en station.

La figure 4.13 illustre le cas où les rames situées en aval de la perturbation sont bloquées en station après un calcul de risque et de non-risque. La rame E2 sur-stationne à la station PNE1 pour éviter de s'arrêter en anticollision derrière la rame E1. De ce fait, la rame E3

est bloquée à la station BOR1 pour éviter d'entrer en conflit avec la rame E2. Le sur-stationnement de E3 a été provoqué par le sur-stationnement de E2.

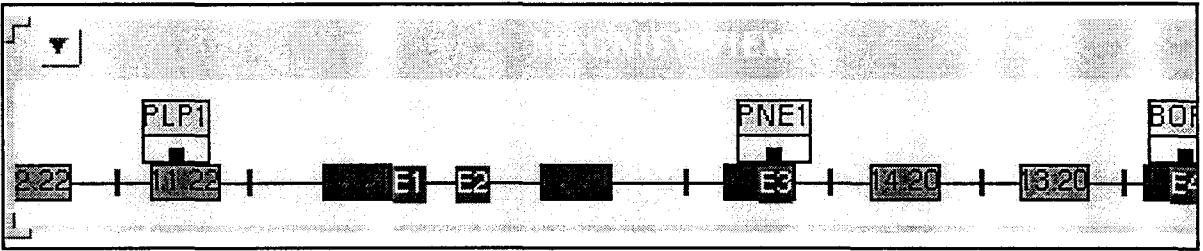


Figure 4.14. Simulation de la ligne 2 sans calcul de risque après une perturbation sur la rame E1

La figure 4.14 illustre le cas où il n'y a pas de calcul de l'indice de risque avant le départ en station. De ce fait, la rame E2 ne sur-stationne pas à la station PNE1. Parce que la perturbation sur la rame E1 se prolonge, elle s'arrête en anticollision sur le canton 12.21, à la limite du canton occupé par la rame perturbée (canton 13.21). Si la perturbation se prolonge, la rame E4 partira de la station BOR1 et s'arrêtera en anticollision sur le canton 14.20 derrière la rame E2. De la même façon, la rame E3 pourra partir de la station PNE1 et de ce fait elle s'arrêtera en anticollision derrière la rame E2.

Nous pouvons regrouper le nombre d'arrêts en anticollision constatés lors d'une perturbation de 200 secondes sur la rame E1 ainsi que leur durée, comme l'illustrent les figures 4.13 et 4.14, pour les trois modes de régulation utilisés, dans le tableau ci-dessous.

Remarque : Cette simulation a été réalisée avec un intervalle entre les rames de 150 secondes.

Type de régulation	Nombre d'arrêts en anticollision	Durée totale des anticollisions
VAL	2	90,3 secondes
Régulation floue sans calcul de risque	2	88,8 secondes
Régulation floue avec calcul de risque	Aucun	0

Pour mesurer l'efficacité du calcul de l'indice de risque, nous diminuons l'intervalle entre les rames. Les cantons ont été calculés pour fonctionner avec un intervalle nominal entre les rames de 60 secondes. Si nous diminuons cet intervalle, des arrêts en anticollision se

produisent. Pour l'algorithme type VAL, plus l'intervalle diminue plus il y a d'arrêts en anticollision. Par contre, grâce au calcul du risque, nous diminuons les arrêts en anticollision en bloquant les rames en station. Ceci permet de limiter les effets d'une perturbation sur le confort des passagers

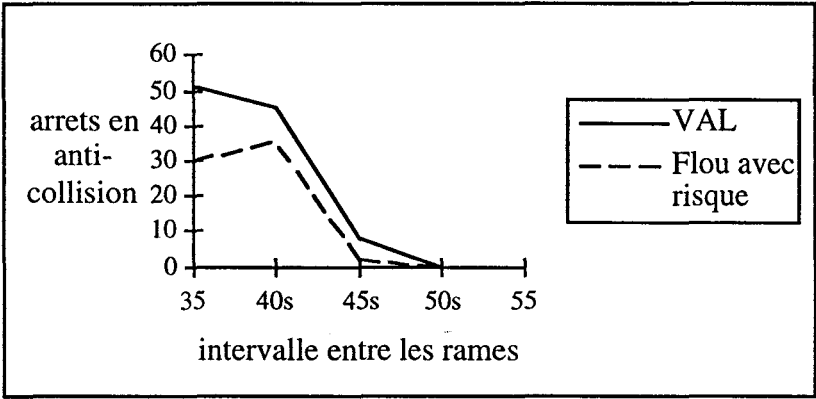


Figure 4.15. Nombre d'arrêts en anticollision en fonction de l'intervalle entre les rames

Nous concluons que l'algorithme du calcul du risque empêche les rames situées en amont de la perturbation de partir de la station où elles se trouvent. Cela évite le déclenchement de la fonction d'anticollision et procure un meilleur confort aux usagers. Le cas présenté figure 4.13 montre l'intérêt de stopper les rames en station si la perturbation se prolonge : nous pouvons facilement arriver à une situation où une rame peut répercuter cette perturbation sur l'ensemble des rames situées derrière elle. Cela est provoqué par l'algorithme de sécurité de la fonction d'anticollision. Un nombre important de passagers se trouveront alors dans une rame bloquée entre deux stations, ce qui est peu appréciable pour leur confort. De plus, le retour au fonctionnement nominal, à la fin de la perturbation sur la rame E1, se fera beaucoup plus difficilement si plusieurs rames se trouvent arrêtées entre deux stations.

Néanmoins, ces résultats sont obtenus lorsque la durée de la panne sur la rame perturbée est connue. Or, par définition, lors du fonctionnement normal, le temps de panne n'est pas connu. Il l'est à la fin de la perturbation lorsqu'on constate le retard provoqué sur la rame. De ce fait, il est nécessaire d'utiliser une estimation de ce temps de perturbation.

Application au service partiel

Pour un bon fonctionnement nominal de la ligne avec service partiel, les horaires de départs aux différents terminus doivent être calculés pour permettre le fonctionnement de la ligne avec un service partiel. Par exemple, si une rame service ligne prend un départ au terminus partiel et, si l'on veut assurer un intervalle entre les rames de 1 minute 15 secondes à

l'intérieur de la boucle partielle, il est nécessaire de construire une table horaire cohérente. Ainsi, une rame service partiel assurera le départ suivant 1 mn 15 secondes après le départ de la rame service ligne.

Cependant, lors de perturbations sur la ligne, il doit toujours y avoir, à l'intérieur du service partiel, une rame dite service partiel précédée et suivie par une rame service ligne. Comme nous l'avons vu au chapitre précédent, nous utilisons l'algorithme de calcul de risque pour synchroniser les rames afin de permettre l'insertion de rames dites service partiel entre les rames dites service ligne lors de perturbations survenant sur la ligne.

Nous avons simulé le fonctionnement de la ligne 2 du VAL de Lille en service partiel avec une table horaire cohérente. Le rebroussement d'une rame sur deux au terminus partiel s'est réalisé naturellement sans perturbation. La rame service partiel est encadrée par des rames service ligne.

Nous perturbons une rame service ligne juste avant son entrée dans la boucle service partiel d'une durée supérieure à l'intervalle entre les rames, ici 200 secondes (l'intervalle entre les rames dans la boucle service partiel est égal à 75 secondes et dans la boucle service ligne à 150 secondes). Sur la figure 4.16 nous voyons que la première rame service partiel devant rebrousser chemin (la rame A3) doit passer après la rame E1 pour maintenir, sur la boucle service partiel, l'alternance des rames service partiel avec les rames service ligne.

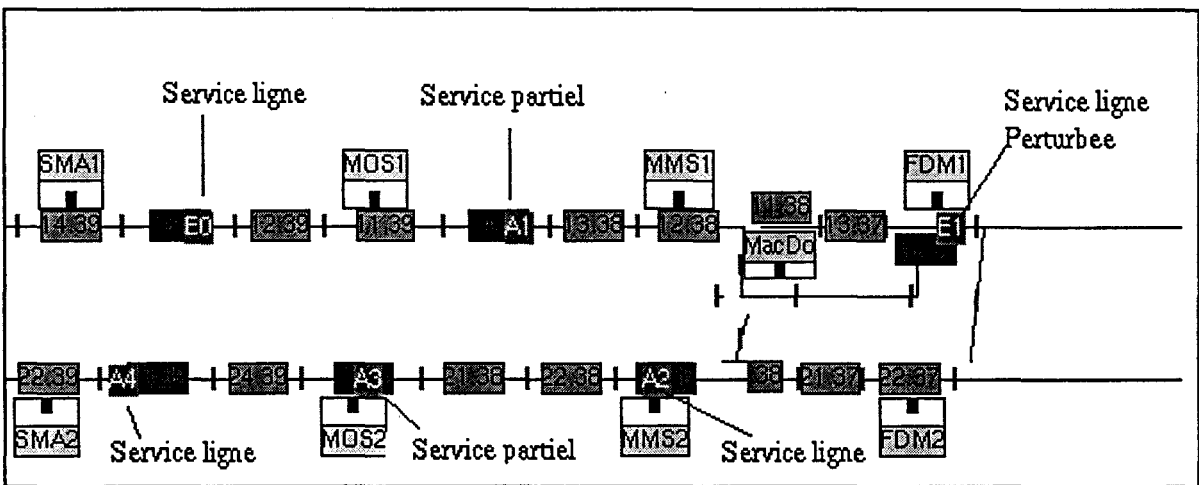


Figure 4.16. Perturbation de la rame service ligne avant son entrée dans la boucle partielle

Lors de l'arrivée de la rame service partiel A3 au terminus partiel MMS2, pour éviter le passage de cette rame avant la rame E1 à la station MMS1, nous introduisons une rame fictive seulement visible par la rame A3 sur le canton après l'aiguillage (canton noté 12.38). La création d'une rame fictive permet d'empêcher éventuellement le départ de la rame A3 suite à un

calcul de risque sur la rame fictive située devant elle. De ce fait la rame A3 sur-stationne à la station MMS2. Dès que le calcul de risque l'autorise, la rame A3 part de la station où elle sur-stationnait. Ce calcul certifie le passage de la rame E1 au terminus partiel MMS1 avant la rame A3.

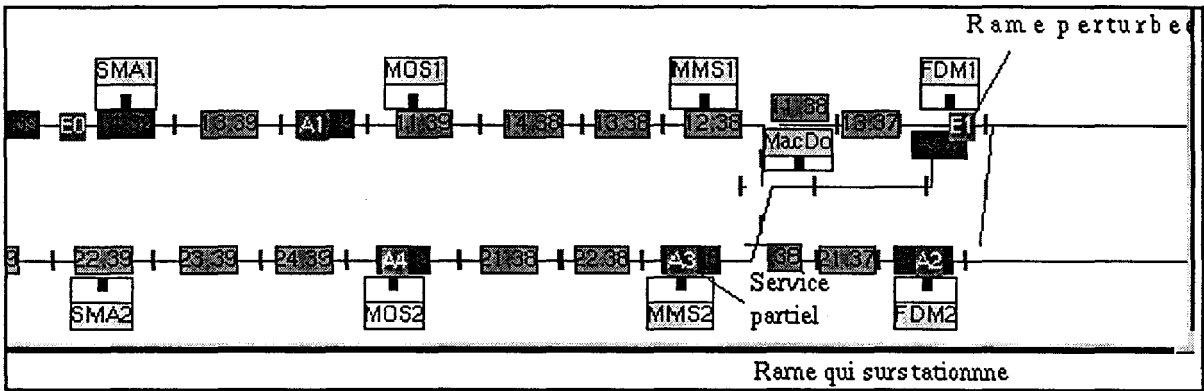


Figure 4.17. Création d'une rame fictive visible seulement pour A3

Aussitôt que la rame E1 est arrivée au terminus départ partiel MMS1, nous supprimons la rame fictive visible pour A3 et nous en ajoutons une autre qui sera visible seulement pour E2, pour permettre à la rame A3 de prendre le départ en MMS1 avant la prochaine rame service ligne (notée E2).

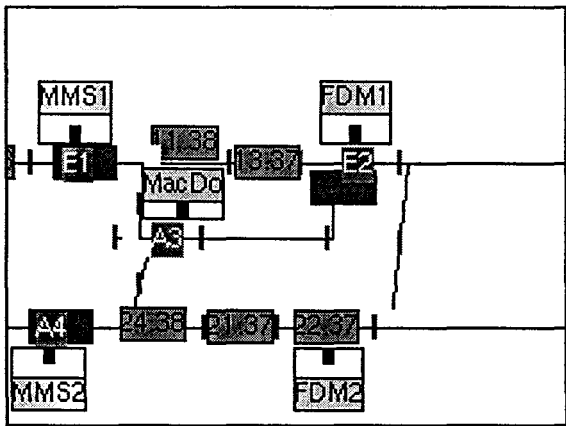


Figure 4.18. Arrivée de la rame E1 au terminus partiel, rebroussement de la rame A3 et création d'une rame fictive pour E2

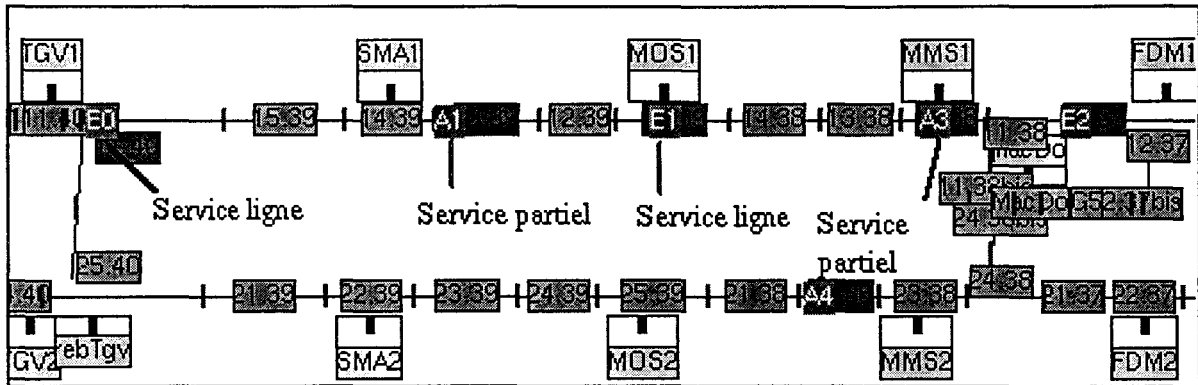


Figure 4.19. Rebroussement des rames service partiel en synchronisation avec les autres rames

Comme nous pouvons le voir sur la figure ci-dessus, les deux rames service partiel A1 et A3 se sont insérées dans le carrousel assurant sur la boucle service partiel une succession de rames étiquetées successivement service ligne et service partiel.

5. Réalisation d'un estimateur de la durée de perturbation d'un train

Deux approches sont envisageables pour déterminer la durée de la perturbation : la première est d'établir un diagnostic précis de la panne à l'aide des informations disponibles sur l'état des infrastructures. Ce diagnostic sera une aide pour l'opérateur car il lui proposera des actions à mener pour la résolution de l'incident. De ce fait, nous aurions une idée assez précise de la durée de cette perturbation puisque l'incident serait parfaitement connu [KHO 94]. L'autre approche est d'utiliser le peu d'informations disponibles sur l'état du système, pour classer ces pannes en types plus ou moins précis et d'utiliser les données statistiques sur les durées de perturbations provoquées par ces pannes [CRO 95]. Cette approche choisie laisse le régulateur au poste de commande effectuer son diagnostic sur la panne.

Par définition, nous pouvons classer les pannes en types différents. Chaque type induit une perturbation qui dure en moyenne un certain temps. Il est facilement concevable que les différents types de perturbations n'aient pas la même durée moyenne. Par exemple, il est raisonnable de penser qu'un passager, empêchant une porte de se fermer, provoque une perturbation moins importante qu'une alerte au feu dans une rame. Dans les lignes de métro automatisées, un certain nombre d'informations sur l'état du système sont disponibles au poste de commande. Nous pouvons regrouper les types de pannes en considérant la première alarme visible sur l'état du système. Le classement de ces pannes se veut large pour empêcher un

diagnostic complet de la panne. Par exemple, nous classerons une panne comme "défaut d'une porte d'une rame" lorsqu'une porte d'une rame est perturbée. Nous ignorons le détail de la cause de la non-fermeture ou de l'absence d'ouverture de cette porte. De ce fait, aussitôt que la perturbation se produit, nous connaissons le type dans lequel elle est représentée. Bien évidemment, si d'autres alarmes sur l'état du système nous fournissent une indication permettant de classer la panne dans un autre type, alors nous modifierons le classement de cette panne et de ce fait l'estimation de sa durée. Le changement de cette estimation modifiera la prise de décision sur le départ de la rame située à la station en amont de la perturbation.

5.1. Différents moyens de réaliser une prédiction de la durée de la panne

Les principaux moyens de prédiction d'une durée de panne sont tirés de l'étude de la sûreté de fonctionnement des systèmes. Ces méthodes proposent de mesurer la fiabilité, la disponibilité, la maintenabilité et la sécurité des systèmes [VIL 88]. De ces définitions, on peut déduire une prédiction de ces paramètres grâce à une étude probabiliste des événements qui se sont déroulés lors du fonctionnement.

Les paramètres du système peuvent être modélisés par différentes lois de probabilités dont les plus courantes sont la loi binomiale, la loi de Poisson, la loi exponentielle, la loi normale, la loi log - normale, la loi de Weibull. Il ne suffit pas, pour modéliser une variable aléatoire, d'estimer les paramètres d'une loi ; il faut également s'assurer que la variable aléatoire considérée est effectivement régie par cette loi. Dans ce but, il existe des tests d'hypothèses : les deux principaux sont le test du khi-deux et le test de Kolmogorov-Smirnov.

Calcul d'estimateurs

Il est important d'affecter chaque estimateur d'un intervalle de confiance à un niveau de confiance donné. Par principe, on cherche ainsi à évaluer les bornes d'un intervalle encadrant la valeur de l'estimateur. Si on affirme que le paramètre appartient à cet intervalle, il y a une probabilité (a) d'erreur. De manière simplifiée, on peut dire que le paramètre cherché aura une probabilité donnée (a) de non-appartenance à cet intervalle. Celui-ci est appelé intervalle de confiance et la quantité (1-a) est appelée niveau de confiance.

Par exemple, si nous calculons l'estimation et l'intervalle de confiance d'un taux de défaillance en fonctionnement. Soit :

$\hat{\lambda}$: Estimateur du taux de défaillance

λ_{sup} : Borne supérieure de l'intervalle de confiance au niveau de confiance $1-\alpha$

λ_{inf} : Borne inférieure de l'intervalle de confiance au niveau de confiance $1-\alpha$

On a :

$$\hat{\lambda} = \frac{N_f}{T_f} = \frac{\text{Nombre de défaillances en fonctionnement}}{\text{Durée cumulée de fonctionnement}}$$

Si on suppose que λ est constant, le nombre de défaillances est distribué suivant une loi de Poisson. On en déduit que les limites de l'intervalle de confiance sont les suivantes :

$$\lambda_{\text{sup}} = \frac{\chi^2_{1-\frac{\alpha}{2}}(2N_f + 2)}{2T_f} \quad \text{et} \quad \lambda_{\text{inf}} = \frac{\chi^2_{\frac{\alpha}{2}}(2N_f)}{2T_f}$$

Avec $\chi^2_{\alpha}(v)$ valeur du test du Khi-deux.

Dans cette approche, il est indispensable de connaître la loi de probabilités qui régit le phénomène ; or, cette loi de probabilités n'est pas toujours évidente à trouver. Dans notre cas, sa détermination est difficile à cause du nombre peu élevé d'informations disponibles sur le comportement du système. En effet, même si la durée d'exploitation du VAL de Lille est de plus de dix ans, le nombre de pannes se produisant sur le système reste trop faible pour avoir des données suffisamment riches. De plus, nous ne nous intéressons qu'à certaines classes de perturbations et non pas à leur cause. Les classes regroupent des événements indépendants n'ayant pas la même cause. De ce fait, la loi régissant les pannes de la classe A peut être la somme de plusieurs lois de probabilités dont les conséquences provoquent des pannes de type A. Par exemple, la non fermeture d'une porte peut être provoquée par la malveillance des voyageurs ou la panne mécanique. Comme nous essayons de prédire la durée de la panne au moment où celle-ci est détectée, nous n'avons pas d'informations à ce moment sur la cause ayant produit le cas "non fermeture d'une porte". La seule information disponible sur cette panne est l'état de la porte en question. Nous avons donc développé un algorithme ne nécessitant pas la connaissance précise des lois de probabilités régissant les différentes pannes.

5.2. Approche retenue

Lors d'une panne sur un train i , certaines informations sont disponibles sur l'état général du système notamment :

- La position du train perturbé en station ou entre deux stations.
- La probabilité d'apparition de chaque type de panne en fonction de l'heure de la journée.
- La densité de probabilités de la durée de la panne pour chaque type de panne.
- Le temps écoulé depuis le début de la panne.

Comme nous l'avons vu précédemment, la densité de probabilités de la durée de la panne pour chaque type n'est pas connue précisément puisque nous ne connaissons pas la loi de probabilités qui modélise ce type de panne. D'après l'étude des perturbations, durant les années précédentes, nous obtenons une figuration graphique de la distribution de la variable sous forme d'histogramme.

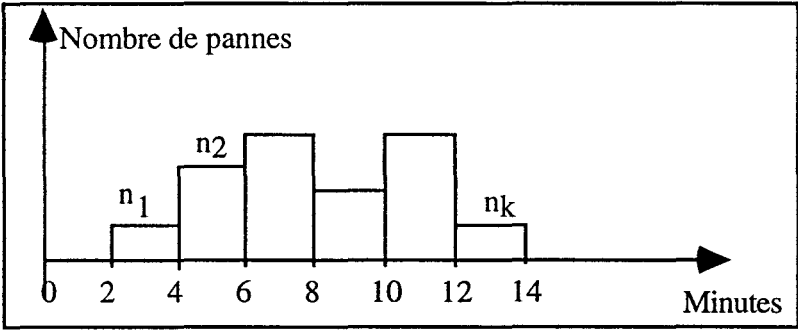


Figure 4.20. Distribution de la variable panne de type 1

Dans cet histogramme, si l'amplitude de la classe - les classes étant supposées égales - est prise pour unité, le premier rectangle de l'histogramme aura pour superficie $n_1, \dots$ le $k^{\text{ième}}$ n_k , en conséquence la surface totale de l'histogramme sera N , somme du nombre de pannes constatées. Nous aurions eu un histogramme de même forme si au lieu de construire des rectangles de hauteurs $n_1, \dots, n_k$, nous avions construit des rectangles de hauteurs $\frac{n_1}{N} = p_1, \dots, \frac{n_k}{N} = p_k$; dans ce cas la surface totale de l'histogramme aurait été égale à 1.

Si la variable continue est telle que l'amplitude des intervalles de classes puisse être réduite jusqu'à devenir presque nulle, la représentation graphique de la variable sera continue et la fonction qui traduira la loi de probabilités de la variable aléatoire x sera $y = f(x)$.

Si l'on s'intéresse à la probabilité contenue dans un intervalle de 6 à 8 minutes, la probabilité sera égale à la surface représentée par la figure ci-dessous :

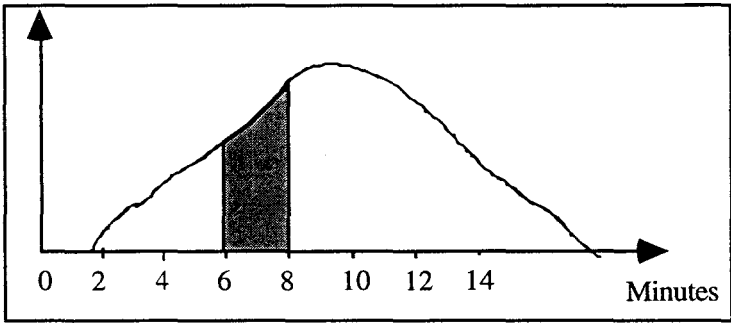


Figure 4.21. Distribution de la variable panne de type 1

Si on réduit l'amplitude de l'intervalle jusqu'à le rendre infiniment petit, égal par exemple à dx , la surface figurant la probabilité correspondant à ce petit intervalle pourra être assimilée à un rectangle de base dx et de hauteur $f(x)$, donc de surface $f(x) \cdot dx$. Ainsi la probabilité élémentaire sera égale à $f(x) \cdot dx$. On appellera $y = f(x)$ densité de probabilités de la variable continue aléatoire x .

Sans la connaître, nous désignerons par densité de probabilités de la variable x , la fonction représentée par un histogramme qui représente la distribution de chaque type de panne. C'est en fait un abus de langage.

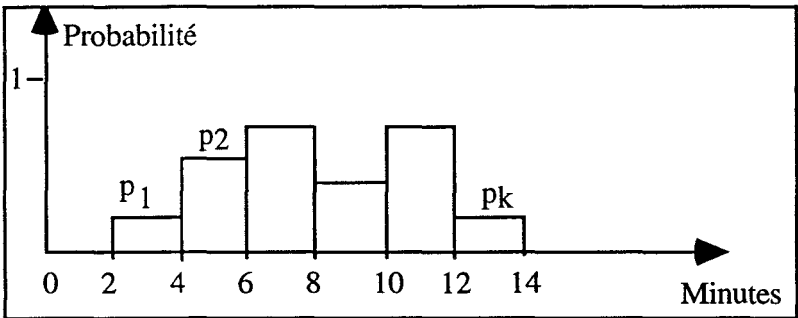


Figure 4.22. Loi de probabilités de la panne de type 1

Bien évidemment, notre algorithme donnera de meilleurs résultats en connaissant la loi de probabilités et donc la densité de probabilités de chaque type de panne. L'algorithme présenté ci-dessous suppose que ces lois sont de forme exponentielle. En fait, comme expliqué précédemment, nous utiliserons les densités de probabilités représentées sous forme d'histogramme.

L'estimation de la durée peut être réalisée selon deux approches : la théorie des probabilités et la théorie des possibilités.

5.2.1. Résolution par la théorie des probabilités.

Certaines informations statistiques peuvent être disponibles grâce à une étude des différentes perturbations survenues lors de l'exploitation de la ligne durant les années précédentes. Nous avons limité volontairement le nombre de types de pannes à trois dans un souci de clarté de présentation. Dans la réalité on pourrait distinguer 17 types différents. Les informations disponibles sont :

- Une probabilité d'apparition pour chaque type de panne en fonction de l'heure.

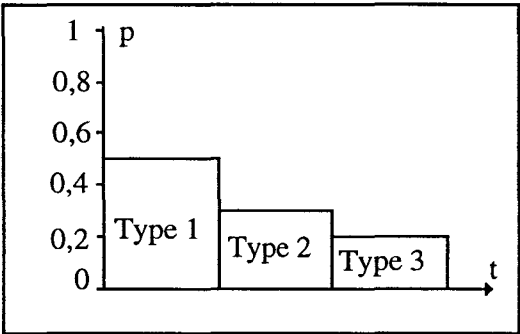


Figure 4.23. Probabilité d'apparition de chaque panne

- Une densité de probabilités de durée de la panne pour chaque type de panne.

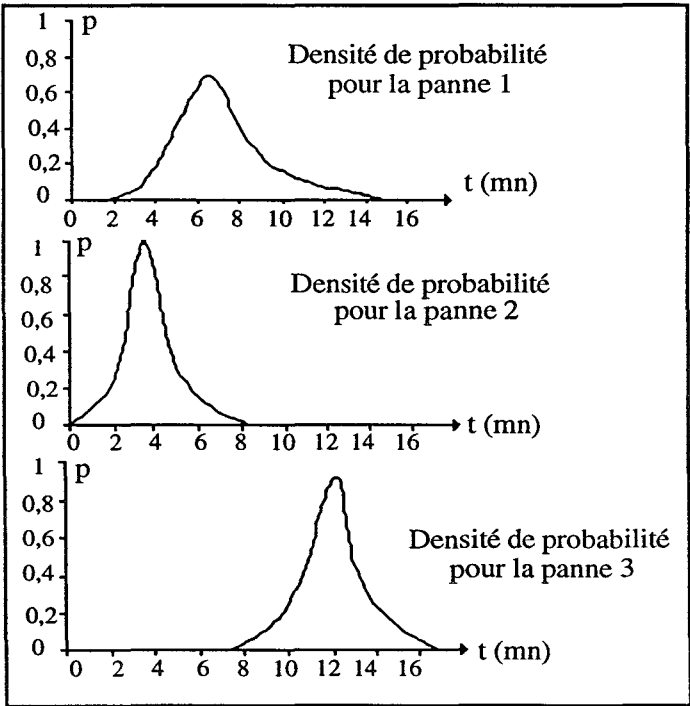


Figure 4.24. Densité de probabilités pour chaque type de panne

La densité de probabilités finale est égale à la somme des densités de probabilités, pondéré par la probabilité d'apparition de chaque panne. La densité résultante est représentée ci-dessous :

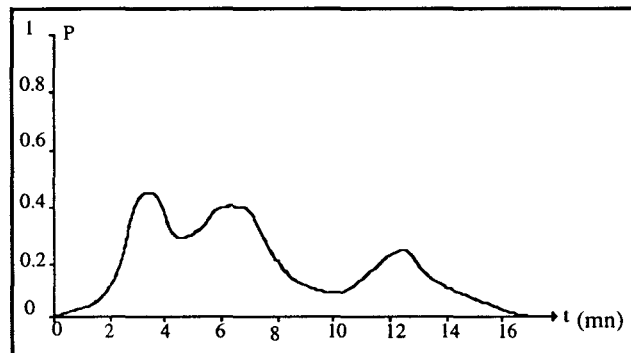


Figure 4.25. Densité de probabilités totale

Une fois cette fonction déterminée, une estimation du temps de panne peut être obtenue en s'intéressant à l'intervalle de temps de longueur fixée (ex : 30 secondes) qui maximise sa probabilité d'apparition. La longueur de l'intervalle pourra jouer un rôle considérable dans le résultat final, et ce, suivant l'allure de la densité résultante obtenue. En effet, le fait de prendre un intervalle de petite taille favorisera la détection des pics. Un intervalle de longueur plus importante aurait tendance à détecter des blocs compacts sous la courbe. Le choix de la longueur de l'intervalle sera facilité, lors de l'implantation du programme sur la simulation, par des essais de différentes valeurs de ce paramètre.

Cependant, des pannes peu fréquentes mais de durée importante existent. Le traitement des informations disponibles par la théorie des probabilités entraîne alors une minimisation du rôle joué par ces pannes. Il peut être judicieux de s'intéresser non plus à une probabilité d'apparition de la durée de la panne mais à sa possibilité.

5.2.2. Résolution par la théorie des possibilités

Grâce aux transformations de Dubois et Prade [DUB 93] et de S. Sandri [SAN 91], il est possible de convertir une mesure de probabilités en une mesure de possibilités.

Soit $X=\{x_1...x_n\}$ l'ensemble de référence et les x_i parties élémentaires. Nous trouvons la distribution de possibilités $\pi(x_i)$ grâce à la densité de probabilités $p(x_i)$ par la formule :

$$\forall i, \pi(x_i) = \sum_{j=1}^n \left\{ p(x_j) / p(x_j) \leq p(x_i) \right\} \quad (1)$$

Les nouvelles données disponibles sont donc des distributions de possibilités des différents types de pannes ainsi qu'une possibilité d'apparition pour chaque type de panne. L'agrégation de ces différentes informations va s'effectuer en deux temps :

- Le calcul de la nouvelle distribution de possibilités pour chaque panne connaissant sa probabilité d'apparition,
- Le calcul de la distribution de possibilités de l'union de toutes les pannes.

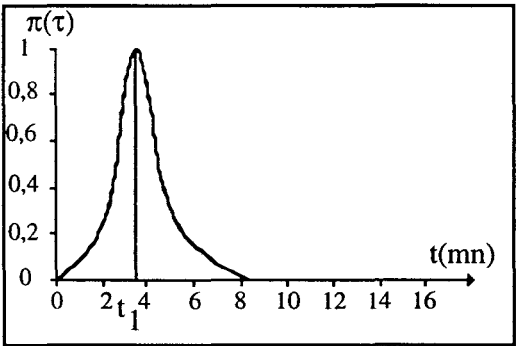


Figure 4.26. Possibilité que la panne dure t_1 minutes sachant qu'elle est de type 2

La courbe précédente doit maintenant être pondérée par la possibilité que la panne soit de type 1 (figure 4.27). La pondération se fera grâce à l'opérateur multiplication car l'opérateur minimum déforme de manière assez nette la courbe. Cette déformation se traduit par une perte de précision sur les ensembles flous.

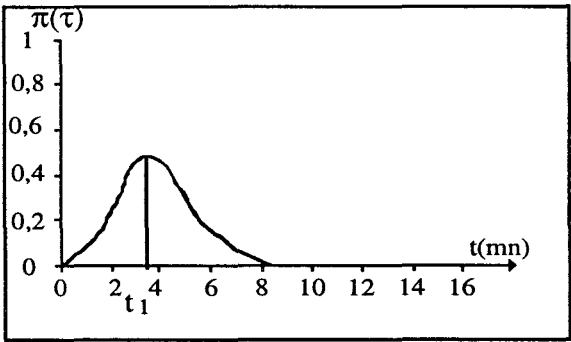


Figure 4.27. Distribution de possibilités pondérées

La distribution de possibilités recherchée est celle de l'union de toutes les pannes possibles. Il suffit donc d'agréger les trois distributions obtenues par le premier calcul par l'opérateur max. Le résultat est le suivant :

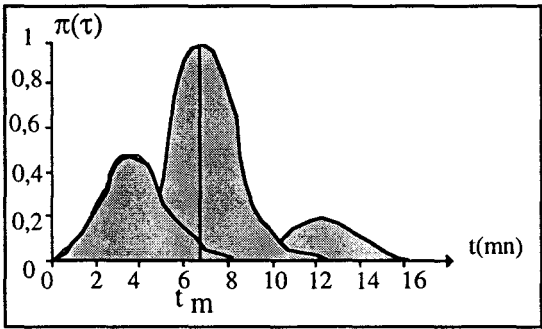


Figure 4.28. Possibilité que la panne rencontrée dure t_m minutes

Il reste à extraire de cette distribution de possibilités une estimation de la durée de la panne. La durée estimée de la panne la plus possible est celle qui a le degré de possibilités le plus élevé : la valeur t_m sur la figure 4.28.

5.2.3. Intervention du temps dans les calculs

Les probabilités d'apparition de chaque panne ne sont pas constantes selon le temps qui s'écoule depuis le début de la perturbation. Il faut donc agir sur les coefficients de répartition des pannes, c'est à dire sur les probabilités d'apparition pour chaque type de panne en fonction du temps écoulé. La probabilité d'apparition de la panne A_i , sachant que t_e minutes se sont écoulées depuis le début de la panne, est égale à l'intégrale de t_e à l'infini de sa densité de probabilités.

Nous calculons les coefficients de pondération A_i de la probabilité statique d'apparition pour chaque type de panne lorsque la perturbation dure depuis $t_e = 6$ minutes :

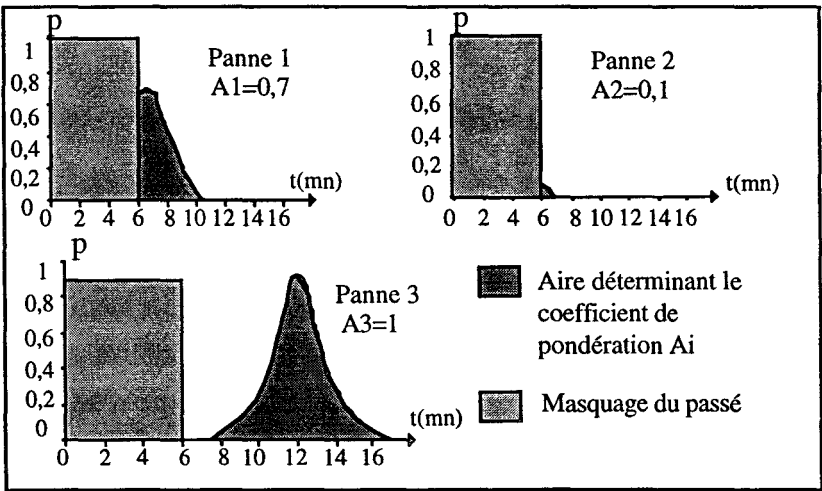


Figure 4.29. Calcul du coefficient de pondération de la probabilité statique d'apparition des trois types de pannes

Nous pondérons cette intégrale A_i par la probabilité statique d'apparition de la panne et nous obtenons la nouvelle probabilité de cette panne (probabilité dynamique).

Les probabilités d'apparition des types de panne 1, 2 et 3, avant le début de la panne, étaient respectivement de $P_1=0,5$, $P_2=0,3$ et $P_3=0,2$ (probabilités statiques).

Les probabilités dynamiques sont donc égales à :

$$P' 1 = 0,5 * 0,7 = 0,35$$

$$P' 2 = 0,3 * 0,1 = 0,03$$

$$P' 3 = 0,2 * 1 = 0,20$$

Soit après normalisation :

$$P' 1 = \frac{0,35}{0,35 + 0,03 + 0,2} = 0,60$$

$$P' 2 = 0,05$$

$$P' 3 = 0,35$$

Ainsi pour estimer la durée de la panne en fonction du temps écoulé depuis le début de la perturbation, nous calculons une nouvelle densité de probabilités finale (figure 4.25). Elle est égale à la somme des densités de probabilités de chaque type de panne (figure 4.24), pondérée par la nouvelle probabilité de chaque panne P'_i .

D'après une étude statistique des pannes survenues sur la ligne 1 du VAL de Lille, nous pouvons regrouper les pannes en types différents, obtenir les probabilités d'apparition et les densités de probabilités de la durée de ces pannes pour chaque type de panne envisagé. Lors de chaque panne détectée, nous sommes en mesure de fournir une durée estimée de cette perturbation. Nous agissons en conséquence sur le départ (ou non), de la rame située en amont de la perturbation.

5.2.4. Simulation de l'estimateur de la durée de la panne

Les algorithmes d'estimation décrits ci-dessus ont été implémentés dans la simulation de lignes de métro. Une base de données a été au préalable incluse dans la simulation et pour correspondre à la réalité, nous avons introduit les différentes perturbations rencontrées lors de l'exploitation de la ligne 1 du VAL de Lille. Pour permettre une estimation cohérente, nous avons réparti ces pannes en différentes classes, d'après les informations disponibles au Poste Central de Contrôle et d'après les principales pannes rencontrées. Un outil de simulation a été créé afin de tester la validité de l'estimateur.

Test de l'estimateur du temps de panne

Type de perturbation

- ☐ défaut fermeture portes paliers
- ☐ défaut fermeture porte véhicule
- ☐ deverrouillage porte paliers
- ☐ detection obstacle
- ☐ Disjoncteur HT
- ☐ Poignée EVAC
- ☐ Défaut isolement
- ☐ Défaut frein permanent
- ☒ FU
- ☐ Défaut suspension
- ☐ défaut compresseur
- ☐ Défaut conduite principale
- ☐ Coupure haute tension
- ☐ Perte reception FS
- ☐ Rame bloquée en terminus
- ☐ ADLC transition
- ☐ Divers

Definition temps Pertu

Temps depuis le debut

Temps depuis le debut de la panne

0 mn

SIMULATION

Statistique

Erreur de probabilité : 0

Erreur de possibilité : -10

Nombre de pannes simulées : 1

Durée de perturbation

Durée de la panne : 4 minutes(s)

BOUCLAGE

Système non boucle

Bouclage oui/non

Durée estimée

Durée proba : 4.0

Durée possi : 3.6

Figure 4.30. Outil de validation de l'estimateur de la durée de la panne

En premier lieu, nous choisissons comme base de données, les perturbations rencontrées durant deux années d'exploitation. Ensuite, nous tenons compte des événements qui se sont déroulés lors d'autres années et nous estimons la durée de chaque perturbation. Ainsi, nous disposons du temps de perturbation réel subi par la rame et de son temps de perturbation estimé. Une comparaison entre ces deux temps permet de mesurer l'efficacité de l'estimateur.

Nous pouvons regrouper les événements qui sont déroulés en classes qui correspondent à différentes causes. Nous pouvons représenter la répartition des perturbations durant une année :

Causes	Cause 1	Cause 2	Cause 3	Cause 4	Cause 5	Cause 6	Cause 7	Cause 9
Durée totale (mn)	52	98	16	160	46	112	45	49
Nombre	10	11	3	16	5	12	5	5

Pour chaque perturbation constatée, nous avons estimé le temps de perturbation lorsque celle-ci est détectée (t_0) puis 4 minutes après la détection. Nous calculons l'erreur totale sur toutes les estimations par rapport à la réalité pour l'approche probabiliste et l'approche possibiliste.

Instant de l'estimation	t_0	$t_0+4\text{ mn}$
Erreur totale de l'estimation par une approche probabiliste (%)	15,64	14,68
Erreur totale de l'estimation par une approche possibiliste (%)	26,06	23,53

Nous pouvons voir que l'estimation de la durée de la panne donne de bons résultats. En effet, dans notre cas, faire une erreur de 15 % sur la durée représente une bonne estimation. Si une panne se produit et qu'elle dure 4 minutes, notre prédiction de sa durée prévoit une durée comprise entre 3,4 minutes et 4,6 minutes. Ceci est largement suffisant pour éviter les arrêts en anticollision.

Les résultats nous montrent que l'approche probabiliste est meilleure que l'approche possibiliste. Ceci est dû aux nombres importants d'échantillons dans la base de données : la théorie des probabilités donne de meilleurs résultats. Par contre, lors d'événements peu fréquents, il apparaît que la théorie des possibilités donne de meilleurs résultats ; cela correspond à la définition de la théorie des possibilités par rapport à la théorie des probabilités.

Au fur et à mesure du prolongement de la perturbation, la qualité de l'estimation s'améliore. Ainsi, lorsque la perturbation est détectée, l'erreur est de 15,64 % pour l'approche probabiliste. A chaque instant, nous recalculons une nouvelle estimation de la durée de perturbation si elle se prolonge. Lorsque 4 minutes se sont écoulées depuis le début de la perturbation, les événements plus courts que cette durée ne sont plus pris en compte dans notre calcul. Nous affinons ainsi le résultat, ce qui nous permet d'améliorer l'efficacité de l'estimateur puisque nous obtenons une erreur de 14,68 %.

Les erreurs de prédiction de la durée montrent les limites d'une telle approche. En effet, si ces résultats sont suffisants dans notre cas, pour avoir une précision meilleure sur le temps de panne, il est nécessaire d'avoir une approche davantage basée sur le diagnostic de panne. Cette approche est plus complexe et fait appel à des notions de fiabilité plus qu'à des approches basées sur les statistiques [VIL 88].

6. Conclusion

Tous les résultats de simulation montrent l'efficacité de notre système de régulation. Le niveau bas, utilisant les contrôleurs flous, permet d'améliorer la régularité du passage des rames en station sans pour autant modifier la ponctualité. Dans le niveau haut, grâce aux calculs de risque et de non-risque, nous évitons les situations de conflits qui provoquent des arrêts en anticollision. Cette régulation permet d'améliorer considérablement la qualité de service car les personnes ne se retrouvent plus bloquées en tunnel. De plus, ce calcul de risque permet la résolution du fonctionnement d'une ligne avec un service partiel. D'après les essais de simulation, notre approche afin de prédire la durée d'une perturbation s'avère suffisante pour notre système de régulation.

Nous concluons que, d'après les simulations réalisées sur un simulateur représentant suffisamment bien la réalité d'une ligne de métro automatisée, notre système de régulation se compose de modules cohérents permettant de réguler au mieux une ligne de métro automatisée. A chaque niveau de régulation, notre module permet d'améliorer la qualité de service par rapport à la régulation existante.

Conclusion générale

Depuis la première application industrielle de la théorie de la logique floue dans les années 1970, le nombre d'applications ne cesse de s'accroître, tout particulièrement dans la commande des systèmes. Ce phénomène s'explique par la possibilité, lors de la construction d'un système flou, d'intégrer l'analyse qualitative d'un expert.

Cette facilité de construction nous a permis de construire une commande floue permettant de réguler les rames d'une façon plus souple, s'approchant fortement d'un contrôle humain. Néanmoins, l'approche par logique floue doit être perçue comme un outil additionnel aux méthodes classiques de commandes déjà existantes.

Nous avons choisi l'approche par logique floue car sa principale caractéristique est sa capacité à la représentation de l'incertitude et de l'imprécision. Or, dans un système de transport, ces deux caractéristiques sont fortement présentes. Nous nous sommes servis de cette représentation de l'incertitude et de l'imprécision pour construire, dans un deuxième temps, un module de supervision qui empêche les rames de partir de la station où elles se trouvent si une rame est perturbée devant elles. Ce module repose sur une estimation du temps de panne de la rame perturbée et d'une estimation des temps de parcours des rames. Bien évidemment, ce module dépend, comme nous l'avons vu, de l'intervalle entre les rames. Si les rames sont trop rapprochées les unes des autres, aucune perturbation ne sera détectée à temps. Néanmoins, ce module permet d'éviter certaines situations de conflits où les passagers se retrouvent bloqués dans une rame entre deux stations. De ce fait, la qualité de service, qui est notre unique critère, s'en trouve considérablement améliorée.

Pour que notre estimateur de durée de la panne soit cohérent, il faut que chaque nouvelle perturbation soit entrée dans une base de données. Si cette base de données n'est pas mise à jour, cela risque d'entraîner des résultats ne représentant pas la réalité. En effet, une ligne de métro évolue : certains incidents fréquents sont résolus grâce à une analyse pertinente de leurs causes et grâce aux services de maintenance. De ce fait, si la base de données n'est pas

mise à jour, les résultats de l'estimateur seront totalement faussés et notre module de supervision n'aura pas un fonctionnement optimal.

L'approche développée pour la réalisation de l'estimateur est une approche essentiellement basée sur l'utilisation des statistiques d'exploitation de la ligne de métro. Nous pourrions nous baser sur une approche tournée vers la prédiction de panne. Ce procédé nous permet d'anticiper sur les prochaines pannes et de limiter ainsi le nombre de pannes et, par conséquent, leurs conséquences sur la ligne.

Le but de l'étude est d'améliorer la qualité de service du système de transport. Or, dans toute notre étude, nous avons pris comme indicateurs de qualité de service ceux utilisés par le système de transport étudié. Pour réaliser une comparaison plus objective, il serait plus raisonnable de construire des indices de qualité de service plus généraux, par exemple le temps de trajet moyen d'un voyageur. Mais ce type de critère nous oblige à construire un modèle de flux de passagers fréquentant le système de transport. Ce modèle de flux de passagers représente, à lui seul, une étude complète. En effet, le nombre de passagers attendant une rame sur un quai dépend :

- de l'heure,
- du nombre total de passagers fréquentant le métro à cette heure,
- de l'emplacement de la station sur la ligne (la station se trouve au milieu de la ligne, en zone de forte densité urbaine, à un terminus de bus...),
- du nombre de passagers présents dans les rames qui sont déjà passées par cette station ou, plus précisément, le nombre de passagers qui ont pu monter dans les rames précédentes,
- du retard du train sur les différentes stations précédentes (si le train a pris trop de retard, le nombre de passagers attendant sur les quais précédents est plus important ; par conséquent, le nombre de passagers dans la rame est plus important),
- du nombre de passagers présents dans les prochaines rames et qui vont descendre à cette station ...etc.

Comme on peut le voir, le trajet moyen d'un voyageur dépend d'un nombre important d'événements interdépendants. Le temps moyen peut être seulement calculé si on se construit un modèle représentatif du flux de voyageurs. Ce critère peut paraître utopique mais un certain

nombre d'autres critères plus représentatifs peuvent être trouvés pour permettre une comparaison plus objective.

Le niveau bas de la régulation prend en compte les retards par groupe de 3 rames, cette approche a été adoptée pour limiter la complexité du contrôleur flou. Néanmoins, cette approche n'optimise pas les commandes appliquées sur les rames. En effet, la commande appliquée sur la rame i dépend du moment où celle-ci s'arrête en station, instant où l'on applique la commande sur la rame, or les retards des rames sont en constante évolution. De ce fait, si la rame s'était arrêtée un instant plus tard, la commande appliquée aurait été différente. Cette différence s'explique par l'évolution des retards des rames amont et aval, eux mêmes dépendants des rames situées devant et derrière elles. Dans notre algorithme, nous supposons que ces commandes "fausses" ne jouent pas un rôle déterminant dans la récupération des retards. Par contre, si nous voulons optimiser l'intervalle entre les rames en tout point de la ligne, il est nécessaire d'ajouter à notre niveau bas de régulation, un module de supervision. Ce superviseur prendrait en compte l'ensemble des retards des rames sur la ligne et ferait une prédiction à court terme sur les conséquences des commandes appliquées sur les rames. De ce fait, l'ensemble des rames serait commandé d'une façon cohérente, respectant mieux l'intervalle entre les rames.

Actuellement, aucun réseau de transport en commun ne prend en compte les informations des autres réseaux. Par exemple, lors de correspondances, nous pourrions utiliser les informations sur les retards à l'arrivée d'un réseau de transport pour réguler un autre réseau et assurer ainsi une correspondance parfaite entre les différents modes de transport pour diminuer les attentes des voyageurs. Ce type de régulation nous amène à considérer non plus un réseau, mais plusieurs systèmes interconnectés les uns aux autres. Dans un système complexe hiérarchisé, on ajouterait un niveau supérieur qui serait l'ensemble des réseaux de transport connectés les uns aux autres. L'ajout d'un niveau supérieur changerait les consignes et les buts de la régulation des niveaux inférieurs.

Bibliographie

- [ADA 93] Adams M.B., Kolitz S.E. "Hierarchical Rail Traffic Flow Control", Computers in Railways, p 211-222, 1993.
- [BAI 93] Bailly E. "Introduction de la commande floue dans la régulation de trafic de lignes de métro automatisées", Rapport de D.E.A, USTL-CAL, 1993.
- [BAI 1-94] Bailly E. "La théorie des sous ensembles flous appliquée à la régulation de trafic de lignes de métro", Tome 1, Rapport interne INRETS 94-79, août 1994.
- [BAI 2-94] Bailly E., Hayat S., Rodriguez J., Jolly D., Desodt AM. "Comparison between determinist and fuzzy logic techniques in automatic subway traffic control", Congrès WCRR, Paris 14-16 Novembre 1994.
- [BAI 1-95] Bailly E., Jolly D., Desodt AM. "Intelligent subway traffic control using information on breakdowns", EUFIT'95, Aachen, Germany, august 28-31, 1995, p 1725-1727.
- [BAI 2-95] Bailly E., Hayat S., Desodt AM., Jolly D. "Subway line one of Lille simulation and regulation based on fuzzy logic", Systems science XII, Wroclaw-Poland, 12-15 September 1995, p 331-338, vol III.
- [BAI 3-95] Bailly E. "La théorie des sous ensembles flous appliquée à la régulation de trafic de lignes de métro", Tome 2, Rapport interne INRETS 95-68, octobre 1995.
- [BAI 1-96] Bailly E., Jolly D., Desodt AM. "Gestion de métro automatique utilisant la théorie des ensembles flous", Journal Européen des Systèmes automatisés, vol 30 N°1/1996, p 103-122, 1996.
- [BAI 2-96] Bailly E., Hayat S., Jolly D., Desodt AM. "Use of the statistics of breakdowns occurred during the exploitation of a line of automatic subway for the amelioration of traffic regulation", CESA'96-IEEE, Lille, France, july 9-12, 1996, p 67-72.

- [BAI 3-96] Bailly E., Hayat S., Jolly D., Desodt AM. "Command and control of automated subway in mode of disrupted march using fuzzy logic", soumis à la revue Intelligent and Fuzzy Systems.
- [BOR 85] Bortolan G., Degani R. "A review of some methods for ranking fuzzy sets", Fuzzy Sets and Systems 15, North-Holland, p 1-19, 1985.
- [BOU 93] Bouchon-Meunier B. "La logique floue", Édition "Que sais-je?", Presses universitaires de France, 1993.
- [BRE 84] Brehmer B. "Organization for decision-making in complex systems", Tasks Errors and Mental Models, p 116-127, 1984
- [CAN 90] Cantelier S., Matkowska F. "Réalisation d'une plate-forme de simulation pour les systèmes de transport guidés", Rapport de stage INRETS CRESTA, juin 1990.
- [COU 88] Couvreur G., Heddebaut M., Helle P., Szelag M., Uster. G. "Bilan de cinq années d'exploitation du 13 juin au 9 juin 1988 VAL de Lille" Synthèse INRETS, 1988
- [CRO 95] Croquette M. "Estimateur de temps de panne pour le métro de Lille", Rapport de D.E.A, USTL-CAL, 1995.
- [DUB 83] Dubois D., Prade H. "Ranking of fuzzy numbers in the setting of possibility theory", Inform. Sci 30, p 183-224, 1983.
- [DUB 93] Dubois D., Prade H., Sandri S. "On the possibility/probability transformations", Ed. R. LOWEN M. ROUBENS Fuzzy logic - Kluwer Academic Publisher, 1993.
- [ERS 76] Erschler J., Roubellat F., Vernhes JP. "A decision process for the real time control of a production unit", Int. J. Prod. Res., Vol 14, N°2, p 275-284, 1976.
- [ERS 84] Erschler J., Fontan G., Merce C., Roubellat F. "Approche hiérarchisée pour la conduite de production", Bulletin de liaison de la recherche en informatique et automatique, p 15-26, 1984.
- [FOR 95] Fortemps P., Roubens M. "Ranking and Defuzzification Methods based on Area Compensation", Publication n°95.014, Institut de Mathématique, Université de Liege, 1995.

- [FOU 94] Foulloy L "Logique floue", Chapitre III, Observatoire français des techniques avancées, Masson, 1994.
- [GOG 94] Goguet J.L. "Développement d'un éditeur graphique de réseau pour systèmes de transports guidés", Rapport de fin d'étude, Université de Technologie de Compiègne, janvier 1994.
- [HAY 93] Hayat .S., R Hartani "Modélisation de la régulation de trafic des lignes de métro basée sur les approches neurofloues", rapport interne INRETS n°72, juin 1993.
- [JAI 76] Jain R. "Decision making in the presence of fuzzy variables", IEEE Trans. Systems Man Cybernet. 6, p 698-703, 1976.
- [JAI 77] Jain R. "A procedure for multiple-aspect decision making using fuzzy sets", Internat. J. Systems Sci. 8, p 1-7, 1977.
- [JAM 77] Jamshidi Mo. "Fuzzy control systems : hierarchy and stability", Groupe commande symbolique et neuromimétique du CNRS, 15 juin 1995, Université Paris 6, Rapport 95-4.
- [KHO 94] Khoualdi K. "Filtrage d'alarmes pour un système automatisé par une approche multi-agents", Thèse de doctorat, Paris 6, 1994.
- [MAM 77] Mamdani E.H. "Applications of fuzzy sets theory to control systems to survey", Fuzzy automata and decision process, Ed Gupta, Saridis, Gaines, p 77-88, 1977.
- [MOR 92] Moreau A. "Public transport waiting times as experienced by customers", Public Transport International, N°3, mars 1992.
- [OSH 88] Oshima H., Yasunobu S. et Sekino S. "Automatic train operation system based on predictive Fuzzy Control", International Workshop on Artificial Intelligence for Industrial applications, 1988.
- [ROD 1-94] Rodriguez J. "Modèle pour la simulation en langage orienté objet de lignes de transport guidé", Rapport interne INRETS 179, mai 1994.
- [ROD 2-94] Rodriguez J. et Vanec P. "Modèle de simulation de la ligne 2 du VAL de Lille. Évaluation de la régulation de trafic", Rapport interne INRETS 109, decembre 1994.

Bibliographie

- [SAN 89] Sanchez E. "Applications industrielles de la commande floue au Japon", AFCET/INTERFACES, n°8, p 19-25, juin 1989.
- [SAN 91] Sandri S. "La combinaison de l'information incertaine et ses aspects algorithmiques", Thèse de doctorat d'Université, Toulouse, 1991.
- [TAS 85] Taskin T., Allan J., A Simmons, Mellitt B. "Automatic service regulation for London Underground's Central Line", Railway Operations, p 63-71.
- [TON 85] Tong R.M. "An annotated bibliography of Fuzzy Control". M. Sugeno Editor- Industrial Applications of Fuzzy Control, Amsterdam, 1985.
- [VAN 89] Van Breusegem V. "Simulation et régulation de lignes de métro automatisées", Thèse de l'Université Catholique de Louvain, Octobre 1987.
- [VIL 88] Villemeur A. "Sûreté de fonctionnement des systèmes industriels. Fiabilité - Facteurs humains. Informatisation", Collection de la Direction des Études et Recherches d'Électricité de France, Edition Eyrolles, 1988.
- [YAS 85] Yasunobu S., Miyamoto S. "Automatic train operation system by predictive Fuzzy Control", M Sugeno edition, p 1-18, North-Holland, 1985.
- [ZAD 65] Zadeh L. " Fuzzy Sets", Information and control, vol. 8, p 338-353, 1965.
- [ZAD 77] Zadeh L. "Fuzzy sets as a basis for a theory of possibility", Fuzzy Sets and Systems 1, p 3-28, 1977.

Annexe 1. Indicateurs de qualité de service

1. Algorithme de Taskin

La fonction de coût à optimiser est de la forme :

$$A(t) * \sum_{i=1}^M WI_i^2 + B(t) * \sum_{i=1}^M \left(WI_i - \frac{1}{M} * \sum_{i=1}^M WI_i \right)^2 + C(t) * \sum_{i=1}^M SI_i^2 + D(t) * \sum_{i=1}^M \left(SI_i - \frac{1}{M} * \sum_{i=1}^M SI_i \right)^2 \quad (1)$$

Où A(t), B(t), C(t) et D(t) sont les zones de régulation et des facteurs de poids. Ils varient selon l'heure de la journée.

M est le nombre de trains en service,

WI_i l'index de l'intervalle de régulation,

SI_i l'index de la table horaire de régulation. Pour le train i c'est la valeur de l'écart entre son horaire réel et son horaire théorique.

On définit WI_i comme le poids total des indices TI_i et PI_i . Ces deux indices, attribués au train i, décrivent les retards causés par les perturbations sur les trains et les stations. Ils sont tels que :

$$WI_i = TI_i + k * PI_i$$

où k est déterminé en utilisant la valeur relative des retards de trajet des clients dans les trains et les stations.

$$\text{On définit } TI_i = \sum_{j=0}^N p_{i,j} * d_{i,j}$$

Avec N nombre de stations sur la route du train i depuis l'endroit de la perturbation jusqu'à la station où le retard des passagers n'est plus significatif.

Pour j différent de 0, $p_{i,j}$ est le nombre de passagers qui montent dans le train i à la station j .

$p_{i,0}$ est le nombre de passagers connaissant le retard sur leur trajet. Ce retard est indépendant de la station où ils embarquent.

$d_{i,j}$ est la moyenne des retards totaux du trajet du $j^{\text{ème}}$ groupe de passagers embarqués dans le train i .

$$\text{On définit } PI_i = \sum_{j=0}^Q p_{i,j} * d_{2_{i,j}}$$

Avec Q le nombre de stations sur le trajet du train i

$p_{i,j}$ le nombre de passagers attendant de monter dans le train i à la station j et ressentant un retard.

$d_{2_{i,j}}$ le retard moyen en station des passagers de la station j ressentant un retard par rapport à leur durée nominale de trajet.

Les composantes de la fonction de coût (1) peuvent être expliquées de la façon suivante. La première composante pénalise les retards sur les temps de trajet imposés à l'ensemble des passagers. Cette composante évalue le service rendu au passager. La seconde composante est complémentaire de la première et pénalise la non distribution des retards du temps de trajet. Ainsi, un groupe de passagers ne peut être sujet à des retards importants. La troisième composante pénalise explicitement le non respect de la table horaire aussi bien pour les rames en retard que celles en avance. La quatrième composante est complémentaire de la troisième et pénalise la distribution inégale des retards par rapport à la table horaire.

2. Indicateurs de qualité de service utilisés dans le VAL de Lille

2.1. Nombre de décalages horaires

Cette demande de décalage horaire est faite par le module de gestion de trafic en station, si celui-ci constate un retard supérieur à un certain seuil (en général 120 secondes) à l'arrivée d'une rame. C'est l'indice le plus important puisqu'il évalue les perturbations de la ligne.

2.2. Nombre de kilomètres commerciaux prévus et réalisés

Chaque jour un nombre total de kilomètres est prévu. Par suite des décalages horaires modifiant l'ensemble de la table horaire, tous les kilomètres ne sont pas effectués. Cet indice dépend directement du nombre de décalages horaires, il est donc équivalent à l'indice ci-dessus.

2.3. Respect de la vitesse commerciale

$$\text{Respect vitesse commerciale} = 1/N * \sum_{i=1}^{i=N} (T_{t_i} / T_{r_i})$$

Avec N : nombre de "services" réalisés.

T_{t_i} , temps de parcours théorique = heure théorique d'arrivée - heure théorique de départ

T_{r_i} , temps de parcours réel = heure réelle d'arrivée - heure réelle de départ

2.4. Durée d'exploitation

La durée d'exploitation est définie comme la durée totale du fonctionnement du métro durant la journée. C'est la différence entre l'heure de la dernière arrivée et l'heure du premier départ.

2.5. Nombre de pannes

La détection des pannes enregistrées dans les statistiques est basée sur la mesure des retards à l'arrivée des rames. Pour la rame i , le retard est défini comme la valeur absolue de la différence entre l'heure d'arrivée réelle et l'heure d'arrivée théorique. Nous pouvons en déduire un retard moyen attribué à la rame i :

$$\text{Retard moyen} = 1/N_i * \sum_{k=1}^{i=N_i} (\text{Retard de la rame } k)$$

Avec N_i représentant le nombre de rames qui ont pris un départ entre l'heure de départ réel et l'heure d'arrivée réelle de la rame i .

Une panne est détectée lorsque le retard moyen dépasse un certain seuil de début de panne. Elle se termine lorsque le retard moyen d'une rame suivante passe en dessous d'un seuil

de fin de panne. Il faut qu'une panne pénalise un certain nombre de rames pour qu'elle soit officiellement comptabilisée.

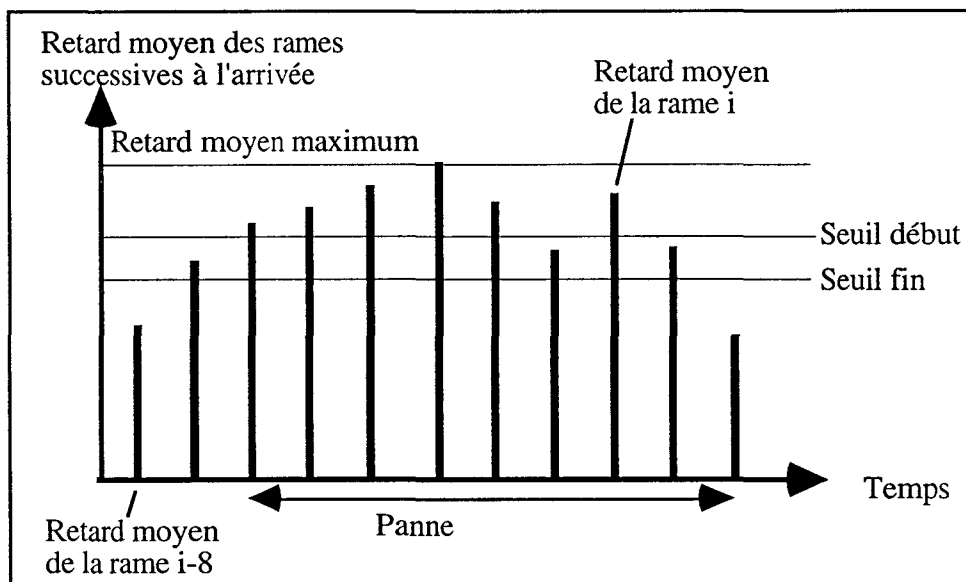


Figure A1.1. Principe de détection d'une panne

Pour chaque panne, on définit un retard moyen maximum.

2.6. Temps moyen entre chaque panne

C'est la durée d'exploitation rapportée au nombre de pannes mesurées.

2.7. Moyenne et total des retards moyens maximum.

Le total des retards moyens est égal à la somme de chaque retard moyen maximum attribué à chaque panne. La moyenne est égale à ce total divisé par le nombre de pannes constatées.

2.8. Assurance ponctualité

$$\text{Assurance ponctualité} = t_m / (t_m + \text{moyenne des retards moyen maximum}) * \exp^{(-T/t_m)}$$

Avec t_m : temps moyen entre chaque panne

T : durée moyenne d'un trajet fait par un voyageur

Il est à noter que la durée moyenne d'un trajet réalisé par un voyageur est une entrée du système. C'est à dire qu'à chaque édition de statistiques, le système demande à l'opérateur d'entrer cette durée.

2.9. Disponibilité

Disponibilité = $t_m / (t_m + \text{moyenne des retards moyens maximum})$

Avec t_m : temps moyen entre chaque panne

2.10. Indice de retard de service

Indice retard service = $\sum \text{Rat}_{>2mn} / \text{Cumul des temps de parcours théoriques}$

Avec $\text{Rat}_{>2mn}$ cumul des retards à l'arrivée supérieurs à 2 minutes

Un temps de parcours théorique est obtenu en calculant l'écart entre l'heure théorique d'arrivée et l'heure théorique de départ.

2.11. Respect de l'intervalle

Ce calcul repose sur la comparaison des intervalles réels réalisés et des intervalles théoriques prévus dans la table de référence sélectionnée. Cet indice est égal au nombre d'intervalles de classe C divisé par le nombre total d'intervalles.

Les intervalles de classe C sont ceux pour lesquels la précision pour le respect des intervalles est vérifiée soit :

$(\Delta \text{ départ réel} - \text{intervalle théorique}) < \text{précision pour le respect des intervalles}$

Avec $\Delta \text{ départ réel} = \text{heure de départ réel de la rame } i - \text{heure de départ de la rame } i - 1$

La précision pour le respect des intervalles est une entrée du système - comme la durée moyenne d'un trajet par un voyageur. Cette donnée sera demandée à l'opérateur.

2.12. Qualité production métro

Cet indice permet de connaître dans quelle mesure la ligne a respecté l'intervalle entre les rames et le temps de parcours soit (1- indice de retard de service). Il est calculé en pondérant l'indice mesurant le respect de l'intervalle et l'indice mesurant le temps de parcours. Cette pondération permet de favoriser un indice par rapport à l'autre et d'en déduire une satisfaction du service rendu en fonction des indices considérés comme les plus importants pour la quantification du service rendu.

Annexe 2. Rappels sur les sous ensembles flous

1. Définition d'un sous ensemble flou

Un sous ensemble flou A de X est défini par une fonction d'appartenance qui associe à chaque élément x de X , le degré $\mu_A(x)$, compris entre 0 et 1, avec lequel x appartient à A :

$$\mu_A: X \rightarrow [0,1]$$

2. Opérations sur les sous ensembles flous

Les opérateurs sur des sous ensembles flous génèrent de nouveaux sous ensembles flous caractérisés par des fonctions d'appartenance. Par conséquent, une opération sur des sous ensembles flous est réalisée par combinaison des fonctions d'appartenance les caractérisant.

2.1. Opérateur d'intersection

Une fonction $\wedge$ est un opérateur d'intersection si et seulement si :

$$\begin{aligned} \wedge: [0;1] \times [0;1] &\rightarrow [0;1], \\ \wedge &\text{ est non décroissante pour chaque argument,} \\ \wedge &\text{ est commutative,} \\ \wedge &\text{ est associative,} \\ \wedge &\text{ a 1 pour élément neutre} \end{aligned}$$

Les fonctions possédant ces propriétés sont appelées des normes triangulaires (T-norme).

Ainsi, l'intersection de deux sous ensembles flous A et B, d'un même univers de discours U définis par leurs fonctions d'appartenance $\mu_A(x)$ et $\mu_B(x)$, est caractérisée par la fonction d'appartenance $\mu_{A \cap B}$ définie par :

$$\begin{aligned}\mu_{A \cap B}: U &\rightarrow [0;1] \\ x &\rightarrow \mu_{A \cap B}(x) = \mu_A(x) \wedge \mu_B(x).\end{aligned}$$

Les T-normes sont aussi utilisées pour définir le produit cartésien. Ainsi, pour deux sous ensembles flous A et B des univers de discours U et V le produit cartésien $A \times B$ est caractérisé par la fonction d'appartenance $\mu_{A \times B}$:

$$\begin{aligned}\mu_{A \times B}: U \times V &\rightarrow [0;1] \\ (x, y) &\rightarrow \mu_{A \times B}(x, y) = \mu_A(x) \wedge \mu_B(y).\end{aligned}$$

2.2. Opérateur d'union

Une fonction $\vee$ est un opérateur d'union si et seulement si :

$$\begin{aligned}\vee: [0;1] \times [0;1] &\rightarrow [0;1], \\ \vee &\text{ est non décroissante pour chaque argument,} \\ \vee &\text{ est commutative,} \\ \vee &\text{ est associative,} \\ \vee &\text{ a 0 pour élément neutre}\end{aligned}$$

Les fonctions possédant ces propriétés sont appelées des conormes triangulaires (T-conorme).

Ainsi, l'union de deux sous ensembles flous A et B, d'un même univers de discours U définis par leurs fonctions d'appartenance $\mu_A(x)$ et $\mu_B(x)$ est caractérisée par la fonction d'appartenance $\mu_{A \cup B}$ définie par :

$$\begin{aligned}\mu_{A \cup B}: U &\rightarrow [0;1] \\ x &\rightarrow \mu_{A \cup B}(x) = \mu_A(x) \vee \mu_B(x).\end{aligned}$$

2.3. T-normes et T-conormes

Il existe un grand nombre de T-normes et de T-conormes, le tableau suivant donne les plus courantes.

T-normes	T-conormes	nom
$\min(x,y)$	$\max(x,y)$	Zadeh
$x.y$	$x+y-xy$	probabiliste
$\max(x+y-1,0)$	$\min(x+y,1)$	Lukasiewicz
$\frac{xy}{\gamma + (1-\gamma)(x+y-xy)}$	$\frac{x+y-xy-(1-\gamma)xy}{1-(1-\gamma)xy}$	Halmacher ($\gamma>0$)
$\begin{cases} x \text{ si } y = 1 \\ y \text{ si } x = 1 \\ 0 \text{ sinon} \end{cases}$	$\begin{cases} x \text{ si } y = 0 \\ y \text{ si } x = 0 \\ 1 \text{ sinon} \end{cases}$	Weber

Remarques :

La fonction min majore toutes les T-normes. Ainsi,

$$\forall (a,b) \in [0;1]^2, \forall \wedge \in \{\text{T-normes}\}; \min(a,b) \geq a \wedge b.$$

La fonction max minore toutes les T-conormes. Ainsi,

$$\forall (a,b) \in [0;1]^2, \forall \vee \in \{\text{T-conormes}\}; \max(a,b) \leq a \vee b.$$

On peut observer que les opérateurs d'union et d'intersection ne sont en fait que des fonctions d'interpolation permettant de relier les valeurs obtenues par logique booléenne. Par cet intermédiaire, on peut étendre les définitions de la logique classique à des valeurs réelles appartenant à l'intervalle [0;1].

2.4. Opérateur de négation

Une fonction $\neg$ est un opérateur de négation si et seulement si :

$$\begin{aligned} \neg: [0;1] &\rightarrow [0;1], \\ \neg &\text{ est continue,} \end{aligned}$$

$\neg$ non croissante,

$$\neg(0) = 1 \text{ et } \neg(1) = 0,$$

de plus, $\neg$ est une négation stricte ssi $\forall a \neg(\neg(a)) = a$.

L'opérateur de négation le plus utilisé en logique floue est :

$$\neg(a) = 1 - a$$

La figure suivante présente un exemple d'application de cet opérateur.

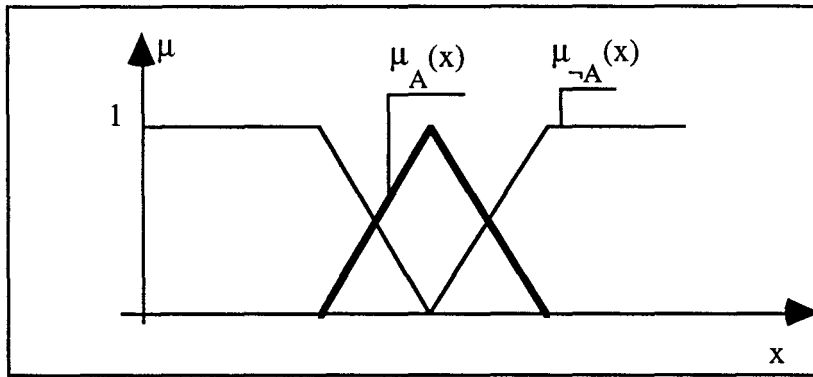


Figure A2.1. Exemple d'application de l'opérateur de négation

2.5. Projection d'un sous ensemble flou

Considérons un sous ensemble flou A d'un univers de discours $U \times V$, caractérisé par sa fonction d'appartenance $\mu_A(x, y)$. On désire établir l'influence de ce sous ensemble sur un univers de discours particulier U ou V . Pour ce faire, on utilise un opérateur de projection.

La projection d'un sous ensemble flou A d'un univers de discours $U \times V$, caractérisé par sa fonction d'appartenance $\mu_A(x, y)$, sur U est le sous ensemble flou $P_U(A)$ de U caractérisé par la fonction d'appartenance :

$$\mu_{P_U(A)}(x) = \sup_y (\mu_A(x, y)).$$

2.6. Règle floue

Une règle floue $R = A \rightarrow B$, $A \subset U$ et $B \subset V$, est un sous ensemble flou de l'univers de discours $U \times V$ caractérisé par la fonction d'appartenance μ_R définie à partir des fonctions d'appartenance $\mu_A(x)$ et $\mu_B(y)$ définissant les sous ensembles flous A et B :

$$\begin{aligned}\mu_R: U \times V &\rightarrow [0;1] \\ (x,y) &\rightarrow \mu_R(x,y) = \mu_{A \rightarrow B}(x,y) = f(\mu_A(x), \mu_B(y)).\end{aligned}$$

L'implication floue caractérisant une règle floue de la forme $R = A \rightarrow B$ est définie comme une extension de la définition de l'implication A implique B utilisée en logique booléenne. La définition de celle-ci est :

$$A \rightarrow B = (\neg A) \text{ OU } B.$$

Ainsi, la définition générale de l'implication floue est :

La fonction d'appartenance μ_R de la règle floue $R = A \rightarrow B$, $A \subset U$ et $B \subset V$ est :

$$\begin{aligned}\mu_R: U \times V &\rightarrow [0;1] \\ (x,y) &\rightarrow \neg(\mu_A(x)) \vee \mu_B(y).\end{aligned}$$

Dans cette équation, $\neg$ et $\vee$ sont des opérateurs de négation et d'union.

2.7. Inférence floue

Une règle floue du type "Si x est A alors y est B " détermine la relation liant des sous ensembles flous parfaitement définis A et B . Peut-on, en utilisant l'implication floue caractérisant cette règle, trouver le sous ensemble flou B' correspondant à un fait " x est A' " où A' est un sous ensemble flou différent de A ?

Le calcul de la fonction d'appartenance de B' est réalisé par l'intermédiaire du mécanisme d'inférence. Ce dernier tend à résoudre le problème si $A \rightarrow B$ alors $A' \rightarrow ?$. La description de ce problème est :

Fait flou : x est A' .
Règle : Si x est A alors y est B .
Conclusion : y est B' .

On veut trouver le sous ensemble flou B' de l'univers de discours V , caractérisé par sa fonction d'appartenance $\mu_{B'}$, à partir de la connaissance du sous ensemble flou A' de l'univers de discours U (caractérisé par $\mu_{A'}$) et de l'implication floue μ_R définie sur $U \times V$. Ce sous ensemble flou B' est la projection sur V du produit cartésien entre le sous ensemble flou A' et la relation floue R :

$$B' = P_V(A' \times R).$$

En reprenant les définitions du produit cartésien et de la projection floue, la fonction d'appartenance $\mu_{B'}(y)$ caractérisant la solution B' est :

$$\mu_{B'}(y) = \sup_y (\mu_{A'}(x) \wedge \mu_R(x, y)),$$

où l'opérateur $\wedge$ est une T-norme.

Exemple

Considérons les fonctions d'appartenance des sous ensembles flous A, A' et B connues et représentées sur la figure A2.2.

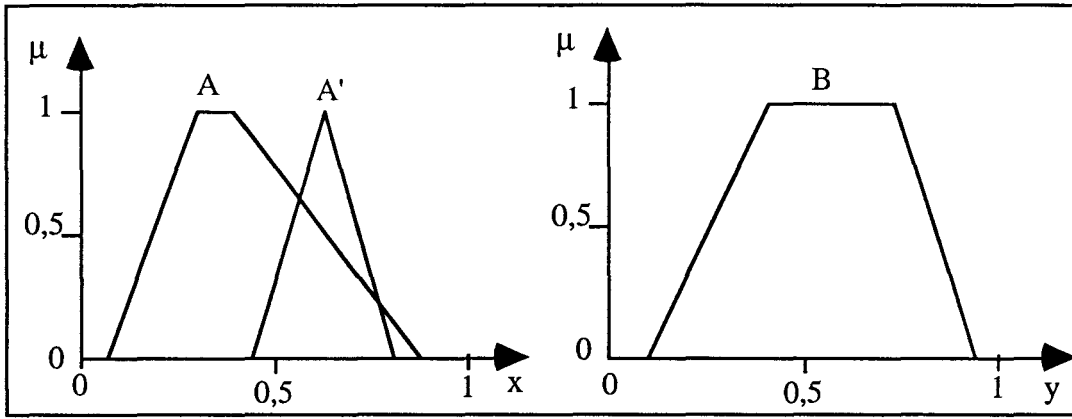


Figure A2.2. Représentation des fonctions d'appartenance de A, A' et B

Nous avons choisi comme opérateur le produit cartésien et l'implication la fonction minimum.

La fonction d'appartenance du sous ensemble flou B' est alors :

$$\begin{aligned} \mu_{B'}(y) &= \sup_x (\mu_{A'}(x) \wedge \mu_R(x, y)) \\ &= \sup_x (\min(\mu_{A'}(x), \min(\mu_A(x), \mu_B(y)))) \\ &= \min(\sup_x (\min(\mu_{A'}(x), \mu_A(x))), \mu_B(y)). \end{aligned}$$

On en déduit la construction graphique du sous ensemble flou B', présentée sur la figure suivante.

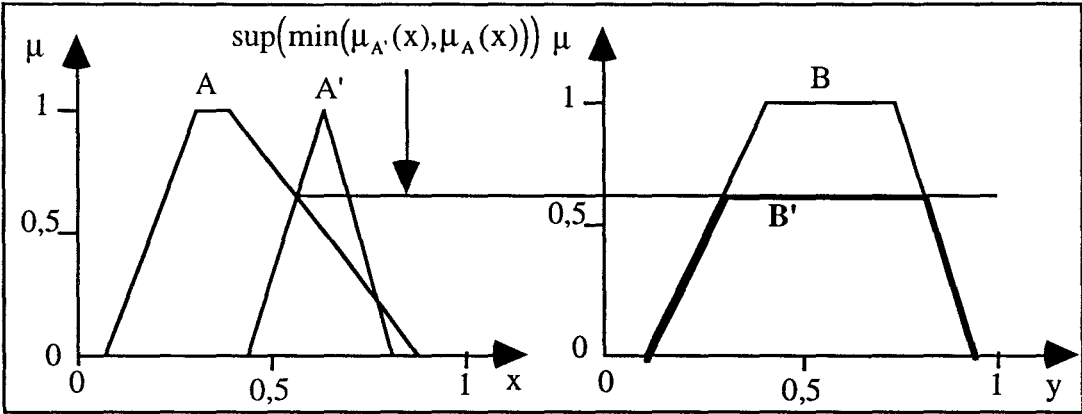


Figure A2.3. Construction du sous ensemble flou B'

2.7. Agrégation de règles floues

La description d'un système par un expert nécessite l'emploi de plusieurs règles floues R^i . Elles sont regroupées au sein d'une base de règles.

Règles R^i : Si x est A^i alors y est B^i ,

où A^i et B^i , $i=1$ à m , sont des sous ensembles flous parfaitement définis sur leurs univers de discours respectifs U et V .

Le problème posé est identique au précédent. En effet, à partir d'une information provenant du monde extérieur et de cette base de règle, on désire déterminer une décision B' .

Fait flou : x est A' .

m Règle R^i : Si x est A^i alors y est B^i .

Conclusion : y est B' .

Le plus souvent, la description sous forme de règles d'un système physique est disjonctive. C'est à dire que le système est décrit comme une union des cas envisagés dans les règles. Si la description du système est disjonctive, la relation globale R le décrivant est définie en utilisant un opérateur d'union sur les règles floues R^i .

$$R = R^i \vee \dots \vee R^m.$$

Ainsi, cette relation globale R se caractérise par la fonction d'appartenance μ_R :

$$\mu_R: U \times V \rightarrow [0;1]$$

$$(x, y) \rightarrow \mu_R(x, y) = \bigcup_{i=1}^m \mu_{R^i}(x, y) = \mu_{R^1}(x, y) \vee \dots \vee \mu_{R^m}(x, y).$$

A partir de cette relation globale R, on peut appliquer le mécanisme d'inférence. La fonction d'appartenance de la décision B' est alors :

$$\begin{aligned} \mu_{B'}(y) &= \sup_x (\mu_{A'}(x) \wedge \mu_R(x, y)) \\ &= \sup_x \left(\mu_{A'}(x) \wedge \left(\mu_{R^1}(x, y) \vee \dots \vee \mu_{R^m}(x, y) \right) \right), \end{aligned}$$

où $\wedge$ est une T-norme caractérisant l'opérateur du produit cartésien, $\vee$ est une T-conorme définissant l'opérateur d'union et la fonction sup est l'opérateur de la projection floue.

Annexe 3. Description de la plateforme de simulation

1. La programmation par objets

Les qualités attendues d'un programme sont la facilité à tester, à améliorer, à réutiliser et à maintenir. Ces qualités sont rarement entièrement vérifiées. Même si l'on ne peut attendre d'un programme qu'il satisfasse à tous ces besoins, il faut que l'environnement de développement du programme facilite la tâche du programmeur.

Au départ, les programmes étaient écrits en langage machine et ne pouvaient donc être réutilisés sur une machine différente. Puis, vinrent les langages comme FORTRAN qui permettaient au programmeur de s'abstraire de la machine. Après plusieurs mises à jour, les programmes obtenus devenaient de plus en plus difficiles à maintenir. Un programme était considéré comme un ensemble de procédures et un ensemble de données sur lesquels agissent les procédures. L'analyse d'une application consistait à la décomposer en sous-tâches indépendantes. Les problèmes de maintenance vinrent non pas de l'évolution des tâches de l'application mais de l'évolution des données. En effet, le moindre changement dans la structure des données provoquait des changements importants sur l'ensemble des procédures.

La particularité des langages à objets est de renfermer, dans une même entité, les procédures et les données, ces dernières sont ainsi protégées par une couche appelée interface. Un objet est donc un ensemble de procédures et un ensemble de données dont les liens avec le monde extérieur passent par l'interface (voir figure A3.1)

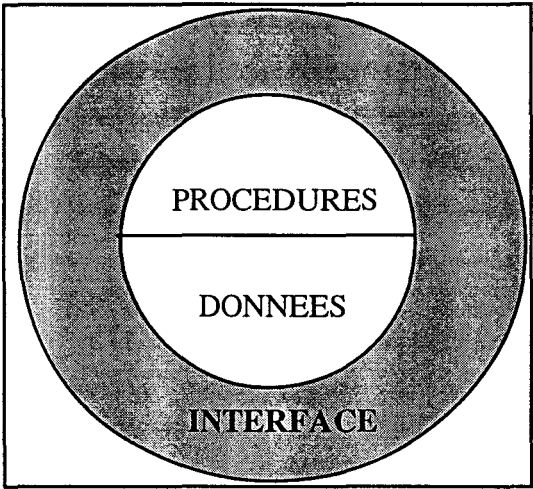


Figure A3.1. Le concept d'objet

Ainsi, pour utiliser les services d'un objet, on peut s'abstraire de la structure des données ou des détails d'implantation des procédures, l'important est de connaître l'interface. L'interface est constituée d'un ensemble de messages qui sont compréhensibles par l'objet. A chaque message correspond une procédure (le terme méthode est plus souvent employé). Lorsqu'un objet reçoit un message défini dans l'interface, il exécute la méthode qui lui correspond.

Les objets qui ont les mêmes méthodes et les mêmes structures de données sont décrits par un objet commun appelé classe. Une classe a la possibilité de créer sur le même moule autant d'objets semblables qu'on le souhaite. Les objets ainsi créés sont appelés instances de la classe. La figure A3.2 illustre l'envoi d'un message de création d'une rame à la classe Rame.

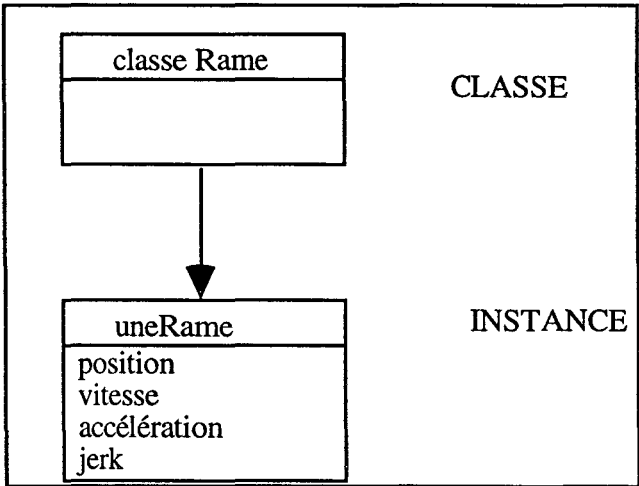


Figure A3.2. Création d'une instance de la classe Rame

Les données d'un objet sont accessibles, de façon interne, par des variables appelées variables d'instance. Par exemple, la rame "R07" est une instance de la classe Rame ; on définit dans la classe Rame une méthode qui modifie les variables d'instance pour arrêter la rame. Lorsqu'on envoie à "R07" le message "arrête", la méthode d'arrêt est exécutée sur l'instance "R07" jusqu'à ce que sa variable d'instance "vitesse" soit égale à 0 (figure A3.3)

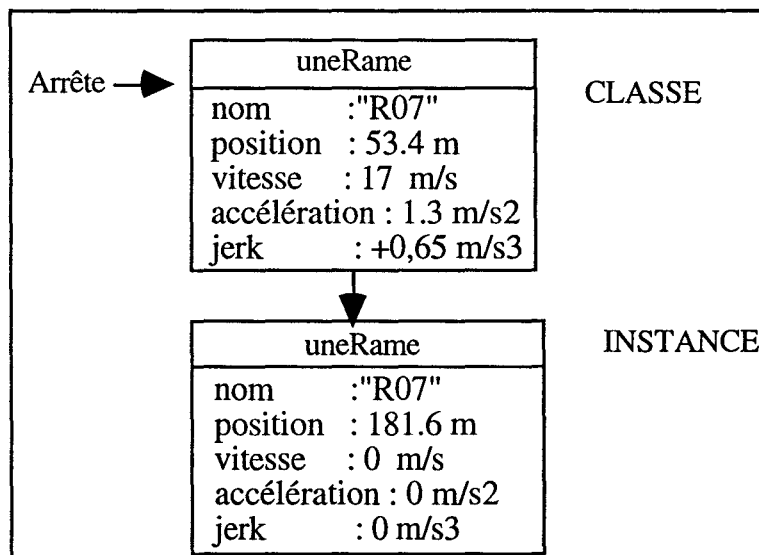


Figure A3.3. Envoi d'un message à une instance de classe

Les messages d'objets sont organisés en hiérarchie qui définit les relations d'héritage existant entre les classes. On appelle super-classe (respectivement sous-classe) d'une classe, la classe se trouvant au-dessus (en dessous) dans la hiérarchie d'héritage.

Les instances d'une classe héritent de toutes les variables d'instance et de toutes les méthodes définies dans les super-classes de la hiérarchie. C'est à dire qu'une instance peut comprendre des messages dont la méthode d'exécution a été définie dans une classe se trouvant dans le chemin qui mène à la racine de l'arbre d'héritage.

2. La simulation

On distingue trois types d'approches en simulation :

- Approche par événements,
- Approche par scrutation des activités ou par périodes,
- Approche par processus.

Dans l'approche par événements, la modélisation consiste à répertorier tous les événements susceptibles de provoquer un changement d'état des objets. La description de la dynamique du système désigne pour chaque événement, les événements ultérieurs qui lui sont liés. Ces descriptions sont appelées les logiques de changement d'état. Les événements issus des logiques de changement d'état sont stockés dans une pile. La simulation consiste à supprimer l'événement situé en haut de la pile, à affecter la variable temps courant à la date d'occurrence de l'événement puis à le traiter.

Dans l'approche par scrutation des activités, on définit une succession d'activités pour chaque objet. La description de la dynamique du système donne pour chaque objet les conditions nécessaires devant être satisfaites au début et à la fin de chaque activité. Une horloge scrute périodiquement chaque objet qui exécute la portion d'activité courante correspondant à la période de scrutation. L'ensemble des activités d'un objet définit son comportement vis à vis de son environnement.

Dans l'approche par processus, la modélisation consiste à assembler des modèles pré-programmés de modules du système. Ces modèles sont eux mêmes basés sur l'une des deux approches précédentes.

L'approche par scrutation a l'inconvénient de requérir de nombreux calculs si l'on veut suivre le système avec précision.

Un point faible de l'approche par événements est la définition du modèle et notamment des liens entre les événements.

Le choix de l'approche par scrutation des activités facilite la modélisation au détriment du coût des calculs nécessaire à la simulation.

La mise en oeuvre de l'approche par scrutation consiste à définir un objet similaire à une horloge. Cette horloge interpelle à une fréquence donnée l'ensemble des objets de simulation au comportement dynamique. Les objets interpellés doivent en réponse décider de l'action ou des actions qu'ils effectueront jusqu'au prochain top de l'horloge. Ce mode de simulation est appelé "Dead time simulation" car il y aura des pas de simulation où rien ne se passera. Mais ces temps morts ont l'avantage de laisser le processeur libre et de permettre ainsi à un utilisateur d'interagir pendant la simulation avec les objets de la simulation.

Les objets composant la simulation peuvent se regrouper en quatre catégories distinctes :

- Les rames étant définies par leurs caractéristiques physiques et cinématiques.

- Le réseau comprenant les aiguillages, les voies, les stations, les terminus et les balises de vitesse.
- La sécurité pouvant être le dispositif d'anticollision réalisé par cantons fixes ou mobiles déformables.
- La régulation.

3. Le mouvement des rames

Le terme "commande de mouvement" regroupe les différents signaux (ou consignes) qu'une rame reçoit et qui influenceront sa cinématique. Un dispositif de pilotage est un composant interchangeable d'une rame, il devra comprendre les commandes de mouvements reçues et contrôler la cinématique en conséquence. Il simulera aussi bien des automatismes ou un pilotage manuel.

Commandes de mouvement des rames

Le type de commande du mouvement des rames le plus répandu est une consigne de vitesse : la rame doit modifier sa vitesse pour satisfaire au plus tôt cette consigne. Un second type de commande est défini par une cible cinématique qui, en général, est un couple (position, vitesse). L'affectation d'une cible à une rame revient à lui demander d'atteindre le plus tôt possible une position avec une accélération et une vitesse données. L'affectation d'une cible est toujours associée à une vitesse limite qu'elle soit issue d'une commande ou liée à une contrainte physique de la rame.

Une instance de dispositif de pilotage est un algorithme de pilotage par le jerk (la dérivée de l'accélération) [CAN 90]. Celui-ci calcule à chaque pas de simulation une valeur du jerk permettant d'atteindre la cible le plus rapidement possible en respectant une vitesse limite admissible sur le tronçon où se trouve la rame. Le principe de cet algorithme consiste à anticiper sur un pas de simulation les conséquences d'un jerk. En fonction des conséquences obtenues, le choix du jerk se portera sur l'une des trois valeurs $-J_{\max}$, 0, $+J_{\max}$. Une fois la valeur du jerk établie les autres paramètres cinématiques sont calculés par intégration.

Les rames ne voient qu'une chaîne de cibles et un tableau de vitesses limites. Le tableau des vitesses limites comporte trois colonnes : les deux premières contiennent le début et la fin du segment d'élément de voie sur lequel porte la limitation de vitesse, la troisième colonne

est la valeur limite de la vitesse. Cette chaîne et ce tableau contiennent les cibles et les vitesses dites actives. La cible et la vitesse que doit prendre en compte une rame dépendent de sa position. A chaque mouvement de la rame, celle-ci doit vérifier que l'une ou l'autre partie de la commande n'a pas changé. S'il y a eu changement, il sera pris en considération pour le choix du jerk lors du prochain pas de simulation

Les chaînes de cibles et le tableau des vitesses limites joueront le rôle d'interface entre les objets qui représentent les automatismes de commande et les rames. Quel que soit le principe d'anticollision adopté par un automate commandant le mouvement des rames, les commandes sont traduites en opérations d'insertion ou de suppression de cibles et de vitesses limites (figure A3.4).

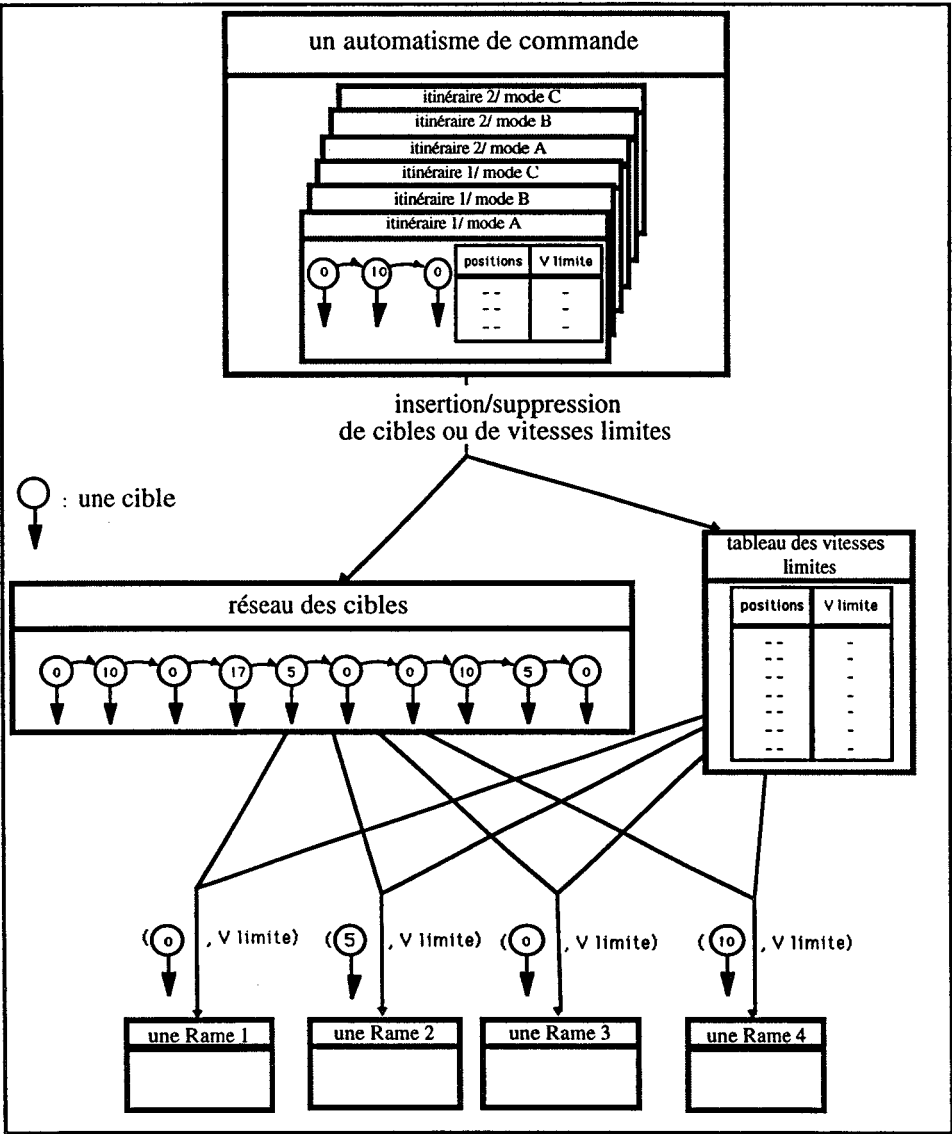


Figure A3.4. Commandes du mouvement des rames

Une instance d'automatisme chargée du mouvement des rames d'une zone contient un dictionnaire de tous les mouvements à commander. Dans le cas où la zone contrôlée est un canton fixe (cas du VAL de Lille), les mouvements sont décomposés par rapport à deux paramètres : l'itinéraire et le mode d'exploitation. Un itinéraire est défini par la liste des cantons dans l'ordre de parcours. La définition d'un mode d'exploitation est similaire à celle employée dans les automatismes du système VAL.

"mode normal" = vitesse maximale et arrêt aux quais,

"mode perturbé" = mode normal et arrêt avant la limite du canton,

"mode accostage" = vitesse constante très faible.

Dans le dictionnaire des mouvements qui peuvent être commandés, la valeur associée à une clé (itinéraire, mode d'exploitation) est une chaîne de cibles fixes et un tableau de vitesses limites qui couvrent la zone en charge par l'automatisme. Lorsqu'une instance d'automatisme choisit un mode d'exploitation pour un itinéraire établi, la chaîne de cibles et les vitesses limites correspondantes sont insérées respectivement dans le réseau des cibles et dans le tableau des vitesses limites activés dans la zone contrôlée.

La modélisation des automatismes à canton fixe à partir des commandes de mouvement des rames venant d'être définis se fait de la façon suivante : pour chaque canton, les courbes de consignes de vitesses correspondant aux différents modes d'exploitation et itinéraires sont traduites en cibles et vitesses limites. Le passage d'un mode à l'autre se fait par substitution des cibles et des vitesses limites correspondantes.

4. Interface utilisateur

Comme démontré précédemment, l'objectif de ces outils de simulation est de développer rapidement une simulation d'un système de transport guidé. Les simulations permettent d'extraire rapidement des caractéristiques d'exploitation (intervalles de passage, vitesse commerciale, influence des perturbations...). Il faut donc, par exemple, introduire n'importe quel type de perturbation à des moments opportuns. De même, il faut suivre visuellement le déroulement de la simulation pour demander des actions et détecter des défauts du système (mauvaise répartition du carrousel, effets de bord des automatismes, ...). Les besoins dans ce domaine sont très variés et dépendent des objectifs visés par l'utilisation de la simulation.

Une interface utilisateur présente des avantages significatifs sur trois points. Tout d'abord, elle permet à l'utilisateur de mieux comprendre le modèle de simulation. En second lieu, elle identifie les faiblesses du système simulé et ainsi elle trouve de meilleures solutions. En dernier lieu, l'utilisateur est complètement impliqué dans le déroulement d'une simulation en contrôlant certaines phases décisionnelles.

L'interface repose essentiellement sur la notion de triade MVC (Modèle-Vue-Contrôleur) qui est disponible dans le langage de programmation par objets Smalltalk 80. Une triade MVC est une architecture d'objets qui relie trois entités :

- Le modèle ("M") qui représente des données de l'application, ces données seront soit directement des objets de la simulation soit des objets intermédiaires avec les objets de la simulation.
- La vue ("V") qui représente une (ou plusieurs) fenêtre(s) de visualisation des données.
- Le contrôleur ("C") qui a pour tâche de gérer le dialogue entre l'utilisateur et les données au sein d'une vue.

La figure A3.5 montre l'interface utilisée pour permettre le dialogue avec la simulation. Cette interface permet de débiter une simulation, de l'arrêter, d'ouvrir une animation graphique de la ligne, d'ouvrir une interface avec les objets de simulation (figure A3.7), d'ouvrir des outils de suivi des objets de simulation, de modifier les heures de départs au terminus, d'ouvrir l'interface de la régulation (figure A3.6), de perturber les rames à chaque instant et d'ouvrir l'interface de l'estimateur du temps de panne (chapitre 4, paragraphe 5.2.4, figure 4.30).

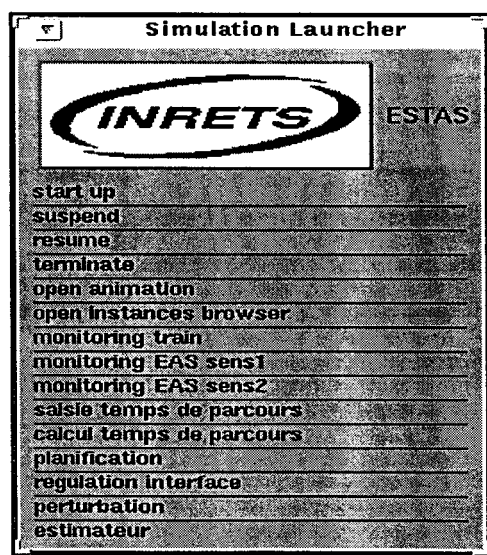


Figure A3.5. Interface avec la simulation

La figure suivante montre l'interface de la régulation. Cette interface permet de visualiser les horaires de départ aux terminus des deux sens, de choisir le type de régulation (régulation floue ou régulation type VAL) et de choisir le type de service que l'on désire (service partiel ou service ligne).

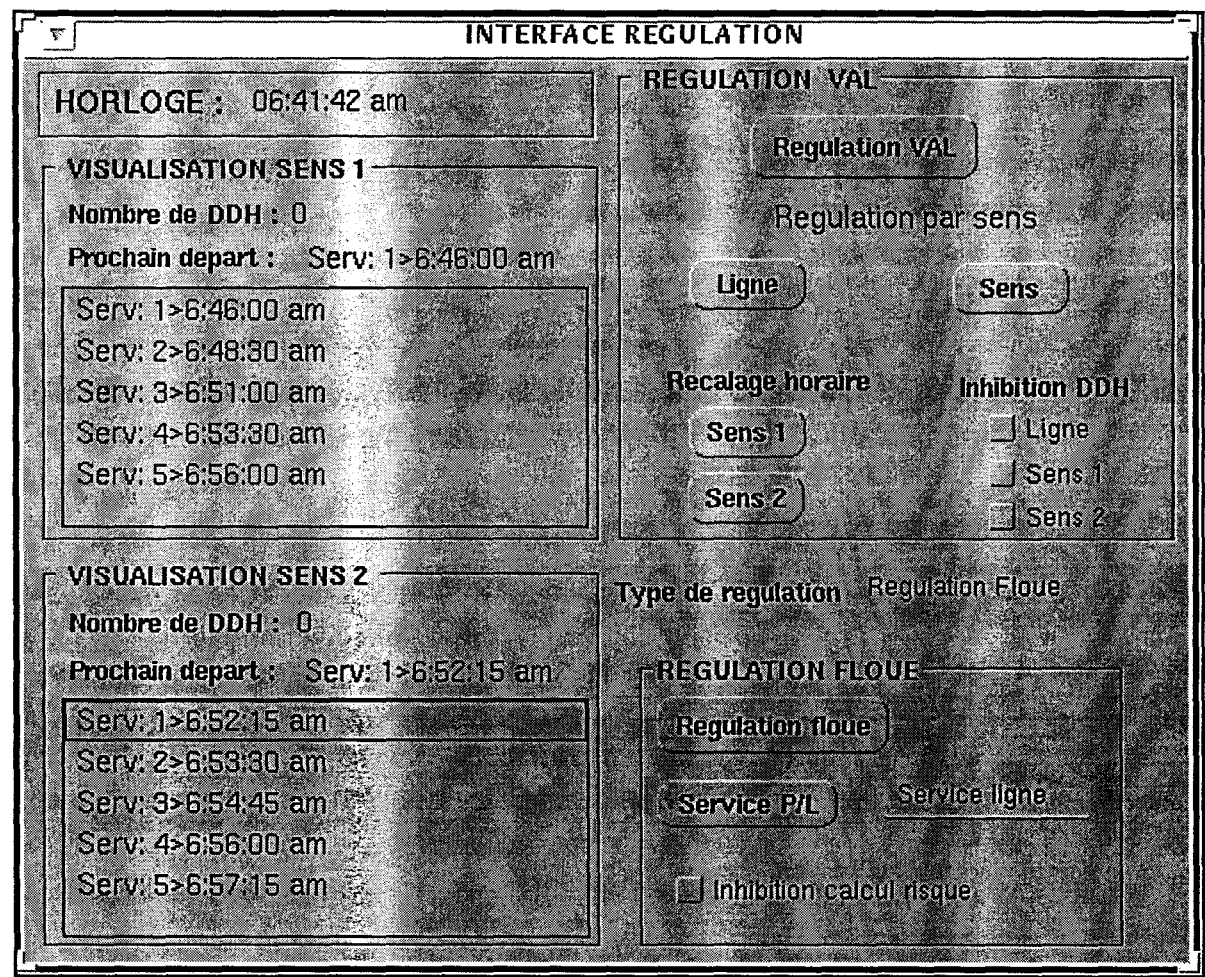


Figure A3.6. Interface de la régulation

Cette interface permet de réaliser des opérations spécifiques à chaque type de régulation. A savoir, pour la régulation type VAL, choisir le type de régulation (par sens ou par ligne), imposer des décalages horaires ou les inhiber ; pour la régulation floue, inhiber le calcul du risque effectué à chaque départ de station.

Interface avec les objets de simulation

L'interaction avec les objets de simulation est réalisée avec un "browser" d'instance. Un "browser" est une architecture d'objets disponibles dans Smalltalk ; nous avons appliqué

son principe à l'accès à n'importe quelle instance d'une liste de classes. Un "browser" d'instance est associé à deux types de fenêtres : des fenêtres qui permettent d'exprimer une requête et des fenêtres d'édition du résultat de la requête. Une requête s'exprime par une relation entre les attributs d'objets et des valeurs, elle permet de sélectionner l'objet qui vérifie ces relations.

Par exemple, la figure A3.7 représente le "browser" d'instance pour la ligne 1. Ici nous pointons la variable d'instance "position" de l'instance de la classe "Platform" représentant la station dont le nom est "Gam2".

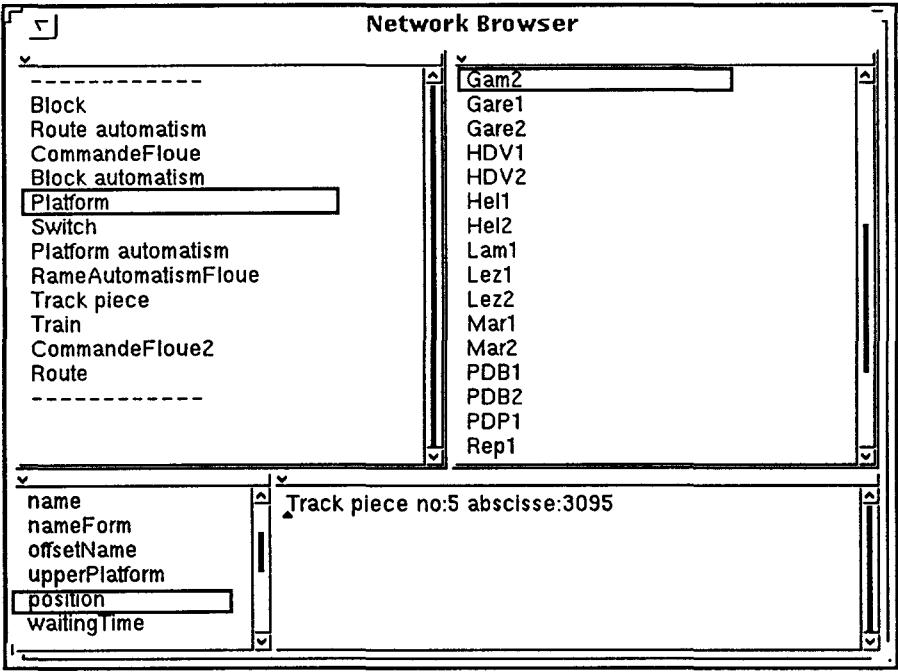


Figure A3.7. "Browser" d'instance pour la simulation de la ligne 1

