

Ministère de l'Enseignement Supérieur et de la Recherche Scientifique

Université Mouloud Mammeri de Tizi-Ouzou



Faculté de Génie Electrique et Informatique

Département Automatique

MEMOIRE DE MAGISTER

En Automatique

Option : Automatique des Systèmes Continus et Productique

**Thème : Approche ensembliste en automatique :
Application à l'estimation et la commande**

Présenté par : M^{elle} Khelouat Lila

Mémoire soutenu le : 08 juillet 2012 devant le jury d'examen composé de:

Président : BENFDILA Arezki Professeur, UMMTO.

Rapporteur : DJENNOUNE Saïd Professeur, UMMTO.

Examineur : KARA Redouane Maître de Conférences 'A', UMMTO.

Examineur : HAMMOUCHE Kamal Maître de Conférences 'A', UMMTO.

Remerciements

Le travail présenté dans ce mémoire a été dirigé sous la direction de Monsieur S.Djennoune que je remercie pour son encadrement, son aide, ses directives et ses précieux conseils.

Je remercie les membres du jury qui m'ont fait l'honneur de participer à l'examen de ce travail.

Je remercie toute personne qui m'a apporté son aide, ne serait-ce que minime.

Sommaire

Notations	v
Liste des figures	viii
Liste des tableaux	xiii

Introduction générale	1
------------------------------------	---

Chapitre 1 : Approche ensembliste

1.1 Introduction.....	4
1.2 Motivations.....	5
1.3 Les fondements de l'approche ensembliste.....	7
1.3.1 Ensembles.....	7
1.3.2 Opérations sur des ensembles réels.....	8
a) Opérations ensemblistes pures.....	9
b) Opérations étendues.....	10
1.3.3 Distance entre deux ensembles compacts.....	12
1.4 Représentation des sous ensembles.....	13
1.4.1 Représentation par un nuage de points	13
1.4.2 Représentation par équation et inéquation	14
1.4.3 Représentation par recouvrement	14
a) Les zonotopes	15
b) Les pavés	16
1.4.4 Représentation par encadrement	16
1.5 Conclusion.....	16

Chapitre 2 : Arithmétique des intervalles

2.1 Introduction.....	17
2.2 Historique.....	18
2.3 Notions sur l'arithmétique des intervalles.....	19
2.3.1 Les intervalles.....	19
2.3.2 Les pavés.....	21
2.3.3 Le pessimisme.....	22
a) Le phénomène de dépendance.....	22
b) Le phénomène d'enveloppement.....	23
2.3.4 Fonctions d'inclusion.....	24

a)	Les différentes fonctions d'inclusion	25
a.1)	Fonction d'inclusion des fonctions élémentaires.....	25
a.2)	Fonction d'inclusion naturelle.....	26
a.3)	Fonction d'inclusion moyenne.....	27
a.4)	Fonction d'inclusion de Taylor.....	27
b)	Propriétés des fonctions d'inclusion.....	28
b.1)	Monotonie d'une fonction d'inclusion.....	28
b.2)	Convergence d'une fonction d'inclusion.....	28
b.3)	Ordre de convergence	28
2.3.5	Partitionnement de pavé	30
2.4	Algorithmes ensemblistes.....	32
2.4.1	Inversion ensembliste par l'arithmétique des intervalles.....	32
2.4.2	Les contracteurs.....	34
a)	Projection de contraintes.....	35
b)	Quelques algorithmes de projection de contraintes.....	35
c)	Propagation de contraintes.....	36
c.1)	Principe de la technique.....	36
c.2)	Problème de consistance locale.....	37
c.3)	Solutions apportées au problème de consistance locale.....	38
d)	Infaisabilité d'un problème de satisfaction de contraintes.....	39
e)	Algorithme de propagation-rétropropagation.....	39
f)	Quelques applications des contracteurs.....	40
f.1)	Exemple d'application à la stabilité robuste	41
f.2)	Application à l'estimation à erreur bornée.....	43
2.5	Résolution des systèmes algébriques linéaires à intervalles.....	45
2.5.1	Définition d'un système algébrique linéaire à intervalles.....	46
2.5.2	Méthodes de résolution des systèmes algébriques linéaires à intervalles.....	47
a)	Méthode de Krawczyk.....	47
b)	Méthode de Hansen-Blier-Rohn-Ning-Kearfott-Neumaier.....	48
c)	Résolution des systèmes algébriques linéaires à intervalles par la fonction « Verifylss » de la toolbox INTLAB	49
2.6	Conclusion.....	50

Chapitre3 : Estimation d'état par l'arithmétique d'intervalles

3.1 Introduction.....	52
3.2 Revue sur les concepts de base de la théorie classique de l'estimation des systèmes linéaires temps invariant	54
3.2.1 Commandabilité.....	54
3.2.2 Observabilité.....	55
3.2.3 Définition d'un observateur.....	55
3.2.4 Méthodes classiques d'estimation.....	55
a) Observateur de Luenberger.....	56
b) Filtre de Kalman continu.....	57
3.3 Test de commandabilité et d'observabilité robuste des systèmes linéaires temps invariant à paramètres incertains par l'arithmétique d'intervalles.....	59
3.3.1 Dépendance et indépendance linéaires des vecteurs intervalles.....	61
3.3.2 Application au test de commandabilité robuste et d'observabilité robuste.....	64
3.4 Estimation d'état des systèmes linéaires temps invariant incertains par l'arithmétique des intervalles.....	73
3.4.1 Observateur intervalle pour les systèmes linéaires temps invariant à paramètres incertains.....	73
a) Application au modèle d'un réacteur chimique continu.....	75
3.4.2 Observateur intervalle pour les systèmes linéaires temps invariant en présence des perturbations.....	86
a) Exemple d'application sur un système d'ordre trois	87
b) Etapes de construction de l'observateur intervalle temps-variant.....	88
b.1) Changement de coordonnées temps-invariant par la forme canonique de Jordan.....	88
b.2) Le changement de coordonnées temps-variant.....	90
b.3) L'observateur intervalle	91
3.5 Conclusion.....	95

Chapitre 4 : Application au pendule inversé

4.1 Introduction.....	96
4.2 Modélisation du pendule inversé.....	97
4.2.1 Equations du mouvement du système.....	97
4.2.2 Linéarisation des équations.....	101
4.2.3 Représentation dans l'espace d'état.....	102

4.3 Synthèse d'un observateur Luenberger classique pour le pendule inversé en absence des perturbations en sortie.....	103
4.4 Estimation d'état en présence d'incertitudes sur le système.....	110
4.4.1 Synthèse d'un observateur Luenberger classique en tenant compte des perturbations en sortie et en absence d'incertitudes.....	112
4.4.2 Synthèse d'un observateur intervalle pour l'erreur dynamique du système.....	115
a) Transformation sous forme canonique de Jordan	116
b) Changement de coordonnées temps-variant.....	116
c) Observateur intervalle pour l'erreur dynamique.....	117
4.4.3 Observateur intervalle du système.....	120
4.5 Commande par retour d'état du pendule inversé.....	125
4.5.1 Description de la méthode.....	125
4.5.2 Commande par retour d'état reconstruit par l'observateur de Luenberger et l'observateur intervalle.....	128
a) Commande par retour d'état reconstruit par l'observateur de Luenberger en absence d'incertitudes	128
b) Commande par retour d'état en présences d'incertitudes.....	130
b.1) L'observateur intervalle de l'erreur dynamique en boucle fermée...	130
b.2) Limites supérieure et inférieure pour l'état corrigé	132
4.6 Conclusion.....	137
Conclusion générale et perspectives...	138
Bibliographie	140
Annexe	143

Notations

$\triangleq$: égale par définition.

$\coloneqq$: opération d'affectation.

$\neg$: complémentation.

R : ensemble des nombres réels.

R^{n_x} : espace produit de R dont les éléments sont des n-uplets $(x = (x_1, x_2, \dots, x_{n_x})^T$ avec $x_i \in R$).

$f : R^{n_x} \rightarrow R^{n_y}$: fonction vectorielle.

X, Y : deux ensembles dans R^{n_x} et R^{n_y} respectivement.

$f(X)$: image de X par f .

$f^{-1}(Y)$: image inverse de Y par f .

$h_\infty^0(X, Y)$: proximité d'un ensemble compact X à l'ensemble compact Y .

$h_\infty(X, Y)$: distance de Hausdorff entre deux ensembles compact.

$[\cdot]$: crochets utilisés pour représenter des intervalles et vecteurs intervalles (Pavés).

$[a] = [\underline{a}, \overline{a}]$: intervalle de borne inférieure $\underline{a}$ et de borne supérieure $\overline{a}$ dans R .

IR : ensemble des intervalles de R .

$[a] = ([\underline{a}_1, \overline{a}_1], \dots, [\underline{a}_n, \overline{a}_n])^T$ de R^n : vecteur intervalle.

Comme nous mettons x^I et X^I respectivement pour représenter un vecteur intervalle et matrice intervalle.

IR^n : ensemble des intervalles de R^n .

$[f]$: fonction d'inclusion.

$[f]^*$: fonction d'inclusion minimale.

$[f_n]$: fonction d'inclusion naturelle.

$[f_m]$: fonction d'inclusion moyenne.

$[x] = \bigcup_{i=1}^n [x_i]$: sous pavage d'un pavé $[x]$.

C_X : contracteur pour l'ensemble X .

CSP : problème de satisfaction de contraintes (Constraint Satisfaction Problem).

$\sqcap$: intersection répétée.

ICP : propagation de contraintes intervalles (Interval Constraint Propagation).

$C_{S_1}^*$: contracteur minimal de l'ensemble S_1 .

A_{ij} : éléments de la matrice A .

$\langle A \rangle$: matrice de comparaison.

$\Sigma(A, b)$: ensemble solution du système algébrique linéaire à intervalle ($Ax = b$) (solution réelle).

$\cup \Sigma(A, b)$: pavé qui est l'union de tous les ensembles solutions du système algébrique linéaire à intervalle $Ax = b$.

LTI : linéaire temps invariant.

$\mathcal{C}$: matrice de commandabilité.

$\mathcal{O}$: matrice d'observabilité.

L : gain d'observateur.

x : état réel

$\hat{x}$: état estimé.

$\dot{e}$: erreur dynamique d'un système.

$X^I \in IR^{m \times n}$: matrice intervalle composée de n vecteur intervalles x_i^I de IR^m .

X_0 : matrice centrée de X^I .

ΔX : matrice rayon de X^I .

μ_{xy} : rapport entre deux vecteurs intervalles $x^I, y^I \in IR^m$.

S^i : sous-matrice carrée de dimension $(n \times n)$ construite à partir de la matrice intervalle X^I .

k : le nombre des sous-matrices carrées S^i construites à partir de la matrice X^I .

S_X : ensemble des sous-matrices carrées S^i .

s^I : ensemble indice des sous-matrices carrées S^i

S_{Xc} : ensemble des matrices centrée S_0^i des sous-matrices carrées S^i de X^I .

ΔS_X : ensemble des matrices rayon ΔS^i des sous-matrices carrées S^i de X^I .

$\rho(A)$: rayon spectrale de la matrice carrée A , qui est égale au plus grand module des valeurs propres de la matrice A .

$x^-(t) \leq x(t) \leq x^+(t)$: $x(t)$ est l'état réel enfermé par l'estimation d'état supérieure $x^+(t)$ et l'estimation d'état inférieure $x^-(t)$ (x^+ et x^- : est l'observateur intervalle de $x(t)$).

$\theta^- \leq \theta \leq \theta^+$: θ est un paramètre incertain borné par une borne supérieure θ^+ et une borne inférieure θ^- .

σ : facteur pondéré.

$w(t)$: est un bruit (où une perturbation) qui est une fonction Lipschitz continue inconnue.

Γ : est un paramètre variable permettant de considérer plusieurs bornes inférieure et supérieure pour les fonctions w^- et w^+ respectivement qui sont les bornes de $w(t)$.

$J = P A P^{-1}$: forme canonique de Jordan de A par la transformation de base P .

$P(t)$: matrice temps variant.

$P^+(t) : \max(0, P(t))$

$M(t)$: matrice temps variant qui est l'inverse de $P(t)$.

$M^+(t) : \max(0, M(t))$

$\mathcal{L}_g$: Lagrangien d'un système

E_c : énergie cinétique d'un système

E_p : énergie potentielle d'un système.

ξ : coordonnées généralisées d'un système

$\varphi(t)$: perturbations en sortie du système LTI qui est une fonction Lipschitz continue inconnue.

K : gain du retour d'état

Liste des figures

Figure 1.1 : Opérations ensemblistes pures.....	9
Figure 1.2 : Fonction f reliant les ensembles A et B	11
Figure 1.3 : Distance de Hausdorff entre deux ensembles compacts X et Y	13
Figure 1.4 : Exemple d'un zonotope.....	15
Figure 2.1 : Représentation du pavé $[Y]$ et de l'ensemble réel y	24
Figure 2.2 : Evaluation des fonctions $f([x])$, $[f]([x_1])$, $[f]([x_2])$ et $[f]([x])$	31
Figure 2.3 : Caractérisation de X par l'algorithme SIVIA.....	33
Figure 2.4 : Système corrigé en boucle fermée.....	41
Figure 2.5 : Circuit électrique comprenant une pile et deux résistances en séries.....	43
Figure 2.6 : Exemple montrant l'encadrement extérieur de la solution réelle $\Sigma(A, b)$ par l'ensemble $\cup \Sigma(A, b)$ qui est un vecteur intervalle.....	47
Figure 3.1 : Principe d'estimation d'état.....	56
Figure 3.2 : Estimation de l'état réel $x(t)$ (en bleu) par l'observateur intervalle (en vert z^- et en rouge z^+) et par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour un choix de gain $L = [0 \ 0 \ 0]^T$	78
Figure 3.3 : Estimation de l'état réel $x(t)$ (en bleu) par l'observateur intervalle (en vert z^- et en rouge z^+) et par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour un choix de gain $L = [0 \ 0 \ 10]^T$	79
Figure 3.4 : Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 4.5$	80
Figure 3.5 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 4.5$	81
Figure 3.6 : Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 5.5$	81
Figure 3.7 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 5.5$	82
Figure 3.8 : Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 6$	82
Figure 3.9 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 6$	83

Figure 3.10: Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 7$.	83
Figure 3.11 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 7$.	84
Figure 3.12 : Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 7.5$.	84
Figure 3.13 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 7.5$.	85
Figure 3.14 : Estimation de l'état réel (en bleu) par un observateur intervalle (l'observateur supérieur en rouge et l'observateur inférieure en violet) et pour $w^+ = w(t) + 0.2$, $w^- = w(t) - 0.2$.	93
Figure 3.15 : Estimation de l'état réel (en bleu) par un observateur intervalle (l'observateur supérieur en rouge et l'observateur inférieure en violet) et pour $w^+ = w(t) + 0.1$, $w^- = w(t) - 0.1$.	94
Figure 3.16 : Estimation de l'état réel (en bleu) par un observateur intervalle (l'observateur supérieur en rouge et l'observateur inférieure en violet), pour $w^+ = w(t) + 0$, $w^- = w(t) - 0$.	94
Figure 4.1: Schéma de l'ensemble chariot-pendule.	98
Figure 4.2 : Schéma de simulation sous SIMULINK du système réel doté d'un observateur Luenberger classique.	106
Figure 4.3 : Réponse en position du chariot, état réel (en rouge) et état estimé par l'observateur classique de Luenberger (en bleu).	106
Figure 4.4 : Réponse en vitesse du chariot, état réel (en rouge) et état estimé (en bleu) par l'observateur classique de Luenberger.	107
Figure 4.5 : Réponse angulaire du pendule, état réel (en rouge) et état estimé (en bleu) par l'observateur classique de Luenberger.	107
Figure 4.6 : Réponse en vitesse angulaire du pendule, état réel (en rouge) et état estimé (en bleu) par l'observateur classique de Luenberger.	108
Figure 4.7 : Evolution de l'erreur dynamique pour les quatre variables d'état du pendule inversé.	109
Figure 4.8 : Estimation de la position du chariot par l'observateur Luenberger classique en présence des perturbations sur la mesure.	112

Figure 4.9 : Estimation de la vitesse du chariot par l'observateur Luenberger classique en présence des perturbations sur la mesure.	113
Figure 4.10 : Estimation de l'angle du pendule par l'observateur Luenberger classique en présence des perturbations sur la mesure.	113
Figure 4.11 : Estimation de la vitesse angulaire du pendule par l'observateur Luenberger classique en présence des perturbations sur la mesure.	113
Figure 4.12 : Evolution de l'erreur dynamique du système perturbé.....	114
Figure 4.13: Limites supérieures et inférieures de l'erreur dynamique obtenues par un observateur intervalle pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$	119
Figure 4.14: Limites supérieures et inférieures de l'erreur dynamique obtenues par un observateur intervalle pour $w^+ = w(t) + 0.01$ et $w^- = w(t) - 0.01$	120
Figure 4.15 : Estimation d'état de la position du chariot (x_1), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par l'observateur intervalle en vert et en violet pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$	121
Figure 4.16 : Estimation d'état de la vitesse du chariot (x_2), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet et pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$	121
Figure 4.17: Estimation d'état de l'angle du pendule (x_3), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et violet, pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$	122
Figure 4.18 : Estimation d'état de la vitesse angulaire du pendule (x_4), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet, pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$	122
Figure 4.19: Estimation d'état de la position du chariot (x_1), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet, pour $w^+ = w(t) + 0.01$, $w^- = w(t) - 0.01$	123
Figure 4.20 : Estimation d'état de la vitesse du chariot (x_2), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par l'observateur intervalle en vert et en violet, pour $w^+ = w(t) + 0.01$ et $w^- = w(t) - 0.01$	123
Figure 4.21: Estimation d'état de l'angle du pendule (x_3), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet et pour $w^+ = w(t) + 0.01$ et $w^- = w(t) - 0.01$	124

Figure 4.22: Estimation d'état de la vitesse angulaire du pendule (x_4), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet et pour $w^+ = w(t) + 0.01$ et $w^- = w(t) - 0.01$	124
Figure 4.23 : Réponses du système réel bouclé en rouge, l'état corrigé reconstruit par l'observateur de Luenberger en bleu. a) et b) Représentent la position et la vitesse du chariot respectivement, c) et d) représentent la position angulaire et la vitesse angulaire du pendule respectivement.....	129
Figure 4.24 : Evolution de l'erreur dynamique en boucle fermée.....	130
Figure 4.25 : Réponse de la position du chariot par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.005$	132
Figure 4.26 : Réponse de la vitesse du chariot par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.005$	133
Figure 4.27 : Réponse de la position angulaire du pendule par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.005$	133
Figure 4.28 : Réponse de la vitesse angulaire du pendule par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.005$	133
Figure 4. 29: Limites supérieure et inférieure de l'erreur en boucle fermée obtenues par l'observateur intervalle en rouge et en violet enfermant l'erreur associée à l'observateur de Luenberger classique en bleu pour une valeur de Γ égale à 0.005.....	134
Figure 4.30: Réponse de la position du chariot par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.002$	134
Figure 4.31: Réponse de la vitesse du chariot par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.002$	134
Figure 4.32 : Réponse de la position angulaire du pendule par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.002$	134

- Figure 4.33 :** Réponse de la vitesse angulaire du pendule par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.002$134
- Figure 4.34:** Limites supérieure et inférieure de l'erreur en boucle fermée obtenues par l'observateur intervalle en rouge et en violet enfermant l'erreur associée à l'observateur de Luenberger classique en bleu pour une valeur de Γ égale à 0.002.....136

Liste des tableaux

Tableau 2.1 : Comparatif des fonctions d'inclusion de f sur deux pavés différents $[x_1]$ et $[x_2]$	29
Tableau 2.2 : Nombre des pavés obtenus pour trois différentes précisions.....	34
Tableau 2.3 : Table de Routh.....	42
Tableau 2.4 : Solution du système linéaire à intervalle.....	50
Tableau 3.1 : Résultats du test de la commandabilité et de la non-commandabilité par variation de α	71
Tableau 3.2 : Valeurs des paramètres du système pour les simulations.....	77

Introduction générale

Introduction générale

Le traitement des données joue un rôle important dans la conception et la maîtrise de notre environnement. En pratique, ces données sont soit directement accessibles à la mesure, soit estimées grâce à des algorithmes d'estimation. Généralement ces données sont sujettes à des incertitudes qui sont de nature déterministes ou aléatoires.

Ces incertitudes inévitables sur les données dans un système physique peuvent provenir de la chaîne d'instrumentation, dont le problème est lié principalement à la qualité des capteurs, aux conditions opératoires et à la précision des signaux parasites. C'est pour cela qu'on cherche à caractériser plus explicitement ces erreurs de mesure influant sur le système.

Généralement, pour modéliser l'évaluation de ces incertitudes, on fait appel à des méthodes statistiques mettant en jeu des lois probabilistes. Une variable incertaine dans une approche probabiliste est représentée par la densité de probabilité qui nous donne un support qui contient à coup sûr la variable incertaine, mais qui exige une information sur la manière dont elle est distribuée sur ce support. La difficulté rencontrée par l'utilisation de cette approche est dans la manipulation des calculs qui exige souvent de poser des hypothèses simplificatrices des lois conjointes des variables incertaines, ce qui n'est pas le cas en pratique. Alors le rajout de telles hypothèses produit des résultats difficiles à interpréter.

En revanche, une reformulation ensembliste des incertitudes garantit une manipulation aisée des calculs et ne réclame aucune information statistique sur les variables incertaines. Il est, fréquemment, que l'erreur de mesure sur une donnée expérimentale soit donnée en terme de tolérance par l'appareil de mesure et non pas en terme de densité de probabilité, alors l'approche ensembliste est une approche naturelle pour traiter ces mesures incertaines.

La majorité des méthodes ensemblistes sont basées sur l'analyse par intervalles. Initialement développée pour quantifier les erreurs dues à la représentation finie des nombres réels par un calculateur [Moore, 1962], [Moore, 1966], l'analyse par intervalles est devenue un outil de résolution des problèmes mathématiques très complexes, tels que : la résolution des systèmes d'équations linéaires et non linéaires à intervalles [Neumaier, 1991], [Neumaier, 1999] et les problèmes d'optimisation globale [Hansen, 1992], [Schichl et Neumaier, 2005].

Le principal attrait qu'offre cet outil réside dans sa capacité à appréhender les incertitudes sur un modèle donné. Alors celles-ci sont modélisées par des intervalles ou vecteurs intervalles, offrant ainsi des bornes finies qui nous garantissent de les contenir.

L'utilisation de l'analyse par intervalles dans le domaine de l'automatique est très récent, Dans le domaine de l'estimation d'état qui est le cadre de notre travail, l'objectif est de déterminer une estimation d'état fiable pour les états du système incertain. Les incertitudes qui interviennent généralement sont soit associées aux paramètres du modèle d'état, soit aux signaux d'entrées qui sont mal connus ou bien aux conditions initiales des variables d'état du système qui sont incertaines.

Grace à cet outil, nous pouvons concevoir des observateurs qu'on nomme *observateurs intervalles*, qui fournissent des limites à tout instant et qui garantissent de renfermer l'évolution des états réels du système incertain par deux trajectoires distinctes. C'est pour cela, que cette technique devient de plus en plus populaire et a intéressé plusieurs chercheurs. Nous citons par exemple [Videau, 2005] dans la surveillance et la détection d'anomalie des systèmes linéaire temps continu, qui a développé des méthodes garanties où l'évolution des variables d'état est représentée par des enveloppes par la technique des intervalles. Et la contribution de [Moisan, 2007] pour la synthèse d'observateurs intervalles pour des systèmes biologiques caractérisés par des incertitudes.

Le problème d'estimation d'état est un problème majeur en automatique pour plusieurs raisons. D'une part, dans la conception et la réalisation d'une commande par retour d'état, l'utilisation d'un observateur ou d'un estimateur d'état est inévitable lorsque les variables d'états sont inaccessibles à la mesure. D'autre part, l'observateur d'état joue un rôle important dans la résolution des problèmes de diagnostic des fautes et de maintenance des procédés industriels.

Les objectifs fixés dans notre travail sont les suivants:

- Etude de la théorie ensembliste et de l'arithmétique des intervalles.
- Applications à certains problèmes d'automatique, en particulier pour la synthèse d'observateurs intervalles pour la classe fondamentale des systèmes linéaires temps invariant incertains.
- Application sur un pendule inversé: synthèse d'un estimateur de la position, de la vitesse du chariot et de la position angulaire et de la vitesse angulaire du pendule.

Le mémoire est organisé de la manière suivante :

Dans le premier chapitre, nous allons aborder les fondements de base de l'approche ensembliste. Motivée par sa simplicité dans la manipulation et dans les calculs, la représentation par des intervalles (ou vecteurs intervalles) est choisie pour modéliser les incertitudes influant sur un système donné.

Le deuxième chapitre est consacré à exposer cette arithmétique des intervalles. Des exemples d'outils faisant intervenir l'analyse par intervalles dans leur méthode sont illustrés, tels que l'inversion ensembliste par arithmétique d'intervalles et la caractérisation d'ensembles par contraction, qui sont généralement utilisés pour la résolution des systèmes d'équations linéaires et non linéaires. Des exemples sur la stabilité robuste et l'estimation à erreur bornée faisant intervenir l'outil de contraction sont présentés.

Dans le troisième chapitre, nous allons développer deux méthodes faisant intervenir l'arithmétique d'intervalles pour l'estimation d'état. La première méthode considère des incertitudes paramétriques dans la représentation d'état du système. Cette méthode est appliquée au cas des systèmes linéaires temps invariant qui doivent vérifier une hypothèse restreinte qui est la condition de coopérativité. Un exemple d'application est illustré sur un modèle incertain d'un réacteur chimique continu. La deuxième méthode, fait face à cette

condition de coopérativité exigée pour l'estimation d'état par l'arithmétique des intervalles, et ceci par l'utilisation de certains changements de coordonnées qui transforment le système en un système coopératif. Cependant, elle ne concerne qu'une classe des systèmes linéaires temps invariant présentant des perturbations additives mal connues. Un exemple montrant toutes ces étapes de transformation sont entreprises sur un système d'ordre trois.

Le dernier chapitre est consacré à mettre en application la deuxième méthode développée dans le troisième chapitre, pour l'estimation d'état d'un modèle du pendule inversé en boucle ouverte, en présence de perturbations additives sur les mesures de sortie et sur les conditions initiales des variables d'état du système. La méthode est combinée à un observateur classique de Luenberger [Luenberger, 1966]. Par la suite, une commande par retour d'état est développée celle-ci est basée sur l'état reconstruit par l'observateur classique de Luenberger et l'observateur intervalle.

Chapitre1:

Approche ensembliste

Chapitre 1 :

Approche ensembliste

1.1 Introduction

De nombreuses méthodes en mathématique ou dans les sciences pour l'ingénieur ont été développées pour résoudre des problèmes dont les données qui interviennent sont ponctuelles et dont la solution est aussi ponctuelle, ce type de problème est appelé « *problème ponctuel* ».

Généralement les données qui interviennent dans l'énoncé d'un problème sont entachées d'incertitudes, introduites par exemple par le calculateur lors de la résolution (incertitudes numériques) ou tout simplement par l'environnement expérimental (appareils de mesure, bruit et perturbations sur les mesures...). À cet effet les méthodes développées pour la résolution des problèmes ponctuels ne permettent pas (sauf pour une classe très réduite de problèmes) de trouver de façon garantie un résultat global aussi proche que désiré de la solution du problème, ainsi nous ne savons pas quelle validité à accorder à la solution obtenue.

Une reformulation ensembliste du problème ponctuel permettra de concevoir de nouvelles méthodes capables de caractériser de façon garantie et globale le résultat de la solution.

Dans cette représentation ensembliste des incertitudes, nous cherchons non pas une solution ponctuelle mais un « *ensemble solution* ». Un problème qui admet pour solution un ensemble est appelé « *problème ensembliste* ».

La théorie de l'approche ensembliste décrite dans ce chapitre est rapportée par plusieurs travaux [Jaulin, 1994], [Jaulin et Walter, 1993], [Videau, 2005], [Jaulin, 2004], [Berger, 1979] et [Lalami, 2008].

Le premier chapitre est organisé de la manière suivante :

Nous allons en premier lieu, présenter nos motivations qui nous ont permis d'entreprendre l'approche ensembliste dans notre travail (section 1.2), en confrontant celle-ci à l'une des approches statistiques les plus répandues pour la manipulation des incertitudes qui est l'approche probabiliste. Nous allons constater que l'approche ensembliste offre une meilleure modélisation des incertitudes. Les fondements de base de cette approche sont présentés dans la section (1.3), tels que les définitions et les opérations entre les ensembles. Enfin les différentes représentations d'ensembles sont illustrées et comparées (section (1.4)) afin de justifier le choix de la représentation des incertitudes par des intervalles et vecteurs intervalles qui sont la représentation ensembliste choisie dans notre étude.

1.2 Motivations

Obtenir un résultat garanti de la solution d'un problème donné, c'est de mettre en œuvre une méthode qui prend en considération les incertitudes omniprésentes sur les données du problème. Dans ce paragraphe, nous allons présenter et confronter deux approches qui prennent en compte ces incertitudes qui sont l'approche probabiliste et ensembliste.

Représentation d'une variable incertaine

L'approche probabiliste permet de représenter une variable scalaire ou vectorielle incertaine $x \in R^n$ par la densité de probabilité qui est une fonction définie de la manière suivante :

$$\pi: R^n \rightarrow R^+ \text{ tel que } \int_{R^n} \pi(x) = 1 \quad (1.1)$$

Cette représentation nous permet d'avoir un support X donné par :

$$X = \{x \in R^n / \pi(x) \neq 0\} \quad (1.2)$$

qui contient obligatoirement la variable incertaine x et nous informe de la manière dont celle-ci est distribuée dans ce support.

Dans une approche ensembliste, la variable incertaine est représentée uniquement par un ensemble X supposé contenir la variable incertaine x . Il est à noter que l'ensemble X peut comprendre des sous-ensembles où la variable x est introuvable. Par exemple :

Supposons que la variable incertaine x de R^n appartient à l'ensemble X_1 qui est dans ce cas représenté par l'union de deux intervalles (la représentation ensemblistes par des intervalles sera l'objet du chapitre 2) :

$$X_1 = [1,2] \cup [3,5]$$

Nous pouvons aussi dire que la variable x appartient à l'intervalle $X_2 = [1,5]$, car l'ensemble X_2 contient X_1 . Sauf que X_2 comprend une zone qui est $]2,3[$ qui ne contient pas la variable incertaine x .

Donc, dans l'aspect de caractérisation d'incertitudes, l'approche probabiliste est mieux avantagée dans ce sens, sauf que dans le cadre de la manipulation des calculs celle-ci est beaucoup plus difficile et souvent inexploitable. En effet l'approche probabiliste exige des connaissances statistiques sur les variables aléatoires et des lois conjointes entre les variables incertaines qui, dans le cas pratique, peuvent être indisponibles. C'est pour cela que dans la majorité des cas des simplifications sont faites, revenant à supposer des indépendances linéaires entre les variables aléatoires afin de mieux manipuler les calculs.

Dans le cas pratique, les mesures expérimentales prises par un capteur sont souvent sujettes à des erreurs. Ces mesures entachées d'incertitudes sont souvent corrélées (c'est-à-dire les erreurs sur les mesures sont dépendantes les unes des autres), alors le rajout de telles hypothèses statistiques telles que l'indépendance linéaire entre les mesures incertaines produit des résultats difficiles à interpréter.

Par contre la représentation par des ensembles des données expérimentales ne nécessite pas d'avoir des informations statistiques sur les mesures incertaines, ainsi une facilité de calcul est ressentie. Et dans le cadre de l'estimation des erreurs des mesures prises par le capteur, elles sont en générale données sous forme de tolérance, ainsi la représentation ensembliste (par exemple des intervalles) s'avère beaucoup plus intéressante et naturelle à entreprendre.

Si les incertitudes sont dues aux erreurs numériques de calcul, celles-ci peuvent être facilement rajoutées lors du calcul ensembliste. Ce qui est difficile à concevoir avec l'approche probabiliste comme le montre l'exemple suivant :

Si $x \in X$ et $y \in Y$, l'ensemble Z contenant $z = x + y$ sera donné par :

$$Z = X + Y + E \triangleq \{z + e = (x + y) + e | x \in X, y \in Y, e \in E\}$$

où E représente les erreurs introduites lors du calcul de Z .

Recouvrement de l'espace de recherche

Un espace de recherche de solutions [Jaulin, 1994] d'un problème donné est représenté par une infinité non dénombrable de points. Les algorithmes utilisés dans les sciences pour l'ingénieur consistent à balayer cet espace de recherche dans le but de trouver des solutions à un problème ou de montrer qu'une propriété est validée pour tout cet espace. Une approche ponctuelle n'analyse qu'une portion de l'espace de recherche, ainsi rien ne peut être conclu avec certitude que la solution a été trouvée.

Cependant, l'approche ensembliste offre une meilleure représentation de l'espace de recherche des solutions, en caractérisant cet espace, en le recouvrant par un nombre fini de sous ensembles simples sur lesquels on sait calculer, et qui peuvent contenir une infinité de points. C'est ce qu'on appelle « *recouvrement de l'espace de recherche* », nous pouvons alors résoudre avec garantie des problèmes en se basant sur les principes suivants [Jaulin, 1994] :

- Un ensemble est composé d'un nombre fini de sous-ensembles constitués d'une infinité d'éléments et les propriétés de cette représentation offre un résultat certain de la solution.
- Etant donné X un sous-ensemble de R^n représentant l'espace de recherche des solutions, cet espace lui-même est constitué de l'union de sous-ensembles. Si une propriété est validée pour tous les sous-ensembles le constituant, alors celle-ci est aussi validée pour tout l'espace de recherche.
- Si une solution ne se trouve pas dans un sous-ensemble de l'espace X , et nous pouvons prouver son inexistence, alors le sous ensemble est éliminé de l'espace de recherche. Si par contre, nous ne pouvons pas prouver en aucune manière son inexistence dans le sous-ensemble, alors il sera décomposé en sous-ensembles et l'opération se répète.

Cette dernière propriété est utilisée dans plusieurs algorithmes ensemblistes, tel que l'algorithme SIVIA (Set Inversion Via Interval Analysis) [Jaulin et Walter, 1993] qui sera illustré dans le deuxième chapitre.

1.3 Les fondements de l'approche ensembliste

Un problème mathématique, posé d'une manière ensembliste, fait intervenir des paramètres représentés par des ensembles, la résolution de ce type de problèmes se base sur des opérations ensemblistes. Nous allons, dans cette section, énoncer les propriétés qui interviennent lors du calcul ensembliste, nous aborderons les différentes opérations entre ensembles et étendre ces opérations à leurs éléments.

1.3.1 Ensembles

Un ensemble X est une collection de points qu'on nomme par éléments. Soit alors x un élément de l'ensemble X , on a alors les notions suivantes :

a) L'appartenance :

L'élément x appartient à l'ensemble X :

$$x \in X \quad (1.3)$$

L'élément x n'appartient pas à l'ensemble X :

$$x \notin X \quad (1.4)$$

L'ensemble des éléments x de X , ayant la propriété P est noté par l'ensemble :

$$\{x \in X \mid P\} \quad (1.5)$$

b) L'inclusion :

Soient deux ensembles X et Y :

Y est inclus dans X (ou X contient Y), s'il vérifie:

$$Y \subset X \text{ (ou } X \supset Y) \Leftrightarrow \forall y \in Y \Rightarrow y \in X \quad (1.6)$$

Y n'est pas inclus dans X :

$$Y \not\subset X \Leftrightarrow \exists y \in Y \text{ et } y \notin X \quad (1.7)$$

c) Egalité

Y est égal à X lorsque tout élément de Y est aussi élément de X et vice-versa :

$$Y = X \Leftrightarrow Y \subset X \text{ et } X \subset Y \quad (1.8)$$

Y n'est pas égal à X si :

$$\exists y \in Y \text{ et } y \notin X \text{ ou } \exists x \in X \text{ et } x \notin Y \quad (1.9)$$

d) Inclusion stricte

Y est strictement inclus dans X si et seulement si Y est inclus dans X sans lui être égal :

$$Y \subsetneq X \Leftrightarrow Y \subset X \text{ et } Y \neq X \quad (1.10)$$

e) Sous-ensemble

Y est un sous-ensemble de X si Y est inclus dans X ($Y \subset X$)

f) Espace produit

Notons par R l'ensemble des nombres réels. Nous allons utiliser par la suite des sous-ensembles de R et nous utilisons la notation d'ensemble aussi pour tout sous-ensemble de R afin d'alléger l'écriture.

Un espace produit de R est un ensemble noté R^{n_x} dont les éléments sont des n-uplets :

$$x = (x_1, x_2, \dots, x_{n_x})^T \quad \text{avec } x_i \in R \quad (1.11)$$

1.3.2 Opérations sur des ensembles réels

Dans cette section, nous allons illustrer les différentes opérations ensemblistes qu'on nomme par « opérations ensemblistes pures » qui sont : les opérations d'intersection, d'union, de produit cartésien, de différence, de complémentation et de projection. Puis, nous allons étendre ces opérations à leurs éléments.

a) Opérations ensemblistes pures

Soient X et $Y \subset \mathbb{R}^n$: la figure suivante illustre les différentes opérations entre ensembles :

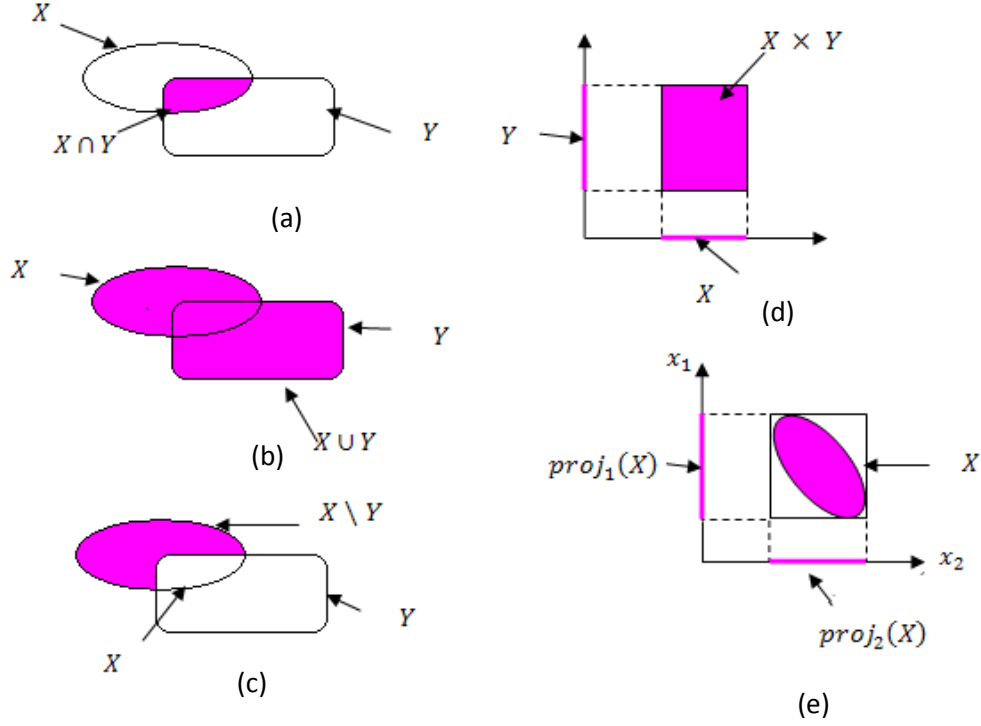


Figure 1.1 : Opérations ensemblistes pures

- L'intersection (a) :

$$X \cap Y \triangleq \{x \in \mathbb{R}^n \mid x \in X \text{ et } x \in Y\} \quad (1.12)$$

- L'union (b):

$$X \cup Y \triangleq \{x \in \mathbb{R}^n \mid x \in X \text{ ou } x \in Y\} \quad (1.13)$$

- La différence (c) :

$$X \setminus Y \triangleq \{x \in \mathbb{R}^n \mid x \in X \text{ et } x \notin Y\} \quad (1.14)$$

- Le produit cartésien de deux ensembles $X \subset \mathbb{R}^{n_x}$ et $Y \subset \mathbb{R}^{n_y}$ (d) :

$$X \times Y \triangleq \{ (x^T y^T)^T \in \mathbb{R}^{n_x+n_y} \mid x \in X \subset \mathbb{R}^{n_x} \text{ et } y \in Y \subset \mathbb{R}^{n_y} \} \text{ avec } n_x = n_y \quad (1.15)$$

- La projection canonique (e) :

$$\text{proj}_i(X) \triangleq \{x_i \in \mathbb{R} \mid \exists x = (x_1, x_2, \dots, x_{n_x})^T \in X\} \quad (1.16)$$

- La complémentation : la complémentation est un cas particulier de la différence notée $\neg X$ tel que :

$$\neg X = R^{n_x}/X \triangleq \{x \in R^{n_x} | x \notin X\} \quad (1.17)$$

b) Opérations étendues

- Soient deux ensembles $X \subset R^{n_x}$ et $Y \subset R^{n_y}$ et une fonction $f : R^{n_x} \rightarrow R^{n_y}$

On définit :

- L'image directe de X par f :

$$f(X) \triangleq \{y \in R^{n_y} | \exists x \in X, f(x) = y\} \quad (1.18)$$

- L'image inverse de Y par f :

$$f^{-1}(Y) \triangleq \{x \in R^{n_x} | \exists y \in Y, f(x) = y\} \quad (1.19)$$

Si un ensemble est défini par $\emptyset$ (l'ensemble vide) on a :

$$f(\emptyset) = f^{-1}(\emptyset) = \emptyset \quad (1.20)$$

- Soient $(X_1, X_2) \subset X$, $(Y_1, Y_2) \subset Y$ et $f : R^{n_x} \rightarrow R^{n_y}$. On alors les propriétés suivantes :

$$\square f(X_1 \cap X_2) \subset f(X_1) \cap f(X_2) \quad (1.21)$$

$$\square f(X_1 \cup X_2) = f(X_1) \cup f(X_2) \quad (1.22)$$

$$\square f^{-1}(Y_1 \cap Y_2) = f^{-1}(Y_1) \cap f^{-1}(Y_2) \quad (1.23)$$

$$\square f^{-1}(Y_1 \cup Y_2) = f^{-1}(Y_1) \cup f^{-1}(Y_2) \quad (1.24)$$

$$\square f(f^{-1}(Y)) \subset Y \quad (1.25)$$

$$\square f^{-1}(f(X)) \supset X \quad (1.26)$$

$$\square X_1 \subset X_2 \Rightarrow f(X_1) \subset f(X_2) \quad (1.27)$$

$$\square Y_1 \subset Y_2 \Rightarrow f^{-1}(Y_1) \subset f^{-1}(Y_2) \quad (1.28)$$

$$\square X \subset (Y_1 \times Y_2) \Rightarrow X \subset (proj_{Y_1}(X) \times proj_{Y_2}(X)) \quad (1.29)$$

Exemple1.1 :

Un aperçu sur les opérations entre ensembles faisant intervenir une fonction f est donné à travers l'exemple ci-dessous [Jaulin, 2004] :

Soit l'ensemble A contenant les éléments $\{a, b, c, e\}$ et B contenant les chiffres $\{1, 2, 3, 4, 5\}$. Et nous disposons de la fonction $f: A \rightarrow B$ qui est illustrée par la figure ci-dessous :

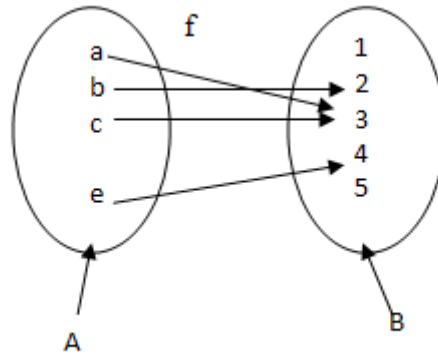


Figure 1.2 : Fonction f reliant les ensembles A et B .

Alors, nous avons les résultats suivants :

L'image de l'ensemble A est donnée par l'ensemble :

$$f(A) = \{2, 3, 4\} \quad (1.30)$$

L'image inverse de l'ensemble B est donnée par l'ensemble :

$$f^{-1}(B) = \{a, b, c, e\} \quad (1.31)$$

L'image inverse de $f(A)$ vérifie :

$$f^{-1}(f(A)) = \{a, b, c, e\} \subset A \quad (1.32)$$

Et l'image inverse de l'image des éléments b et c de l'ensemble A est :

$$f^{-1}(f(\{b, c\})) = \{a, b, c\} \quad (1.33)$$

Si nous prenons par exemple la fonction $f(x) = x^2$, alors :

$$f([2, 3]) = [4, 9]$$

$$f^{-1}([4, 9]) = [-3, -2] \cup [2, 3]$$

Ce qui est confirmé par la propriété (1.25): $f(f^{-1}(Y)) \subset Y$

1.3.3 Distance entre deux ensembles compacts

Définition 1.1 [Videau, 2005] :

Un ensemble compact de R^n est un sous-ensemble clos et borné de R^n . La distance entre deux points par exemple, ou entre un point et une droite est la longueur minimale des chemins possibles d'un point à un autre ou d'un point à une droite.

Dans le cas de deux ensembles $X, Y \subset R^n$, la distance entre ces deux ensembles est un scalaire positif (d) donné par:

$$d(X, Y) = \inf\{d(x, y) \mid x \in X \text{ et } y \in Y\} \quad (1.34)$$

Cette notion de distance permet de mesurer la similitude entre deux ensembles :

- Si la distance d est grande alors les deux ensembles sont différents.
- Si la distance d est petite alors les deux ensembles se ressemblent.
- Si la distance d égale à 0 c'est que les deux ensembles sont identiques.

Définition 1.2 (Distance de Hausdorff [Berger, 1979]) :

La distance de Hausdorff entre deux ensembles X et $Y \subset R^n$ est donnée par :

$$h_\infty(X, Y) \triangleq \max(h_\infty^0(X, Y), h_\infty^0(Y, X)) \quad (1.35)$$

Où $h_\infty^0(X, Y)$ et $h_\infty^0(Y, X)$ sont respectivement la proximité de X à Y et la proximité de Y à X données par :

$$\begin{aligned} h_\infty^0(X, Y) &\triangleq \inf(r \in R^+ / X \subset Y + r U) \\ h_\infty^0(Y, X) &\triangleq \inf(r \in R^+ / Y \subset X + r U) \end{aligned} \quad (1.36)$$

Où U est une sphère unité de (R^n, L_∞) et L_∞ est la distance entre deux points, donnée par :

$$L_\infty(x, y) = \max_{i=\{1 \dots n\}} |y_i - x_i| \quad (1.37)$$

Exemple 1.2

L'exemple suivant [Videau, 2005] considère deux ensembles $X, Y \in R^2$ représentés par la figure 1.3. Pour obtenir $h_\infty^0(X, Y)$, on effectue une inflation sur Y de rU jusqu'à ce que $X \subset Y + rU$ (de la même manière on retrouve $h_\infty^0(Y, X)$).

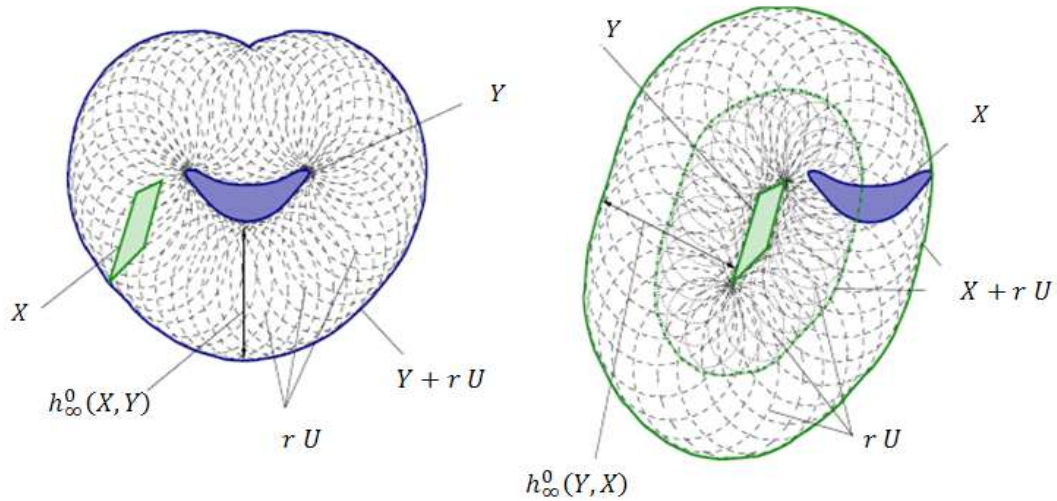


Figure 1.3 : Distance de Hausdorff entre deux ensembles compacts X et Y

Cette distance est définie dans l'ensemble des compacts de R^n et elle vérifie les relations suivantes :

- | | |
|-----------------------------|---|
| 1. La séparabilité | $h_\infty(X, Y) = 0 \Rightarrow X = Y$ |
| 2. La symétrie | $h_\infty(X, Y) = h_\infty(Y, X)$ |
| 3. L'inégalité triangulaire | $h_\infty(X, Z) \leq h_\infty(X, Y) + h_\infty(Y, Z)$ |

1.4 Représentation des sous ensembles

Il existe plusieurs caractérisations des sous-ensembles de R^n , ces représentations ensemblistes doivent être exactes ou le plus possible approchées, et elles doivent aussi être efficaces dans le cadre de la manipulation des calculs des sous-ensembles et riches en matière de propriétés afin d'avoir de meilleurs résultats. Nous allons citer dans cette partie quelques unes de ces représentations ensemblistes :

1.4.1 Représentation par un nuage de points

Dans cette représentation l'ensemble X est sous forme d'un nuage constitué d'un ensemble de points qui est contenu dans X (le nuage de point est alors un sous-ensemble de X). La manipulation des nuages de points est très facile, sauf que la méthode est approximative à cause de la manière dont les points ont été obtenus et l'interprétation des

résultats est souvent arbitraire. Il existe deux méthodes pour pouvoir obtenir ce nuage de points :

- ❖ La première méthode est « l'approche Monté Carlo »: les points représentant l'ensemble X sont obtenus par un tirage aléatoire.
- ❖ La deuxième méthode : les points sont obtenus par un balayage déterministe.

Un exemple montrant cette difficulté dans l'interprétation des résultats par cette représentation ensembliste est lorsque nous disposons d'un nuage de points qui est vide. Dans ce type de problèmes rencontrés, nous ne pouvons rien conclure sur les propriétés de l'ensemble X contenant ce nuage de points, la seule propriété que nous pouvons tirer, est qu'il existe des points qui ne sont pas dans l'ensemble X .

1.4.2 Représentation par équation et inéquation

Ce type de représentation permet d'obtenir la ou les solutions (numériques ou formelles) d'un problème de façon exacte. Par exemple trouver l'ensemble des réels qui sont les solutions de l'équation $x^2 - 1 = 0$, les solutions -1 et 1 peuvent être obtenus avec exactitude.

L'inconvénient rencontré est que cette représentation s'applique à une classe restreinte d'équations et d'inéquations. Les algorithmes qui sont développés pour cette approche sont complexes et souvent inexploitable à cause du temps de calcul et d'espace mémoire.

1.4.3 Représentation par recouvrement

La représentation par recouvrement consiste à représenter un ensemble X de R^n par une union d'éléments qu'on nomme récipients (qui sont des sous-ensembles de R^n) qui recouvrent complètement l'ensemble X .

Ces récipients peuvent être un ensemble d'**ellipsoïdes**, de **polytopes**, de **zonotopes** ou des **pavés**..., cette représentation par recouvrement de l'ensemble X nous fournit d'intéressantes informations caractérisant cet ensemble. Par exemple, si nous disposons d'un recouvrement qui est vide alors forcément l'ensemble X est aussi vide.

Nous présentons ci-dessous quelques représentations de ces réceptifs :

a) Les zonotopes

Les zonotopes sont une classe particulière des polytopes, ils correspondent à la projection par une approche linéaire d'un hypercube de dimension élevée sur un espace considéré de dimension plus faible.

Ces zonotopes sont modélisés par des matrices intervalles qui décrivent l'application linéaire de la transformation de l'hypercube.

Exemple 1.3

La projection d'un cube de dimension 3D sur un espace de dimension 2D comme le montre la figure suivante :

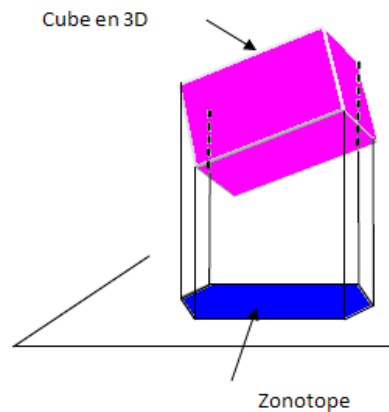


Figure1.4 : Exemple d'un zonotope.

Le zonotope retrouvé par la projection sur 2D, est représenté par une matrice intervalle 2 lignes (projection sur l'espace 2D) et 3 colonnes (c'est l'image linéaire d'un cube 3D).

- ❖ Cette propriété décrivant la transformation linéaire par des matrices permet de justifier l'application des zonotopes dans le domaine ensembliste, car elle permet de recouvrir une large zone d'éléments, et permet de faire des manipulations matricielles.
- ❖ Les zonotopes sont souvent utilisés pour représenter les domaines atteignables des systèmes dynamiques à entrées bornées. Plus généralement, ils sont utilisés pour caractériser les relations de dépendances linéaires des variables incertaines.

Plusieurs domaines d'études s'intéressent à cette représentation par les zonotopes, nous citons par exemple son utilisation dans le diagnostic [Lalami, 2008].

b) Les pavés

Soit un sous-ensemble X de R^n . Une représentation de cet sous-ensemble par un pavé est donnée par le vecteur intervalle suivant :

$$[X] = ([x_1] \times [x_2] \times \dots \times [x_n]) \quad (1.38)$$

qui est le produit cartésien de n intervalles $[x_i] \in R$ ($i = 1 \dots n$).

- ❖ Cette représentation nous permet de manipuler que des intervalles. Cette caractéristique nous fournit une simplicité de représentation d'ensembles et le calcul ensembliste est ainsi plus aisé.

1.4.4 Représentation par encadrement

Cette approche est développée par L. Jaulin [Jaulin, 1994]. Elle consiste à encadrer un ensemble compact X de R^n par un couple de récipients $[\underline{X}, \bar{X}]$ qui représentent respectivement un encadrement **intérieur** et **extérieur** de l'ensemble X comme le montre la relation suivante :

$$\underline{X} \subset X \subset \bar{X} \quad (1.39)$$

La connaissance des encadrements $\underline{X}$ et $\bar{X}$, nous donne de précieuses informations sur l'ensemble X . Cette propriété d'encadrement d'ensembles est utilisée dans plusieurs algorithmes ensemblistes, tels que l'algorithme SIVIA (Set Inversion Via Interval Analysis [Jaulin et Walter, 1993]) où $\underline{X}$ et $\bar{X}$ sont des vecteurs intervalles.

Vue la simplicité ressentie lors de la manipulation des intervalles, la majorité de ces représentations ensemblistes utilise l'outil de l'analyse par intervalles dans les calculs.

1.5 Conclusion

Ce chapitre a été consacré en premier lieu à relever les motivations de l'utilisation de l'approche ensembliste. Celles-ci sont dues principalement à la capacité de mieux prendre en compte les incertitudes sur les données d'un problème et sur les méthodes et les moyens de résolution.

En second lieu, nous avons présenté les fondements de base de cette approche et les notions très importantes reliées aux manipulations des ensembles. Parmi les différents outils développés dans cette théorie, ceux basés sur la représentation par intervalles ou vecteurs intervalles est la plus utilisée et ceci en raison de la simplicité dans la manipulation des différentes opérations ensemblistes (l'arithmétique) et de la façon la plus naturelle de représenter les incertitudes. Cette arithmétique sera l'objet du deuxième chapitre.

Chapitre 2:

Arithmétique des intervalles

Chapitre2 :

Arithmétique des intervalles

2.1 Introduction

La majorité des méthodes ensemblistes développées pour l'estimation sont basées sur un outil de calcul qui est l'arithmétique des intervalles. Initialement, cette arithmétique est développée pour quantifier les erreurs numériques dues à la représentation finie des nombres réels sur un ordinateur [Moore, 1962] et [Moore, 1966]. L'analyse par intervalles est devenue un outil de résolution numérique globale et garanti des problèmes d'estimation.

Durant ces dernières décennies, différentes méthodes basées sur l'analyse par intervalles ont été développées dans le but d'étudier et de quantifier les incertitudes (numériques ou physiques) influant sur les données manipulées.

Le deuxième chapitre est structuré comme suit :

Dans la section (2.2), nous allons dresser un historique de l'arithmétique des intervalles. Son origine, son utilité et ses débuts d'applications en automatique sont présentés.

L'analyse par intervalles est l'outil de base de notre étude, dont les notions élémentaires sont instanciées dans la section (2.3). Nous présentons dans cette partie les définitions générales des intervalles (section 2.3.1) et les vecteurs intervalles ou ce qu'on nomme par pavés (section 2.3.2). Dans la section (2.3.3), nous allons illustrer l'un des problèmes majeur de l'algèbre des intervalles qui est « le pessimisme » dû au fait qu'un résultat d'une suite d'opération entre intervalles n'est pas minimale. Des solutions à cet obstacle rencontré sont illustrées dans la section (2.3.4) et (2.3.5).

Enfin, dans les sections (2.4) et (2.5), des méthodes issues de l'analyse par intervalles sont présentées. Celles-ci permettent de résoudre des problèmes d'équations linéaires, non linéaires ou encore des systèmes d'équations linéaires. Quelques exemples d'application de l'arithmétique des intervalles sont illustrés tels que la stabilité robuste et l'estimation paramétrique à erreur bornée.

Les différents exemples sont développés dans la toolbox « INTLAB » de MATLAB, ainsi qu'aux applications telles que « INTERVAL PEELER » et « Simulateur analyse par intervalle » (Annexe A).

2.2 Historique

Le début de l'arithmétique par intervalles apparaît de manière confidentielle vers les années soixante. À cette époque, le calcul scientifique n'était pas puissant et l'arithmétique flottante n'était pas encore bien spécifiée, les calculs effectués dans les machines numériques étaient arrondis.

La question qui s'est alors posée, est de savoir comment ces incertitudes numériques évoluent au cours des calculs et comment quantifier l'erreur sur le résultat de calcul. C'est en s'attachant à résoudre ce problème qu'en 1962 R. Moore ouvrit la voie de l'analyse par intervalles [Moore, 1962], en publiant notamment en 1966 un ouvrage de référence mettant en place les fondements de base de cet outil [Moore, 1966].

Ce n'est que bien plus tard, vers les années 1980 que l'arithmétique par intervalles vit son essor. En effet, les premiers à avoir introduit cet outil est la firme IBM en introduisant un jeu d'instructions et un compilateur intégrant cette arithmétique.

À partir de l'année 1990, Neumaier développe plus spécifiquement la résolution des systèmes d'équations linéaires et non linéaires en utilisant cette arithmétique [Neumaier, 1991].

Hansen développe en 1992 des méthodes pour la résolution des problèmes d'optimisation globale [Hansen, 1992], c'est-à-dire rechercher de manière garantie tous les minimiseurs d'un critère connexe, alors que les méthodes classiques non linéaires (Gradient, Newton-Raphson,...) ne renvoient qu'un seul optimum pouvant n'être que locale.

En utilisant l'outil intervalles, Schichl et Neumaier développent une nouvelle technique au problème d'optimisation globale sous contraintes (au moyen de la propagation des contraintes qui est issue de la programmation logique) [Schichl et Neumaier, 2005].

Si la montée en puissance de l'analyse par intervalles est due dans un premier temps aux difficultés rencontrées dans les calculateurs numériques, puis à l'étude des problèmes complexes tels que les problèmes d'optimisation globale. Ces débuts dans le domaine de l'automatique est très récent. L'utilisation de l'arithmétique des intervalles en automatique devient de plus en plus populaire car le principal atout qu'offre cette technique réside dans la capacité à concevoir une meilleure modélisation des incertitudes influant sur les systèmes physiques.

La source de ces incertitudes est due principalement aux appareils de mesure, aux perturbations et bruits agissant sur les systèmes. Ainsi, le fait de représenter une donnée expérimentale incertaine dans un intervalle, nous garantit de renfermer cette donnée entre des bornes qui sont connues *a priori*.

À l'aide de l'analyse par intervalles, beaucoup de concepts en automatique ont été redéfini, tels que : l'estimation paramétrique, l'estimation d'état, la commandabilité et l'observabilité.

Plusieurs travaux s'inspirent de cette technique, citons par exemple : L. Jaulin qui utilise la technique de propagation de contraintes par l'arithmétique d'intervalles en estimation paramétrique à erreur bornée, dans la stabilité robuste, en estimation d'état et dans la calibration robotique [Jaulin, 2004]. Comme, il utilise l'arithmétique des intervalles dans des algorithmes d'inversion ensembliste (l'ensemble solution est défini comme l'inverse d'un ensemble par une fonction) en estimation paramétrique à erreur bornée [Jaulin et Walter, 1993] et dans la caractérisation des domaines de stabilité [Walter et Jaulin, 1994].

Cette arithmétique est aussi utilisée dans la conception des estimateurs d'état pour la classe des systèmes non linéaires continus, en se basant sur l'approche prédiction-correction, utilisant l'inversion ensembliste par une analyse intervalles [Raissi et al., 2003].

Il ne faut pas oublier aussi certains résultats dans le domaine biologique, pour la synthèse d'observateurs intervalles des systèmes biologiques caractérisés par des incertitudes [Gouzé et al., 2000], [Bernard et Gouzé, 2004] et [Moisan, 2007] .

2.3 Notions sur l'arithmétique des intervalles

2.3.1 Les intervalles

Définitions 2.1

- Un intervalle $[a]$ est un ensemble connexe, borné, fermé de bornes inférieure et supérieure données par $\underline{a}$ et $\bar{a}$ respectivement qui sont des nombres réels. L'intervalle $[a]$ est représenté de la manière suivante :

$$[a] = [\underline{a}, \bar{a}] = \{a \in R \mid \underline{a} \leq a \leq \bar{a}\} \quad \text{avec} \quad \underline{a}, \bar{a} \in R \quad (2.1)$$

- Soit $\mathbf{IR}$ l'ensemble des intervalles de R , alors l'intervalle $[a]$ est un élément de $\mathbf{IR}$ ($[a] \in \mathbf{IR}$) et nous pouvons designer par **support** tout intervalle $[a]$ ($[a] \subset \mathbf{IR}$).
- Un intervalle $[a]$ est dit dégénéré si la borne supérieure et inférieure de cet intervalle sont identiques.

Exemple 2.1

$[1,3], \{1\},]-\infty, 6], R, \emptyset$ sont des intervalles, alors que $]1,3[, [4,2], [1,2] \cup [4,5]$ ce ne sont pas des intervalles car ils ne vérifient pas les définitions données ci-dessus.

Les opérations mathématiques sur les réels sont étendues aux intervalles. Lorsque nous travaillons avec des intervalles à bornes finies, nous ne manipulons que les bornes de ces derniers.

Nous disposons de l'opération élémentaire $\circ \in \{\oplus; \ominus; \otimes; \oslash\}$ qui représente respectivement l'addition, la soustraction, la multiplication et la division. Et des deux intervalles $[a], [b] \in IR$.

Définitions 2.2

- Pour les opérations élémentaires $\circ \in \{\oplus; \ominus; \otimes\}$ le résultat des opérations pour les deux intervalles est donné par la relation suivante:

$$[a] \circ [b] = \{a \circ b \mid a \in [a], b \in [b]\} \quad (2.2)$$

Nous avons alors pour l'addition, la soustraction et la multiplication des deux intervalles $[a]$ et $[b]$ les opérations suivantes:

$$[a] \oplus [b] = [\underline{a} + \underline{b}, \bar{a} + \bar{b}] \quad (2.3)$$

$$[a] \ominus [b] = [\underline{a} - \bar{b}, \bar{a} - \underline{b}] \quad (2.4)$$

$$[a] \otimes [b] = [\min(\underline{a}\underline{b}, \underline{a}\bar{b}, \bar{a}\underline{b}, \bar{a}\bar{b}), \max(\underline{a}\underline{b}, \underline{a}\bar{b}, \bar{a}\underline{b}, \bar{a}\bar{b})] \quad (2.5)$$

- Pour la division ($\oslash$), la relation n'est plus valable dans le cas où $0 \in [b]$. Dans ce dernier cas, le résultat n'est pas un intervalle mais la réunion de deux intervalles. Il n'existe pas de définition unique lorsque $[b] = [0, 0]$.

Le résultat de la division peut être donné par les relations suivantes [Jaulin et al., 2001]:

$$\begin{aligned} * [a] \oslash [b] &= \emptyset \text{ si } [b] = [0, 0] \\ * [a] \oslash [b] &= [a] \otimes [1/\bar{b}, 1/\underline{b}] \text{ si } 0 \notin [b] \\ * [a] \oslash [b] &= [a] \otimes [1/\bar{b}, +\infty[\text{ si } \underline{b} = 0 \text{ et } \bar{b} > 0 \\ * [a] \oslash [b] &= [a] \otimes] -\infty, 1/\underline{b}] \text{ si } \underline{b} < 0 \text{ et } \bar{b} = 0 \\ * [a] \oslash [b] &=] -\infty, +\infty[\text{ si } \underline{b} < 0 \text{ et } \bar{b} > 0 \end{aligned} \quad (2.6)$$

Définitions 2.3

- Soit l'intervalle $[a] \in IR$, les définitions suivantes sont données :

$$\begin{aligned} * \text{La borne inférieure : } \inf([a]) &= \underline{a} \\ * \text{La borne supérieure : } \sup([a]) &= \bar{a} \\ * \text{La largeur : } w([a]) &= \bar{a} - \underline{a} \geq 0 \\ * \text{Le milieu : } \text{mid}([a]) &= (\bar{a} + \underline{a})/2 \\ * \text{Le rayon : } \text{rad}([a]) &= (\bar{a} - \underline{a})/2 \geq 0 \end{aligned} \quad (2.7)$$

NB : Par convention la borne inférieure et supérieure de l'ensemble vide $\emptyset$ est $+\infty$ et $-\infty$ respectivement.

- Considérons une variable incertaine $\mathbf{a}$ appartenant au support $[a] = [\underline{a}, \bar{a}]$, la variable $\mathbf{a}$ peut être estimée par le $\text{mid}([a])$ et son incertitude par $\text{rad}([a])$.

2.3.2 Les pavés

Soit IR^n l'ensemble des intervalles de R^n , un pavé $[a]$ est un vecteur intervalle de dimension n (c'est-à-dire c'est le produit cartésien de n intervalles) représenté de la manière suivante :

$$[a] = ([\underline{a}_1, \bar{a}_1] \times [\underline{a}_2, \bar{a}_2] \times \cdots \times [\underline{a}_n, \bar{a}_n]) \quad (2.8)$$

Les opérations élémentaires définies précédemment pour les intervalles sont également redéfinies pour les vecteurs intervalles.

Définition 2.4

Soit le pavé $[a] \in IR^n$, on définit :

- * La borne inférieure du pavé $[a]$: $\inf([a]) = (\underline{a}_1, \underline{a}_2, \dots, \underline{a}_n)^T$
- * La borne supérieure du pavé $[a]$: $\sup([a]) = (\bar{a}_1, \bar{a}_2, \dots, \bar{a}_n)^T$
- * La largeur du pavé $[a]$: $w([a]) = \max_{j=1}^n (\bar{a}_j - \underline{a}_j) \geq 0$
- * Le milieu du pavé $[a]$: $\text{mid}([a]) = (\text{mid}([a_1]), \text{mid}([a_2]), \dots, \text{mid}([a_n]))^T$

(2.9)

Exemple 2.2

Soit le pavé $[a] = ([1,2] \times [-1,3])$, sa largeur est $w([a]) = \max(1,4) = 4$.

NB : Par convention la largeur de l'ensemble vide $w(\emptyset)$ est $-\infty$

Définition 2.5

Soient deux vecteurs intervalles $[a]$ et $[b]$ dans IR^n , l'inégalité $[a] \geq [b]$ (ou $[a] \leq [b]$) est calculée élément par élément de la manière suivante:

$$[a] \geq [b] \text{ si } [a_i] \geq [b_i], \forall i = \overline{1, n} \quad (2.10)$$

Avec $[a_i], [b_i]$ sont les intervalles des pavés $[a]$ et $[b]$ respectivement.

Définition 2.6

La norme infinie d'un vecteur intervalle $[a]$ est calculée par le maximum de la valeur absolue des intervalles $[a_i]$ du pavé, qui est donnée par la relation suivante :

$$\|[a]\|_{\infty} = \max\{|[a_i]|, i = 1, \dots, n\} \quad (2.11)$$

Avec la valeur absolue d'un intervalle $[a] = [\underline{a}, \bar{a}]$ est calculée comme suit :

$$|[a]| = \max\{|\underline{a}|, |\bar{a}|\} \quad (2.11)'$$

2.3.3 Le pessimisme

Lors de la manipulation des vecteurs intervalles, nous pouvons avoir un résultat d'une suite d'opération entre pavés qui ne soit pas minimal, alors le résultat obtenu est dit pessimiste. Le problème est principalement lié à deux phénomènes qui sont la dépendance et l'enveloppement qui sont illustrés ci-dessous :

a) Le phénomène de dépendance

Définition 2.7 [Moore, 1966]

Soit un intervalle non dégénéré $[a] = [\underline{a}, \bar{a}]$ et l'opération élémentaire $\circ \in \{\oplus; \ominus; \otimes; \oslash\}$. Effectuer une opération élémentaire entre l'intervalle $[a]$ et lui-même est réalisé de la manière suivante :

$$[a] \circ [a] = \{a \circ b \mid a \in [a], b \in [a]\} \quad (2.12)$$

Avec les éléments a et b sont indépendants. Donc, malgré que l'on fasse cette opération avec le même intervalle, nous prenons toujours un élément a de $[a]$ et un autre élément b du même intervalle qui sont indépendants, c'est ce qu'on appelle le phénomène de dépendance.

Exemple 2.3

- Soit l'intervalle $[a] = [-5, 4]$ et soit l'opération élémentaire $\{\ominus\}$. Alors, effectuer la soustraction $[a] \ominus [a]$ ne donne pas un zéro, mais la soustraction est faite de la manière suivante :

$$[a] \ominus [a] = [-5, 4] \ominus [-5, 4] = [-5 - 4, 4 - (-5)] = [-9, 9]$$

- Si maintenant nous disposons de l'opération élémentaire suivante $[a] \otimes [a]$ Nous avons :

$$[a] \otimes [a] = [\min(25, -20, -20, 16), \max(25, -20, -20, 16)] = [-20, 25]$$

Or, si nous évaluons $[a]^2$ nous aurons réduit le pessimisme. Nous trouvons $[a]^2 = [0, 25]$

avec $[a]^2$ est calculé de la manière suivante :

$$[a]^2 = [0, \max((-5)^2, 4^2)] = [0, 25]$$

b) Le Phénomène d'enveloppement

Ce phénomène est lié principalement à la représentation des ensembles par des vecteurs intervalles (pavés), ceci est expliqué à travers l'exemple ci dessous :

Exemple 2.4

Nous disposons d'une matrice ponctuelle carrée :

$$A = \begin{pmatrix} 1 & 2 \\ 0 & 1 \end{pmatrix} \quad (2.13)$$

et d'un pavé donné par :

$$[x] = \begin{pmatrix} [-1,0] \\ [1,2] \end{pmatrix} \quad (2.14)$$

qui peut aussi être réécrit de la manière suivante :

$$[x] = ([-1,0] \times [1,2]) \quad (2.15)$$

Si nous calculons :

$$[Y] = A[x] \quad (2.16)$$

Nous obtenons le pavé suivant (à noter que la multiplication est calculée par l'arithmétique d'intervalles) :

$$[Y] = \begin{pmatrix} [1, 4] \\ [1, 2] \end{pmatrix} \quad (2.17)$$

Le vecteur intervalle $[Y]$ est représenté par le rectangle en bleu dans figure 2.1. La représentation exacte de la multiplication précédente est donnée par l'ensemble suivant :

$$y = \{Ax \mid x \in [x]\} \quad (2.18)$$

qui est la représentation réelle (le parallélogramme en vert dans figure 2.1) celle-ci est contenue dans le pavé $[Y]$. Nous remarquons alors qu'un pessimisme est introduit.

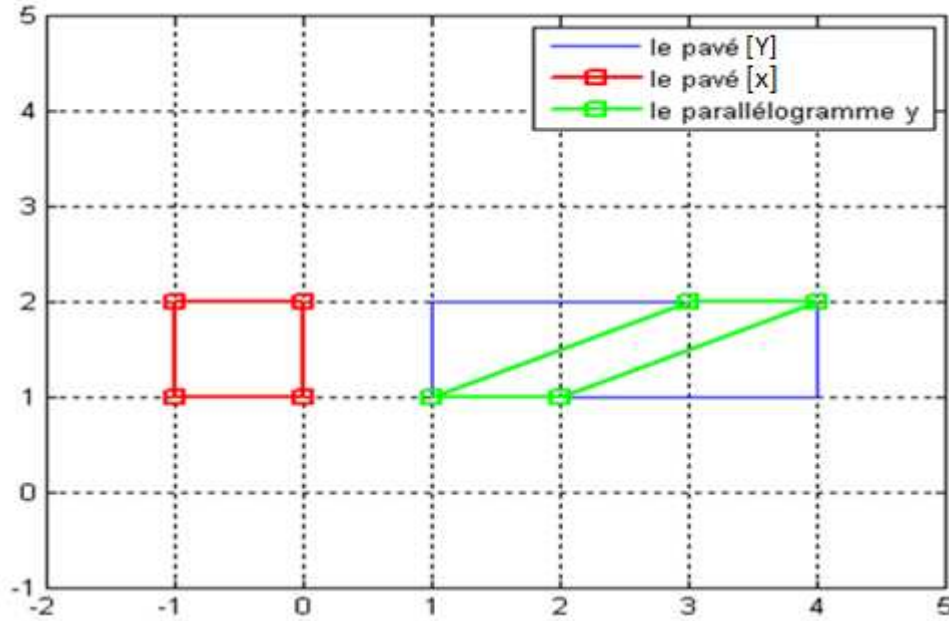


Figure 2.1 : Représentation du pavé $[Y]$ et de l'ensemble réel y .

Donc, le fait de manipuler des vecteurs intervalles, le résultat obtenu n'est pas minimale c'est ce que nous appelons « pessimisme dû à l'effet l'enveloppement ».

2.3.4 Fonctions d'inclusion

L'une des techniques utilisées pour réduire l'effet du pessimisme défini précédemment est la représentation par des fonctions d'inclusion dont l'efficacité est explicitée en s'appuyant sur des exemples illustrant cet effet :

Soit f une fonction vectorielle définie par $f : D \subset \mathbb{R}^{n_x} \rightarrow \mathbb{R}^{n_y}$, l'évaluation de f sur un pavé $[x]$ est donnée par l'ensemble suivant :

$$f([x]) = \{f(x) \mid x \in [x]\} \quad (2.19)$$

Une fonction d'inclusion de f notée $[f]$ est définie par :

$$[f] : \mathbb{R}^{n_x} \rightarrow \mathbb{R}^{n_y} \quad (2.20)$$

vérifiant l'inclusion suivante

$$f([x]) \subset [f]([x]) \quad (2.21)$$

Cette inclusion nous informe que la fonction d'inclusion $[f]$ est une extension aux intervalles d'une fonction f . Cette représentation contribue à la réduction du phénomène de pessimisme. Selon la manière dont la fonction f est écrite, nous pouvons avoir plusieurs fonctions d'inclusion.

Une fonction d'inclusion est dite peu pessimiste si la différence donnée par la taille suivante :

$$([f]([x]) - f([x])) \quad (2.22)$$

est minimale (assez petite).

Une fonction d'inclusion est minimale si pour tout pavé $[x]$, la fonction d'inclusion $[f]([x])$ est le plus petit pavé qui contient f , elle est notée par $[f]^*$ et elle est unique.

a) Les différentes fonctions d'inclusion

a.1) Fonction d'inclusion des fonctions élémentaires

Soit une fonction élémentaire définie par $f : D \subset \mathbb{R} \rightarrow \mathbb{R}$ (les fonctions élémentaires sont $\exp, \log, \cos, \sin, \dots$)

La fonction d'inclusion d'une fonction élémentaire est basée sur le principe de monotonie suivant :

- o Si f est une fonction continue et croissante sur le domaine D alors l'évaluation de sa fonction d'inclusion $[f]$ sur un pavé $[x] = [\underline{x}, \overline{x}] \in D$ est un pavé définie par :

$$[f] = [f(\underline{x}), f(\overline{x})] \quad (2.23)$$

- o Si f est continue et décroissante sur $[x]$ alors $[f]$ est le pavé suivant :

$$[f] = [f(\overline{x}), f(\underline{x})] \quad (2.24)$$

Exemple 2.5

✓ La fonction $\log(x)$

Nous disposons de la fonction $f(x) = \log(x)$ qui est définie dans le domaine $D =]0, +\infty[$. La fonction est continue et croissante sur ce domaine alors $\forall [x] \in D$, sa fonction d'inclusion est donnée par :

$$[f]([x]) = [\log(\underline{x}), \log(\overline{x})]$$

✓ **La fonction cos(x)**

Soit $f(x) = \cos(x)$ définie sur $D = [0, 2\pi]$.

Sur le domaine $D1 = [0, \pi]$ la fonction est décroissante alors l'évaluation de la fonction d'inclusion sur un pavé $[x] \in D1$ est donnée par :

$$[f] = [\cos(\bar{x}), \cos(\underline{x})]$$

et sur le domaine $D2 = [\pi, 2\pi]$ la fonction est croissante, alors $\forall [x] \in D2$ sa fonction d'inclusion est la suivante:

$$[f]([x]) = [\cos(\underline{x}), \cos(\bar{x})]$$

a.2) Fonction d'inclusion naturelle

Soit la fonction vectorielle $f : D \subset \mathbb{R}^{n_x} \rightarrow \mathbb{R}^{n_y}$, la fonction d'inclusion naturelle $[f_n]$ de f s'obtient de la manière suivante :

- Chaque élément x_i de la fonction f est remplacé par le pavé $[x_i]$.
- Chaque opération arithmétique est remplacée par son extension intervalle.

Dans le cas où la dimension n_y égale à 1, la fonction d'inclusion est minimale si nous avons les deux conditions suivantes qui sont vérifiées :

- La fonction f est composée de fonctions continues.
- Toutes les variables sont monooccurences (c'est-à-dire elles n'apparaissent qu'une fois dans f) comme le montre l'exemple suivant :

Exemple 2.6

Nous considérons ces quatre écritures de la fonction f

- $f_1(x) = x(x + 1)$
- $f_2(x) = x \times x + x$
- $f_3(x) = x^2 + x$
- $f_4(x) = \left(x + \frac{1}{2}\right)^2 - \frac{1}{4}$

L'évaluation de ces quatre fonctions dans le pavé $[x] = [-1, 1]$ donne les fonctions d'inclusion naturelle suivantes :

- $[f_{n1}]([x]) = [-2, 2]$
- $[f_{n2}]([x]) = [-2, 2]$
- $[f_{n3}]([x]) = [-1, 2]$
- $[f_{n4}]([x]) = [-0.25, 2]$

Nous remarquons à partir des résultats trouvés, que la taille des intervalles obtenus par ces quatre fonctions d'inclusion naturelles, dépend de l'expression utilisée pour l'écriture de la fonction f . La fonction d'inclusion $[f_{n4}]$ est la fonction d'inclusion minimale car l'intervalle obtenu est le plus petit et la fonction f satisfait les conditions citées auparavant, c'est à dire la continuité et le nombre d'occurrences des variables qui est égal à 1.

a.3) Fonction d'inclusion moyenne

Soit la fonction $f : D \subset R^{n_x} \rightarrow R^{n_y}$ de classe $C^1(D)$ et un pavé $[x] \subset D$ et un point $x_o \in [x]$. La fonction d'inclusion moyenne de f est donnée par la relation suivante :

$$[f_m]([x]) \triangleq f(x_o) + [J_f]([x])([x] - x_o) \quad (2.25)$$

avec $[J_f]$ est la fonction d'inclusion du jacobien de f et le point x_o est pris généralement comme le milieu de $[x]$ ($x_o = mid([x])$).

Si la fonction f est une fonction scalaire à une seule variable alors la fonction d'inclusion moyenne est calculée par:

$$[f_m](\tilde{x}) \triangleq f(x_o) + [f']([x])(\tilde{x} - x_o) \text{ avec } \tilde{x} \in [x] \quad (2.26)$$

a.4) Fonction d'inclusion de Taylor

La fonction d'inclusion de Taylor [Hansen, 1992] est une généralisation de la fonction d'inclusion moyenne, ce type de fonction est obtenu en utilisant un ordre de dérivation plus élevé. La fonction d'inclusion de Taylor de second ordre est donnée par la relation suivante :

$$[f_T]([x]) \triangleq f(x_o) + J_f(x_o)([x] - x_o) + \frac{1}{2}([x] - x_o)^T [H_f]([x])([x] - x_o) \quad (2.27)$$

avec

$J_f(x_o)$ est le jacobien de f en x_o et $[H_f]$ est la fonction d'inclusion du Hessian de la fonction f .

Afin de montrer la meilleure représentation par une fonction d'inclusion, nous allons, dans l'exemple suivant, comparer les fonctions d'inclusion naturelle, moyenne et de Taylor d'une fonction donnée sur un intervalle.

Exemple2.7

Soit la fonction $f = x(1 - x)$, nous allons calculer les fonctions d'inclusion naturelle moyenne et de Taylor dans le pavé $[x] = [0,1]$

Les résultats obtenus pour les trois fonctions d'inclusion sont les suivants :

- La fonction d'inclusion naturelle de f est :

$$[f_n]([x]) = [0.0000, 1.0000]$$

- La fonction d'inclusion moyenne de f est :

$$[f_m]([x]) = [-0.2501, 0.7501]$$

- La fonction d'inclusion de Taylor de f est :

$$[f_T]([x]) = [0.0000, 0.2501]$$

Selon les résultats obtenus, nous pouvons conclure que la fonction d'inclusion de Taylor est moins pessimiste que les fonctions d'inclusion moyenne et naturelle. Donc afin de réduire l'effet du pessimisme, la meilleure représentation est l'utilisation de la fonction d'inclusion de Taylor.

b) Propriétés des fonctions d'inclusion

Nous allons citer ci-dessous les propriétés suivantes:

b.1) Monotonie d'une fonction d'inclusion

La fonction d'inclusion $[f]$ est dite monotone au sens de l'inclusion [Moore, 1966] si la relation suivante est satisfaite :

$$[x] \subset [y] \Rightarrow [f]([x]) \subset [f]([y]) \quad (2.28)$$

b.2) Convergence d'une fonction d'inclusion

La fonction d'inclusion de $[f]$ est convergente [Moore, 1966], si pour toute suite $[x](k)$ nous avons:

$$\lim_{k \rightarrow \infty} w([x](k)) = 0 \Rightarrow \lim_{k \rightarrow \infty} w([f]([x](k))) = 0 \quad (2.29)$$

Pour un intervalle $[x]$ dégénéré, la fonction d'inclusion est égale à la fonction en elle-même

$$[f]([x]) = [f](x) = f(x) \quad (2.30)$$

b.3) Ordre de convergence

L'ordre de convergence de la fonction d'inclusion $[f]$ est le plus grand entier α tel que :

$$\exists \beta > 0 \mid (w([f]([x])) - w(f([x]))) \leq \beta w([x])^\alpha \quad (2.31)$$

- α est égale à l'infini si la fonction d'inclusion est minimale.
- $\alpha \geq 1$ si la fonction d'inclusion est naturelle (convergence linéaire).
- ($\alpha \geq 2$) si la fonction d'inclusion est moyenne ou de Taylor (convergence quadratique).

Exemple 2.8 [Jaulin et al., 2001]

L'exemple suivant donne un comparatif de l'évaluation des fonctions d'inclusion de la fonction $f = x^2 + \sin(x)$ dans deux intervalles de largeur différente :

$$[x_1] = \left[\frac{2\pi}{3} \quad \frac{4\pi}{3}\right] \quad \text{et} \quad [x_2] = \left[\frac{99\pi}{100} \quad \frac{101\pi}{100}\right]$$

avec $w([x_1]) > w([x_2])$ et le point x_o vérifie :

$$x_o = \text{mid}([x_1]) = \text{mid}([x_2])$$

Un comparatif de la normalisation qui définit la convergence d'une fonction d'inclusion est calculée, donnée par:

$$w([f]([x])) - w(f([x])) \quad (2.32)$$

Les résultats sont résumés dans le tableau suivant :

Le pavé	$[x_1] = [2\pi/3 \quad 4\pi/3]$	$[x_2] = [99\pi/100 \quad 101\pi/100]$
Fonction d'inclusion naturelle $[f_n]([x])$	[3.5204, 18.4120]	[9.6417, 10.0993]
Fonction d'inclusion moyenne $[f_m]([x])$	[1.6202, 18.1190]	[9.7016, 10.0375]
Fonction d'inclusion de Taylor $[f_T]([x])$	[4.3370, 16.9737]	[9.7036, 10.0365]
Fonction d'inclusion minimale $[f]^*([x])$	[5.2525, 16.6799]	[9.7046, 10.03657]
$w([f_n]([x])) - w(f([x]))$	3.4641	0.1256
$w([f_m]([x])) - w(f([x]))$	5.0713	0.0039
$w([f_T]([x])) - w(f([x]))$	1.2091	9.9212 10^{-4}
$w([f]^*([x])) - w(f([x]))$	0	0

Tableau 2.1 : Comparatif des fonctions d'inclusion de f sur deux pavés différents $[x_1]$ et $[x_2]$

D'après les résultats du tableau 2.1, nous remarquons que l'évaluation des fonctions d'inclusion dans $[x_2]$ sont meilleures que celles dans l'intervalle $[x_1]$, ceci est relié à la largeur de l'intervalle considéré. Donc pour un résultat meilleur, il faudra avoir une largeur de pavé petite ($w([x_2]) < w([x_1])$).

Et en comparant la taille (2.32), le résultat trouvé par l'utilisation de la fonction d'inclusion moyenne et de Taylor est meilleur que celui obtenu par la fonction d'inclusion naturelle, car la convergence quadratique est meilleure que la convergence linéaire.

2.3.5 Partitionnement de pavé

Le partitionnement de pavé est l'une des techniques pour réduire le pessimisme dû à l'effet d'enveloppement d'une fonction d'inclusion $[f]$ [Videau, 2005].

L'évaluation d'une fonction d'inclusion de f sur un sous-pavage de $[x]$ qui est donné par l'union des pavés $[x_i]$ de $[x]$ donné comme suit :

$$[x] = \bigcup_{i=1}^n [x_i] \quad (2.33)$$

est calculée par la relation suivante :

$$\bigcup_{i=1}^n [f]([x_i]) \quad (2.34)$$

De plus, si la fonction d'inclusion $[f]$ est monotone au sens de l'inclusion c'est-à-dire, elle vérifie :

$$[f]([x_i]) \subset [f]([x]), \forall i \in \{1, \dots, n\} \quad (2.35)$$

alors

$$[f]([x]) \supset \bigcup_{i=1}^n [f]([x_i]) \quad (2.36)$$

Le partitionnement du pavé $[x]$ donné par :

$$[x] = ([\underline{x}_1, \bar{x}_1] \times [\underline{x}_2, \bar{x}_2] \times \dots \times [\underline{x}_k, \bar{x}_k] \times \dots \times [\underline{x}_{n_x}, \bar{x}_{n_x}])$$

suivant la direction $k \in \{1, \dots, n\}$ permet de générer deux pavés :

$$[x_L] = ([\underline{x}_1, \bar{x}_1] \times [\underline{x}_2, \bar{x}_2] \times \dots \times [\underline{x}_k, \frac{\underline{x}_k + \bar{x}_k}{2}] \times \dots \times [\underline{x}_n, \bar{x}_n]) \quad (2.37)$$

$$[x_R] = ([\underline{x}_1, \bar{x}_1] \times [\underline{x}_2, \bar{x}_2] \times \dots \times [\frac{\underline{x}_k + \bar{x}_k}{2}, \bar{x}_k] \times \dots \times [\underline{x}_n, \bar{x}_n]) \quad (2.38)$$

La direction optimale k peut être déduite à partir du critère d'optimalité suivant :

$$k = \arg \left(\max_{i \in \{1, \dots, n_x\}} (D(i)) \right) \quad (2.39)$$

$D(i)$ est la fonction métrique qui est choisie selon plusieurs manières, nous trouvons par exemple dans [Moore, 1962].

$$D(i) = w([x]) \quad (2.40)$$

Exemple 2.9

Dans l'exemple (2.4) précédent, nous effectuons un partitionnement du pavé $[x]$ (2.14) par bisection de celui-ci suivant la direction $k = 2$, ceci donne :

$$[x] = \left(\begin{bmatrix} -1, 0 \\ 1, 2 \end{bmatrix} \right) = [x_1] \cup [x_2] = \left(\begin{bmatrix} -1, 0 \\ 1, \frac{1+2}{2} \end{bmatrix} \right) \cup \left(\begin{bmatrix} -1, 0 \\ \frac{1+2}{2}, 2 \end{bmatrix} \right)$$

Par la suite, nous évaluons la fonction d'inclusion (la fonction considérée est $[Y]$ donnée par (2.16)) sur les pavés $[x], [x_1]$ et $[x_2]$, nous obtenons les fonctions d'inclusion suivantes :

- $[f]([x]) = \left(\begin{bmatrix} 1, 4 \\ 1, 2 \end{bmatrix} \right)$
- $[f]([x_1]) = \left(\begin{bmatrix} 1, 3 \\ 1, 1.5 \end{bmatrix} \right)$
- $[f]([x_2]) = \left(\begin{bmatrix} 2, 4 \\ 1.5, 2 \end{bmatrix} \right)$

Les pavés trouvés sont représentés dans la figure ci-dessous :

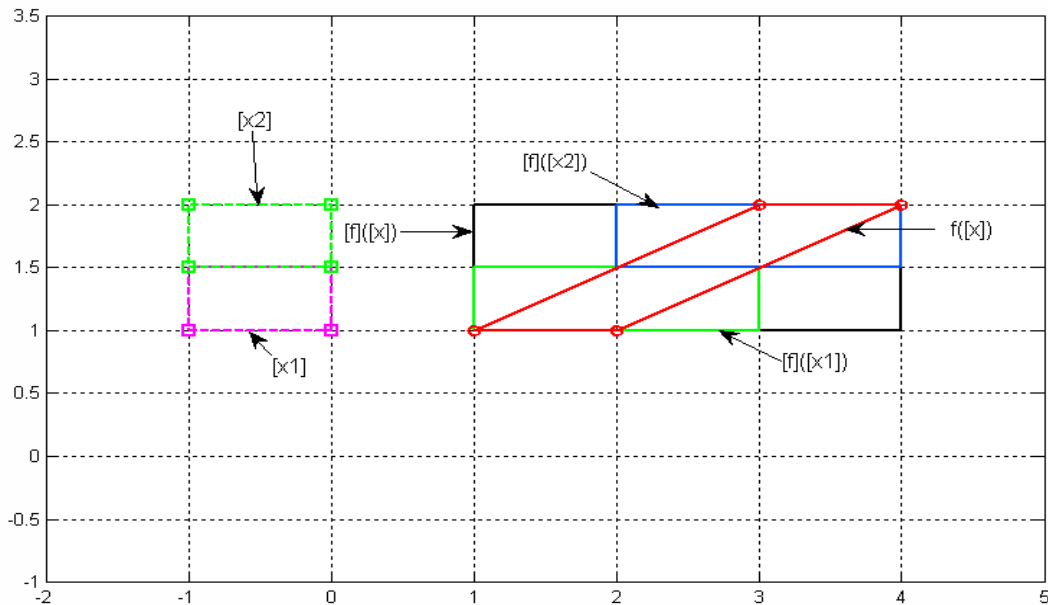


Figure2.2 Evaluation des fonctions $f([x]), [f]([x_1]), [f]([x_2])$ et $[f]([x])$

Donc, d'après les résultats obtenus, nous constatons que le partitionnement du pavé $[x]$ en $[x_1]$ et $[x_2]$ aide à réduire l'effet du pessimisme qui est, dans notre cas, dû à l'effet d'enveloppement de la représentation réelle $f([x])$ représentée par le parallélogramme en rouge dans la figure 2.2. Nous avons l'inclusion suivante qui est vérifiée :

$$f([x]) \subset ([f]([x_1]) \cup [f]([x_2])) \subset [f]([x]) \quad (2.41)$$

2.4 Algorithmes ensemblistes

Dans cette partie, nous allons citer quelques algorithmes ensemblistes utilisant comme outil le calcul par l'arithmétique des intervalles, comme l'inversion ensembliste par l'arithmétique des intervalles et les algorithmes s'appliquant aux contracteurs qui seront expliqués à travers des exemples.

2.4.1 Inversion ensembliste par l'arithmétique des intervalles

Soit une fonction vectorielle $f : R^{n_x} \rightarrow R^{n_y}$, l'image inverse d'un ensemble Y est donnée par l'ensemble suivant :

$$X = \{x \in X_0 \mid f(x) \in Y\} = f^{-1}(Y) \cap X_0 \quad (2.42)$$

Avec $X_0 \subset R^{n_x}$ est le domaine initial de recherche des solutions, l'ensemble $Y \subset R^{n_y}$ est connue à priori, sachant aussi que la fonction f peut ne pas être inversible au sens classique.

L'image inverse est réalisée grâce à plusieurs algorithmes, citons par exemple, l'algorithme SIVIA (Set Inversion Via Interval Analysis) [Jaulin et Walter, 1993] et [Walter et Jaulin, 1994]. Cet algorithme consiste en réalité à une approximation par encadrement de l'ensemble X (présentation ensembliste vue dans le premier chapitre).

La technique peut être expliquée de la manière suivante : partant de l'ensemble initial de recherche X_0 et en effectuant l'opération de partitionnement sur celui-ci, nous obtenons trois sortes de pavés : le premier est l'ensemble des pavés acceptables noté par $\underline{X}$ (l'encadrement intérieur de X), le second est l'ensemble des pavés inacceptables et le dernier est l'ensemble des pavés ambigus noté par ΔX .

Les pavés ambigus sont à leurs tours partitionnés jusqu'à n'avoir que des pavés acceptables. Une limite pour le sous pavage est définie car un pavé ne peut être partitionné à l'infini. Cette limite est appelée précision de sous pavages η .

Alors, l'encadrement de l'ensemble X est donné par la relation suivante:

$$\underline{X} \subset X \subset \underline{X} \cup \Delta X \quad (2.43)$$

avec $\underline{X} \cup \Delta X$ est l'approximation extérieure de l'ensemble X .

Test d'inclusion

C'est le test d'inclusion qui permet de conclure si un pavé est acceptable, inacceptable ou ambigu. A partir de la fonction d'inclusion $[f]([X])$ et pour un pavé initial $[x_0]$, nous avons :

1. $[f]([x_0]) \subset Y \Rightarrow [x_0] \subset X$: le pavé $[x_0]$ est acceptable.
2. $[f]([x_0]) \cap Y = \emptyset \Rightarrow [x_0] \cap X = \emptyset$: le pavé $[x_0]$ est inacceptable.
3. $[f]([x_0]) \cap Y \neq \emptyset$: le pavé $[x_0]$ est ambigu.

Exemple 2.10

Dans l'exemple suivant, nous cherchons à caractériser l'ensemble X par encadrement. L'ensemble est donné par :

$$X = \{x \in X_0 \mid x_1^2 + x_2^2 \in Y\} = f^{-1}(Y) \cap X_0$$

avec

$X_0 = (x_1, x_2)^T = [-2, 2], [-2, 2]^T$ dans R^2 est le pavé initial de recherche, Y est donné à priori par le pavé $[1, 2]$.

La caractérisation de l'ensemble X est un problème d'inversion ensembliste et sa résolution avec l'algorithme SIVIA grâce au simulateur réalisé par Luc Jaulin « simulateur : analyse par intervalle » (Annexe A) est donné par la figure suivante avec différentes précisions :

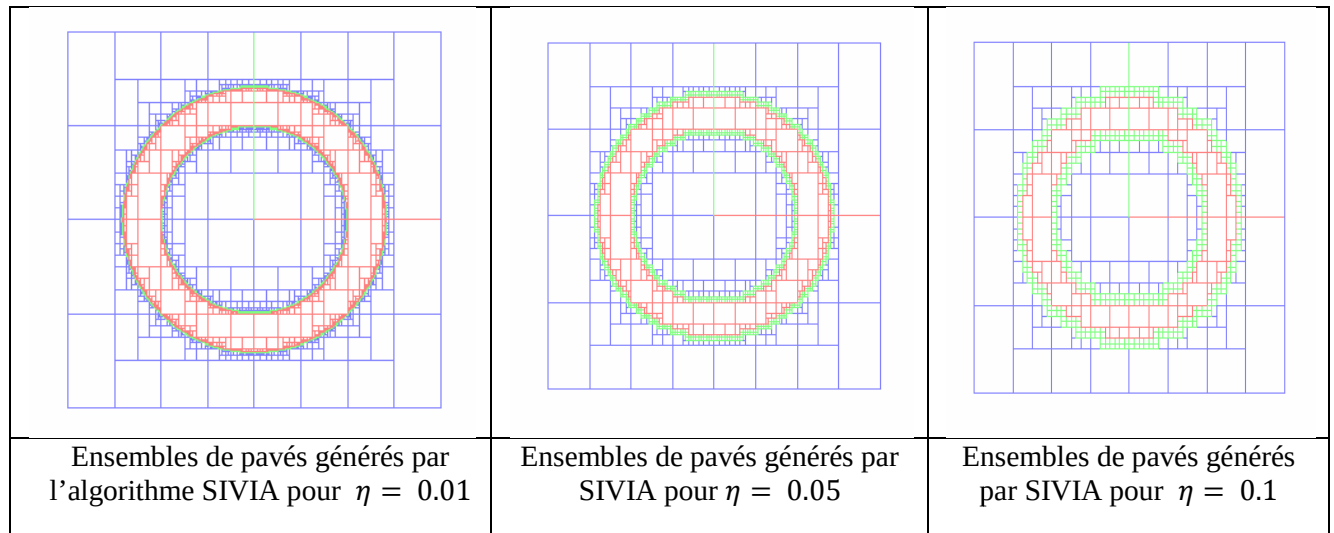


Figure 2.3 : Caractérisation de X par l'algorithme SIVIA

X est encadré par un encadrement intérieur qui est représenté par les pavés rouges dans la figure 2.3. L'encadrement extérieur est représenté par l'union des pavés ambigus de couleur verte et des pavés acceptables de couleur rouge.

Le tableau suivant donne le nombre de pavés générés par SIVIA pour trois précisions :

précision	Nombre total des pavés	Nombre des pavés solutions (rouge)	Nombre des pavés intermédiaires (vert)	Nombre des pavés exclus (bleu)
$\eta = 0.01$	6092	1060	3704	1328
$\eta = 0.05$	1468	212	920	336
$\eta = 0.1$	704	92	456	156

Tableau 2.2 : Nombre des pavés obtenus pour trois différentes précisions

Nous remarquons d'après les résultats du tableau et du tracé de l'ensemble X par la figure 2.3, que réduire la précision de sous pavage η implique forcément un nombre de pavés générés par le partitionnement qui est plus élevé et nous avons un meilleur encadrement de l'ensemble X .

2.4.2 Les contracteurs

L'outil de contraction ou de réduction est apparu à cause de la difficulté rencontrée lors des opérations de bisection d'un pavé, par exemple dans l'algorithme SIVIA pour un nombre de direction $k = 20$, on aura à faire $2^{20} = 1048576$ bisections. Cet outil permet de réduire le nombre de bisections à effectuer.

L'une des applications conçue pour cette technique est « INTERVAL PEELER » (voir Annexe A).

Définition 2.8 [Jaulin, 2004]

L'opérateur $C_X: IR^n \rightarrow IR^n$ est un contracteur pour le sous ensemble clos X de IR^n s'il satisfait les propriétés de contractance et de complétude suivantes :

$$\forall [x] \in IR^n, \begin{cases} C_X([x]) \subset [x] & (contractance) \\ C_X([x]) \cap X = [x] \cap X & (complétude) \end{cases} \quad (2.44)$$

Les propriétés suivantes sont définies :

- C_X est monotone si : $[x] \subset [y] \Rightarrow C_X([x]) \subset C_X([y])$.
- C_X est minimale si : $\forall [x] \in IR^n, C_X([x]) = [x] \cap X$.
- C_X est fin si : $\forall x \in R^n, C_X(\{x\}) = \{x\} \cap X$.
- C_X est idempotent si : $\forall [x] \in IR^n, C_X(C_X([x])) = C_X([x])$
- C_X est dit convergent si : pour tous points x et pour toute suite $[x](k)$ on a :

$$[x](k) \rightarrow x \Rightarrow C_X([x](k)) \rightarrow \{x\} \cap X$$

a) Projection de contraintes

L'objectif et le principe de cet outil sont donnés par l'exemple ci-dessous qui illustre la définition de valeurs consistantes et inconsistantes.

Exemple 2.11

Soient trois variables x, y et z appartenant aux domaines qui sont des intervalles $[1,5]$, $[2,4]$ et $[6,10]$ respectivement et soit la contrainte suivante : $z = x + y$

La valeur de $z = 10$ est **inconsistante** car $\forall x \in [1,5], \forall y \in [2,4], x + y \neq 10$ et la valeur 1 pour x par exemple est aussi **inconsistante** car $\forall z \in [6,10], \forall y \in [2,4], 1 + y \neq z$

L'objectif de la projection de contrainte est le calcul du plus petit pavé qui contient toutes les valeurs consistantes.

b) Quelques algorithmes de projection de contraintes

Afin de projeter un ensemble donné, nous considérons ci-dessous quelques algorithmes qui peuvent être utilisés :

L'algorithme PPLUS(sorties : $[x], [y], [z]$)

$$\begin{aligned} [z] &:= [z] \cap ([x] + [y]); \\ [x] &:= [x] \cap ([z] - [y]); \\ [y] &:= [y] \cap ([z] - [x]) \end{aligned} \quad (2.45)$$

L'algorithme PMULT(sorties : $[x], [y], [z]$)

$$\begin{aligned} [z] &:= [z] \cap ([x] * [y]); \\ [x] &:= [x] \cap \left([z] * \frac{1}{[y]} \right); \\ [y] &:= [y] \cap \left([z] * \frac{1}{[x]} \right); \end{aligned} \quad (2.46)$$

L'algorithme PEXP(sorties : $[x], [y], [z]$)

$$\begin{aligned} [y] &:= [y] \cap (\exp([x])); \\ [x] &:= [x] \cap (\log(y)); \end{aligned} \quad (2.47)$$

Pour l'exemple 2.11 l'ensemble S à projeter est donné par l'ensemble suivant :

$$S = \{(x, y, z) \in [1,5] \times [2,4] \times [6,10] \mid z = x + y\} \quad (2.48)$$

La méthode pour la projection de S basée sur « l'algorithme PPLUS » est donnée ci-dessous :

$$z = x + y \Rightarrow z \in [6,10] \cap ([1,5] + [2,4]) = [6,10] \cap [3,9] = [6,9]$$

$$x = z - y \Rightarrow x \in [1,5] \cap ([6,10] - [2,4]) = [1,5] \cap [2,8] = [2,5]$$

$$y = z - x \Rightarrow y \in [2,4] \cap ([6,10] - [1,5]) = [2,4] \cap [1,9] = [2,4]$$

c) Propagation de contraintes

Afin d'illustrer la technique de propagation de contraintes nous donnons en premier lieu la définition d'un CSP (problème de satisfaction de contraintes).

Définition 2.9

Un problème de satisfaction de contraintes (CSP) est composé :

- D'un ensemble de variables : $v = \{x_1, x_2, \dots, x_n\}$
- D'un ensemble de contraintes : $C = \{C_1, C_2, \dots, C_m\}$
- D'un ensemble de domaines d'intervalles : $\{[x_1], [x_2], \dots, [x_n]\}$

c.1) Principe de la technique

Le but de la technique de propagation de contraintes est de contracter les domaines des variables le plus possible sans perdre aucune solution.

Le pavé $[x]$ est défini par le produit cartésien de tous ses intervalles et par son **intersection répétée** avec les contraintes:

$$[x] \sqcap C_j \quad (2.49)$$

avec, $\sqcap$ est l'intersection répétée (square intersection) et elle est commutative.

Le principe de contraction du domaine $[x]$ est donné par :

$$((((([x] \sqcap C_1) \sqcap C_2) \sqcap \dots \sqcap C_m) \sqcap C_1) \sqcap C_2) \quad (2.50)$$

jusqu'à l'obtention d'un pavé régulier appelé « point fixe ».

Etant donné n variables $x_1, x_2, \dots, x_n$ liées à m contraintes $C_1, C_2, \dots, C_m$, nous supposons que chaque variable $x_i (i = 1 \dots n)$ est associée au domaine $]-\infty, +\infty[$ et ceci si les variables ne présentent pas d'informations.

Le principe de propagation de contraintes d'intervalles (ICP :Interval Constraint Propagation) est d'exécuter la procédure de projection pour toutes les contraintes et cette opération est répétée jusqu'à ce qu'il ne reste plus de contrainte significative à exécuter.

Nous illustrons cette technique par l'exemple suivant :

Exemple 2.12

Soient les trois contraintes suivantes :

$$\begin{aligned}(C_1): & y = x^2 \\(C_2): & xy = 1 \\(C_3): & y = -2x + 1\end{aligned}$$

Le domaine $]-\infty, +\infty[$ est associé à chaque variable. La propagation consiste à une projection de toutes les contraintes jusqu'à **l'équilibre**

Les résultats sont donnés ci-dessous :

$$\begin{aligned}(C_1): & y \in]-\infty, +\infty[^2 = [0, +\infty[\\(C_2): & x \in \frac{1}{[0, +\infty[} = [0, +\infty[\\(C_3): & y \in [0, +\infty[\cap (-2[0, +\infty[+ 1) = [0, +\infty[\cap]-\infty, 1] = [0, 1] \\& x \in [0, +\infty[\cap \left(\frac{-1}{2}\right)[0, 1] + \left(\frac{1}{2}\right) = \left[0, \frac{1}{2}\right] \\(C_1): & y \in [0, 1] \cap \left[0, \frac{1}{2}\right]^2 = \left[0, \frac{1}{4}\right] \\(C_2): & x \in \left[0, \frac{1}{2}\right] \cap \frac{1}{\left[0, \frac{1}{4}\right]} = \phi \\& y \in \left[0, \frac{1}{4}\right] \cap \frac{1}{\phi} = \phi\end{aligned}$$

Donc, nous sommes arrivés au point fixe qui est l'ensemble vide ϕ dans notre cas, ou nous ne pouvons plus continuer la propagation.

c.2) Problème de consistance locale

Si on a deux contracteurs minimaux $C_{S_1}^*$, $C_{S_2}^*$ des deux ensembles S_1 et S_2 respectivement alors la relation suivante :

$$C_S = C_{S_1}^* \circ C_{S_2}^* \circ C_{S_1}^* \circ C_{S_2}^* \circ \dots \quad (2.51)$$

représente le contracteur de l'ensemble $S_1 \cap S_2$, mais ce contracteur peut être non optimal. Ce phénomène est connu sous le nom de consistance locale.

Nous illustrons cet effet par l'exemple suivant :

Exemple 2.13

Soient les deux CSP suivants :

$$CSP1: \begin{pmatrix} x_1 + x_2 = 0 \\ x_1 - x_2 = 0 \\ x_1 \in [-10,10] \\ x_2 \in [-10,10] \end{pmatrix}$$

$$CSP2: \begin{pmatrix} x_1 + x_2 = 0 \\ x_1 - 2x_2 = 0 \\ x_1 \in [-10,10] \\ x_2 \in [-10,10] \end{pmatrix}$$

La contraction du problème $CSP1$ est donnée par :

$$\begin{cases} x_1 = x_1 \sqcap (-x_2) \\ x_2 = x_2 \sqcap (-x_1) \\ x_1 = x_1 \sqcap x_2 \\ x_2 = x_2 \sqcap x_1 \end{cases}$$

Avec

$\sqcap$: est l'intersection répétée.

Cette contraction donne toujours les deux domaines $x_1 \in [-10,10]$ et $x_2 \in [-10,10]$, nous concluons alors que le contracteur est non optimal. Donc nous avons le phénomène de consistance locale qui se produit pour ce CSP.

Or pour le problème $CSP2$, la contraction est donnée par :

$$\begin{cases} x_1 = x_1 \sqcap (-x_2) \\ x_2 = x_2 \sqcap (-x_1) \\ x_1 = x_1 \sqcap (2x_2) \\ x_2 = x_2 \sqcap (\frac{1}{2}x_1) \end{cases}$$

Les résultats de la contraction pour une intersection répétée de 20 itérations sont donnés ci-dessous :

$$\begin{aligned} x_1 &\in [-0.1908 \cdot 10^{-4}, 0.1908 \cdot 10^{-4}] \\ x_2 &\in [-0.9537 \cdot 10^{-5}, 0.9537 \cdot 10^{-5}] \end{aligned}$$

c.3) Solutions apportées au problème de consistance locale

Des solutions au problème de consistance locale sont citées ci-dessous :

▪ **L'ajout de contrainte redondante**

Lors de la présence de l'effet de consistance locale dans un CSP donné, la propagation des contraintes est incapable de le contracter. Afin d'éviter ceci, des contraintes redondantes sont ajoutées au problème. Ces contraintes peuvent être des compositions des contraintes du problème principal.

- **Décomposition en contraintes primitives**

Pour des contraintes complexes, la décomposition en contraintes primitives s'avère nécessaire pour la contraction, l'exemple suivant montre la manière dont les contraintes sont décomposées en contraintes primitives.

Exemple 2.14

Soit le CSP suivant :

$$CSP: \begin{cases} x \in [-1,1] \\ y \in [-1,1] \\ z \in [-1,1] \\ x + \sin(y) - xz \leq 0 \end{cases}$$

La décomposition en contraintes primitives de la contrainte du CSP donne :

$$\begin{cases} a = \sin(y) & y \in [-1,1] , & a \in]-\infty, +\infty[\\ b = x + a & x \in [-1,1] , & b \in]-\infty, +\infty[\\ c = x * z & z \in [-1,1] , & c \in]-\infty, +\infty[\\ d = b - c & & d \in]-\infty, 0] \end{cases}$$

d) Infaillibilité d'un problème de satisfaction de contraintes

Un CSP est dit infaillible, si l'ensemble solution est le produit cartésien de tous ses domaines. Pour montrer qu'un CSP est infaillible, il suffit de montrer que sa négation est un ensemble vide comme le montre l'exemple suivant :

Exemple 2.15

Soit le CSP suivant :

$$csp: \begin{cases} \mathcal{V} = \{x, y\} \\ \mathcal{D} = \{[x], [y]\} \\ C = \{f(x, y) \leq 0, g(x, y) \leq 0\} \end{cases}$$

Le CSP est infaillible si

$$\forall x \in [x], \forall y \in [y], f(x, y) \leq 0 \text{ et } g(x, y) \leq 0$$

C'est-à-dire sa négation va vérifier:

$$\{(x, y) \in [x] \times [y], |f(x, y) > 0 \text{ ou } g(x, y) > 0\} = \emptyset$$

Ou bien avoir:

$$\{(x, y) \in [x] \times [y], |\max(f(x, y), g(x, y)) > 0\} = \emptyset$$

Cette procédure est accomplie efficacement en utilisant la technique de propagation de contraintes.

e) Algorithme de propagation-rétropropagation

Cet algorithme sélectionne des contraintes primitives pour effectuer la contraction dans un ordre optimal. La contraction est faite selon la taille des domaines obtenus.

Exemple 2.16

Soit le CSP suivant :

$$CSP \left\{ \begin{array}{l} \mathcal{V} = \{x_1, x_2, x_3\} \\ \mathcal{D} = \{[x_1], [x_2], [x_3]\} \\ \mathcal{C} = \{f(x) = x_1 \exp(x_2) + \sin(x_3) \in [y]\} \end{array} \right.$$

Les étapes de l'algorithme de propagation-rétropropagation sont les suivantes :

La 1^{ère} étape correspond à l'écriture du CSP sous forme de contraintes primitives.

La 2^{ème} étape est l'écriture des contraintes dans un contexte ensembliste (c'est-à-dire le remplacement des variables par leurs domaines primaires) et en même temps effectuer la propagation des contraintes définie auparavant.

La 3^{ème} étape consiste à effectuer la procédure de rétropropagation.

Nous aurons alors :

1^{ère} étape :

$$\begin{aligned} a_1 &:= \exp(x_2) \\ a_2 &:= x_1 a_1 \\ a_3 &:= \sin(x_3) \\ y &:= a_2 + a_3 \end{aligned}$$

2^{ème} étape :

$$\begin{aligned} 1. [a_1] &:= \exp([x_2]) \\ 2. [a_2] &:= [x_1] [a_1] \\ 3. [a_3] &:= \sin([x_3]) \\ 4. [y] &:= [y] \cap ([a_2] + [a_3]) \end{aligned}$$

3^{ème} étape :

Si $[y]$ obtenu à l'étape 4 de la deuxième étape est un ensemble vide, alors le CSP ne possède pas de solution. Il faudra donc effectuer la procédure de rétropropagation. Cette dernière consiste à calculer les différents domaines du CSP en partant de l'étape 4 en arrière comme suit :

5. $[a_2] := ([y] - [a_3]) \cap [a_2]$
6. $[a_3] := ([y] - [a_2]) \cap [a_3]$
7. $[x_3] := \sin^{-1}([a_3]) \cap [x_3]$
8. $[a_1] := ([a_2]/[x_1]) \cap [a_1]$
9. $[x_1] := ([a_2]/[a_1]) \cap [x_1]$
10. $[x_2] := \log([a_1]) \cap [x_2]$

f) Quelques applications des contracteurs

Nous allons présenter dans cette section deux exemples d'utilisation des contracteurs, le premier exemple est appliqué à la stabilité robuste et le deuxième exemple à l'estimation à erreurs bornées.

f.1) Exemple d'application à la stabilité robuste

L'exemple suivant [Jaulin, 2004], illustre l'application des contracteurs à la stabilité robuste.

Le problème consiste en une moto roulant à une vitesse constante et faible de $1m/s$. L'entrée du système est l'angle θ du guidon et la sortie est l'angle de roulis ϕ de la moto.

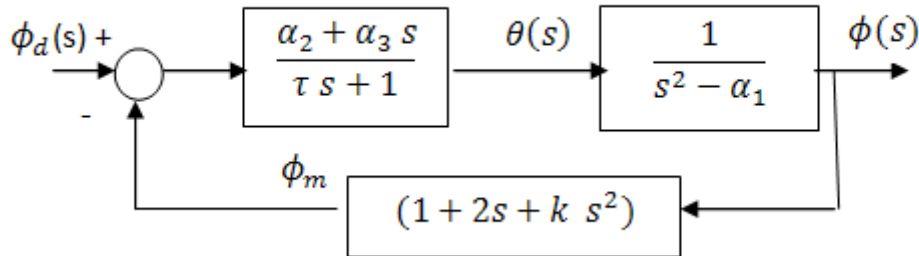


Figure 2. 4 : Système corrigé en boucle fermée

La fonction de transfert du système en boucle ouverte est $\phi(s) = \frac{1}{s^2 - \alpha_1} \theta(s)$

La vitesse de la moto étant faible, alors l'effet gyroscopique qui maintient la moto stable à vive allure a été négligé et le système est donc instable.

Afin de commander l'angle de guidon, un régulateur de 1^{er} ordre est inséré, sa fonction de transfert est la suivante :

$$\theta(s) = \frac{\alpha_2 + \alpha_3 s}{\tau s + 1} (\phi_d(s) - \phi_m(s))$$

Où :

ϕ_d : Angle de roulis désiré.

ϕ_m : Angle de roulis mesuré.

L'angle de roulis mesuré ϕ_m n'est pas exactement l'angle réel ϕ de la moto, il est relié à celui-ci par la relation suivante :

$$\phi_m(s) = (1 + 2s + k s^2) \phi(s)$$

Donc, ϕ_m est une composition d'angle, de vitesse et d'accélération.

La fonction de transfert en boucle fermée du système corrigé est donné par :

$$\phi(s) = \frac{\alpha_2 + \alpha_3 s}{(s^2 - \alpha_1)(\tau s + 1) + (\alpha_2 + \alpha_3 s)(1 + 2s + k s^2)} \theta(s)$$

Les coefficients du capteur et du régulateur sont donnés sous forme de domaines :

$$\alpha_1 \in [8.8, 9.2], \alpha_2 \in [2.8, 3.2], \alpha_3 \in [0.8, 1.2], \tau \in [1.8, 2.2], k \in [-3.2, -2.8]$$

Le polynôme caractéristique du système est le suivant :

$$(s^2 - \alpha_1)(\tau s + 1) + (\alpha_2 + \alpha_3 s)(1 + 2s + k s^2) = a_3 s^3 + a_2 s^2 + a_1 s + a_0$$

avec :

$$a_3 = \tau + \alpha_3 k, a_2 = \alpha_2 k + 2\alpha_3 + 1, a_1 = \alpha_3 - \alpha_1 \tau + 2\alpha_2 \text{ et } a_0 = -\alpha_1 + \alpha_2$$

La table de Routh est donnée ci-dessous :

a_3	a_1
a_2	a_0
$b = \frac{a_2 a_1 - a_3 a_0}{a_2}$	0
a_0	0

Tableau 2.3 : Table de Routh

Le système est dit stable si $a_3, a_2, \frac{a_2 a_1 - a_3 a_0}{a_2}$ et a_0 ont le même signe.

Soit quatre nombres b_1, b_2, b_3, b_4 . Ces quatre nombres sont de même signe s'ils vérifient :

$$\begin{cases} \min(b_1, b_2, b_3, b_4) > 0 \\ \text{ou} \\ \max(b_1, b_2, b_3, b_4) < 0 \end{cases}$$

ou bien :

$$\max(\min(b_1, b_2, b_3, b_4), -\max(b_1, b_2, b_3, b_4)) > 0 \quad (2.52)$$

Dans l'exemple de la moto, en utilisant la propriété de l'infailibilité du CSP, il revient à démontrer que la négation de (2.52) est toujours fausse. C'est-à-dire que la stabilité robuste du système est vérifiée si :

$$\exists \alpha_1 \in [\alpha_1], \exists \alpha_2 \in [\alpha_2], \exists \alpha_3 \in [\alpha_3], \exists \tau \in [\tau], \exists k \in [k]$$

$$\max(\min(a_3, a_2, a_0, b), -\max(a_3, a_2, a_0, b)) \leq 0 \quad (2.53)$$

est toujours fausse.

Le CSP de la stabilité robuste se traduit par :

$$CSP: \left\{ \begin{array}{l} v = \{a_0, a_1, a_2, a_3, b, \alpha_1, \alpha_2, \alpha_3, \tau, k\} \\ C1: a_3 = \tau + \alpha_3 k \\ C2: a_2 = \alpha_2 k + 2 \alpha_3 + 1 \\ C3: a_1 = \alpha_3 - \alpha_1 \tau + 2 \alpha_2 \\ C4: a_0 = -\alpha_1 + \alpha_2 \\ C5: b = \frac{a_2 a_1 - a_3 a_0}{a_2} \\ C6: \max(\min(a_3, a_2, a_0, b), -\max(a_3, a_2, a_0, b)) \leq 0 \\ D = \{[a_0], [a_1], [a_2], [a_3], [b], [\alpha_1], [\alpha_2], [\alpha_3], [\tau], [k]\} \end{array} \right. \quad (2.54)$$

La solution du CSP est facilement résolu sur INTLAB qui nous donne un ensemble vide. Donc la stabilité robuste du système est vérifiée, comme nous pouvons aussi utiliser INTERVAL PEELER (Annexe A).

f.2) Application à l'estimation à erreur bornée

Soit le circuit électrique représenté par la figure ci-dessous [Jaulin et Chabert, 2008] :

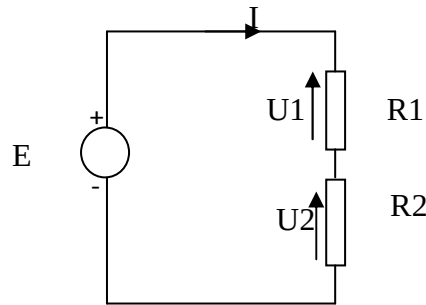


Figure 2.5 : Circuit électrique comprenant une pile et deux résistances en séries.

Ce circuit est composé d'une pile délivrant une tension E , de deux résistances $R1, R2$ en séries dont les tensions aux bornes sont respectivement $U1, U2$. Désignons par I le courant circulant dans le circuit et par P la puissance délivrée par la pile.

Des mesures ont été prises pour chaque variable. Ces mesures se représentent par des relations d'appartenance à des intervalles. Nous avons alors les relations suivantes :

$$E \in [23V, 26V], I \in [4A, 8A], U1 \in [10V, 11V], U2 \in [14V, 17V], P \in [124w, 130w]$$

Aucune information connue concernant les résistances $R1, R2$, mais elles sont prises positives. Alors leur domaine d'appartenance se traduit par l'intervalle $[0, +\infty[$

Les contraintes du CSP, que nous allons donner par la suite, sont prises en tant que les relations électriques reliant les différents composants.

Alors le CSP se traduit par :

$$CSP: \left\{ \begin{array}{l} v = \{E, I, U1, U2, P, R1, R2\} \\ D = \{[23V, 26V], [4A, 8A], [10V, 11V], [14V, 17V], [124watts, 130watts], [0, +\infty[, [0, +\infty[\} \\ C: \left\{ \begin{array}{l} C1: P = E * I \\ C2: E = (R1 + R2) * I \\ C3: U1 = R1 * I \\ C4: U2 = R2 * I \\ C5: E = U1 + U2 \end{array} \right. \end{array} \right. \quad (2.55)$$

Dans ce CSP, nous remarquons que la contrainte $C2$ est redondante, car celle-ci peut être retrouvée en utilisant la contrainte $C3, C4$ et $C5$, ce qui rend la résolution du CSP plus performant.

La solution du CSP est résolue en utilisant l'algorithme de propagation suivant :

1. PPLUS(sorties : [R], [R1], [R2])

1. $[R] := [R] \cap ([R1] + [R2])$;
2. $[R1] := [R1] \cap ([R] - [R2])$;
3. $[R2] := [R2] \cap ([R] - [R1])$;

2. PMULT(sorties : [P], [E], [I])

4. $[P] := [P] \cap ([E] * [I])$;
5. $[I] := [I] \cap \left([P] * \frac{1}{[E]}\right)$;
6. $[E] := [E] \cap \left([P] * \frac{1}{[I]}\right)$;

3. PMULT(sorties : [E], [R], [I])

7. $[E] := [E] \cap ([R] * [I])$;
8. $[I] := [I] \cap \left([E] * \frac{1}{[R]}\right)$;
9. $[R] := [R] \cap \left([E] * \frac{1}{[I]}\right)$;

4. PMULT(sorties : [U2], [R2], [I])

10. $[U2] := [U2] \cap ([R2] * [I])$;
11. $[I] := [I] \cap \left([U2] * \frac{1}{[R2]}\right)$;
12. $[R2] := [R2] \cap \left([U2] * \frac{1}{[I]}\right)$;

5. PMULT(sorties : [U1], [R1], [I])
 13. [U1] := [U1] $\cap$ ([R1] * [I]);
 14. [I] := [I] $\cap$ ([U1] * $\frac{1}{[R1]}$);
 15. [R1] := [R1] $\cap$ ([U1] * $\frac{1}{[I]}$);
6. PPLUS(sorties : [E], [U1], [U2])
 16. [E] := [E] $\cap$ ([U1] + [U2]);
 17. [U1] := [U1] $\cap$ ([E] - [U2]);
 18. [U2] := [U2] $\cap$ ([E] - [U1]);

Les résultats obtenus sur INTLAB pour les variables après contraction sont les suivants :

$$\begin{aligned}
 R1 &= [\quad 1.7692 \quad \Omega, \quad 2.3065 \quad \Omega] \\
 R2 &= [\quad 2.4769 \quad \Omega, \quad 3.5646 \quad \Omega] \\
 I &= [\quad 4.7692 \quad A, \quad 5.6522 \quad A] \\
 U1 &= [\quad 10.0000 \quad V, \quad 11.0000 \quad V] \\
 U2 &= [\quad 14.0000 \quad V, \quad 16.0000 \quad V] \\
 E &= [\quad 24.0000 \quad V, \quad 26.0000 \quad V] \\
 P &= [\quad 124.0000 \quad w, \quad 130.0000 \quad w]
 \end{aligned}$$

Ainsi, les domaines obtenus pour les deux résistances $R1$ et $R2$ sont réduits et nous avons aussi une contraction pour les variables $I, U2$ et E pourtant celles-ci sont accessibles à la mesure.

Nous n'avons pas de changement pour P et $U1$.

2.5 Résolution des systèmes algébriques linéaires à intervalles

Dans cette partie, nous allons illustrer les différentes méthodes pour la résolution des systèmes d'équations lineaires à intervalles. Nous définissons pour cet objet la matrice intervalle, la matrice de comparaison, la M-matrice et la H-matrice dont nous avons besoin dans les différentes méthodes.

Par la suite, une méthode de résolution dans INTLAB est présentée et un exemple d'un système algebrique lineaire à intervalles est résolu en utilisant les differentes méthodes .

Définitions 2.10 [Hargreaves, 2002]

- Une matrice intervalle est une matrice dont les éléments sont des intervalles. L'ensemble de toutes les matrices intervalles sont définies dans l'espace $IR^{n \times n}$.
- Une matrice de comparaison d'une matrice A notée $\langle A \rangle$ est obtenue en remplaçant les éléments de la diagonale de la matrice A par :

$$\langle A \rangle_{ii} = \min\{|\alpha|, \alpha \in A_{ii}\} \quad (2.56)$$

et les éléments hors diagonale par :

$$\langle A \rangle_{ik} = -\max\{|\alpha|, \alpha \in A_{ik}\}, k \neq i \quad (2.57)$$

- Une matrice A est dite une M-matrice si et seulement si la relation suivante est vérifiée :

$$\forall i \neq j, A_{ij} \leq 0 \text{ et } \exists u \in \mathbb{R}^n | A u > 0 \quad (2.58)$$

- Une matrice A est dite une H-matrice si sa matrice de comparaison est une M-matrice.

2.5.1 Définition d'un système algébrique linéaire à intervalles

Un système algébrique linéaire à intervalles est sous la forme :

$$Ax = b \quad (2.59)$$

Avec

A est une matrice intervalle dans l'espace $\mathbb{IR}^{n \times n}$.

b est un vecteur intervalle dans $\mathbb{IR}^n$.

L'ensemble solution du système algébrique linéaire à intervalles est le suivant :

$$\Sigma(A, b) = \{ \tilde{x} : \tilde{A} \tilde{x} = \tilde{b} \text{ pour } \tilde{A} \in A; \tilde{b} \in b \} \quad (2.60)$$

Exemple 2.17

Soit le système algébrique linéaire à intervalles suivant $Ax = b$

avec

$$A = \begin{pmatrix} [1,3] & [-1,1] \\ [-1,0] & [2,4] \end{pmatrix} \text{ et } b = \begin{pmatrix} [-2,2] \\ [-2,2] \end{pmatrix} \quad (2.61)$$

L'ensemble solution $\Sigma(A, b)$ est représenté dans la figure 2.6 par l'ensemble en rouge et l'union de tous les ensembles solution qui est noté par $\cup \Sigma(A, b)$ est représenté par le pavé en trait bleu discontinu. Cet ensemble encadre la solution réelle.

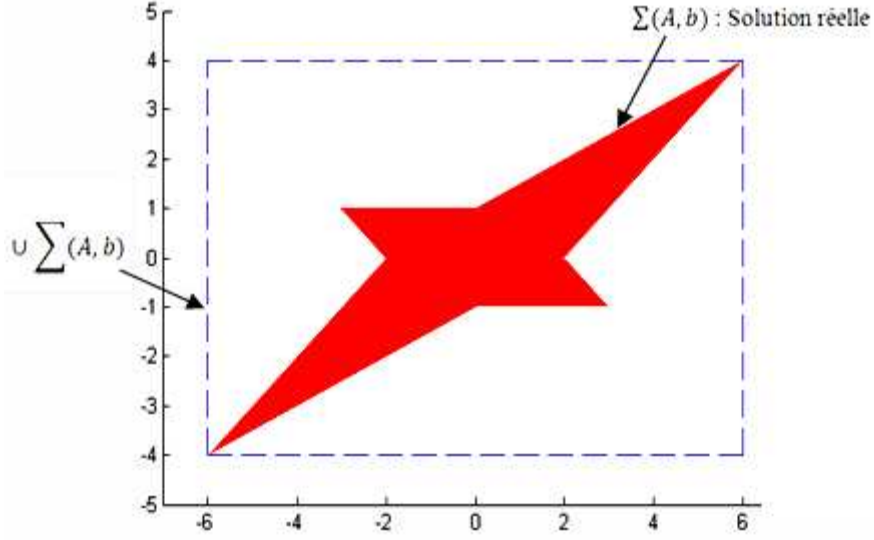


Figure 2.6 : Exemple montrant l'encadrement extérieur de la solution réelle $\Sigma(A, b)$ par l'ensemble $\cup \Sigma(A, b)$ qui est un vecteur intervalle.

2.5.2 Méthodes de résolution des systèmes algébriques linéaires à intervalles

a) Méthode de Krawczyk [Hargreaves, 2002]

Le système algébrique linéaire à intervalles $Ax = b$ peut être résolu en prenant comme pré-condition que celui-ci est multiplié par une matrice $C \in R^{n \times n}$. Celle-ci est prise comme l'inverse de la matrice $(mid(A))$ avec le produit des deux matrices C et A forme une H-matrice.

Nous supposons qu'un vecteur intervalle $x^{(i)}$ est connu tel que :

$$\cup \Sigma(A, b) \subseteq x^{(i)}$$

Alors

$$\tilde{A}^{-1}\tilde{b} = C\tilde{b} + (I - C\tilde{A})\tilde{A}^{-1}\tilde{b} \in Cb + (I - CA)x^{(i)}$$

$\forall \tilde{A} \in A$ et $\forall \tilde{b}$ on a :

$$\cup \Sigma(A, b) \subseteq x^{(i)} \Rightarrow \cup \Sigma(A, b) \subseteq ((Cb + (I - CA)x^{(i)}) \cap x^{(i)}) \quad (2.62)$$

Ce qui donne l'itération de **Krawczyk** :

$$x^{(i+1)} = (Cb + (I - CA)x^{(i)}) \cap x^{(i)} \quad (2.63)$$

Pour commencer l'itération, nous avons besoin d'un vecteur initial $x^{(0)}$ tel que la solution $\tilde{x} \in x^{(0)}$ et $x^{(0)} \supseteq (\cup \Sigma(A, b))$.

Le vecteur intervalle initial $x^{(0)}$ peut être choisi de la manière suivante :

$$x^{(0)} = ([-\alpha, \alpha], \dots, [-\alpha, \alpha])^T \quad (2.64)$$

Avec, $\alpha = \frac{\|C A\|_{\infty}}{1-\beta}$ et $\beta = \|I - C \tilde{A}\| < 1$

Le critère d'arrêt pour cette méthode est basé sur l'évolution de la somme des rayons des éléments du vecteur intervalle $x^{(i)}$. Dès que cette somme n'évolue plus ou diminue lentement, on arrête le critère.

La somme de ces rayons peut être calculée après chaque itération et comparée à la somme précédente. Cette somme est donnée par $S = \text{sum}(\text{rad}(x^{(i)}))$. Initialement, elle prend la valeur de $S_0 = \infty$ et nous définissons un facteur $M = (1 + \beta)/2$.

Le critère d'arrêt est alors donné comme suit :

Tant que $S < M S_0$

$$x^{(i+1)} = (Cb + (I - CA)x^{(i)}) \cap x^{(i)}$$

$$S_0 := S$$

$$S := \text{sum}(\text{rad}(x^{(i)}))$$

Cette méthode est donnée par la fonction `kraw.m` développée sur INTLAB (Annexe B).

b) Méthode de Hansen-Bliek-Rohn-Ning-Kearfott-Neumaier [Hargreaves, 2002]

L'ensemble solution $\cup \Sigma(A, b)$ du système algébrique linéaire à intervalles donné par Hansen est calculé dans le cas où $\text{mid}(A)$ est une matrice identité. Ce résultat a été aussi retrouvé par Rohn et Bliek.

Ning et Kearfott généralisent le résultat dans le cas où A est une H-matrice.

La méthode est basée sur le théorème [Neumaier, 1999] ci-dessous qui utilise la matrice de comparaison $\langle A \rangle$.

Théorème 2.1 [Neumaier, 1999]

Soient $A \in \mathbb{IR}^{n \times n}$ une H-matrice et b un vecteur intervalle dans $\mathbb{IR}^n$:

$$u = \langle A \rangle^{-1} |b| \text{ et } d_i = (\langle A \rangle^{-1})_{ii}$$

et

$$\alpha_i = \langle A \rangle_{ii} - \frac{1}{d_i}, \beta_i = \frac{u_i}{d_i} - |b_i|$$

Alors, $\cup \Sigma(A, b)$ est incluse dans le vecteur x dont les éléments sont donnés comme suit :

$$x^{(i)} = \frac{b_i + [-\beta_i, \beta_i]}{A_{ii} + [-\alpha_i, \alpha_i]} \quad (2.65)$$

Cette méthode est donnée par la fonction `hsolve.m` (Annexe B).

c) Résolution des systèmes algébriques linéaires à intervalles par la fonction « `verifylss` » de la toolbox INTLAB [Hargreaves, 2002]:

INTLAB fournit la fonction "**verifylss**" [Hargreaves, 2002] pour résoudre les systèmes algébriques linéaires à intervalle ou pour donner des bornes à la solution d'un système ponctuel. Comme nous trouvons aussi la commande `\` qui est équivalente à faire appel à la fonction **verifylss**.

La solution d'un système linéaire à intervalles est résolue avec la méthode itérative basée sur l'opérateur de Krawczyk, mais la méthode est différente de celle décrite auparavant. En utilisant comme pré-condition la matrice C qui est égale à :

$$C = inv(mid(A))$$

Une solution approximative x_s est calculée par la multiplication de la matrice C par le vecteur $mid(b)$.

En appliquant l'équation (2.62) à l'encadrement $d^{(i)}$ de $\cup \Sigma(A, b - Ax_s)$ donne l'itération résiduelle de Krawczyk suivante :

$$d^{(i+1)} = (C(b - Ax_s) + (I - C A)d^{(i)}) \cap d^{(i)} \quad (2.66)$$

Exemple 2.18

Dans l'exemple suivant, nous disposons du système algébrique linéaire à intervalles $Ax = b$
Avec :

$$A = \begin{pmatrix} [2,4] & [-1,0] \\ [0,1] & [1,2] \end{pmatrix} \text{ et } b = \begin{pmatrix} [-2,5] \\ [1,2] \end{pmatrix}$$

La solution du système algébrique linéaire à intervalles, en utilisant les fonctions « verifylss », « kraw » and « hsolve » décrites auparavant est donnée par le tableau suivant :

Méthode	Krawczyk « kraw »	Hansen-Bliek-Rohn-Ning-Kearfott-Neumaier « hsolve »	« Verifylss »
Solution	$x = \begin{pmatrix} [-3.6474, 4.7465] \\ [-3.1052, 6.1843] \end{pmatrix}$	$x = \begin{pmatrix} [-3.4697, 4.7328] \\ [-3.9696, 5.5485] \end{pmatrix}$	$x = \begin{pmatrix} [-2.0715, 4.6429] \\ [-2.2143, 5.4286] \end{pmatrix}$
Temps de calcul (en s)	0.000043	0.000064	0.000042

Tableau2.4 : Solution du système linéaire à intervalle

On compare le temps de calcul et la précision par les trois méthodes, nous utilisons comme type de machine : Intel(R) Pentium(R) Dual CPU, E 2200, 2.20GHz, 960 Mo de RAM

Nous remarquons d’après les résultats du tableau, par ce simple exemple, que la solution est mieux encadrée par la fonction « verifylss » et le temps de calcul est réduit en utilisant cette méthode.

2.6 Conclusion

Ce chapitre a été consacré à la présentation de l’outil fondamental qui est l’arithmétique des intervalles pour la suite de notre travail. En premier lieu, nous nous sommes intéressés à l’arithmétique des intervalles ainsi qu’aux principaux problèmes rencontrés lors de la manipulation des intervalles. Nous avons souligné le fait que nous manipulons des vecteurs intervalles (pavés) engendre un pessimisme dû au phénomène d’enveloppement ainsi qu’au phénomène de dépendance. Nous avons remarqué qu’en utilisant des fonctions d’inclusion plus élaborées telles que les fonctions d’inclusion moyenne et de Taylor contribuent à réduire ce pessimisme. Pour le problème d’enveloppement, il est souhaitable d’avoir recours au partitionnement des pavés afin de réduire ce pessimisme. Puis, nous avons présenté des outils mettant en œuvre l’arithmétique des intervalles pour manipuler et caractériser de façon approchée et garantie des ensembles de valeurs. Nous avons présenté le principe de l’inversion ensembliste par arithmétique des intervalles (SIVIA), souvent utilisé pour des problèmes non linéaires.

L’inversion ensembliste devient complexe dès que la taille des pavés à partitionner devient grande. Alors, la caractérisation d’ensembles par contraction vient en aide en utilisant un principe de réduction ou de contraction basé sur les notions de consistance. L’idée étant d’éliminer, pour un domaine initial, des parties inconsistantes par rapport à une contrainte. Cette approche permet alors, de réduire le temps de calcul de l’algorithme SIVIA. Nous avons présenté des exemples d’utilisation des contracteurs sur la stabilité et l’estimation à erreur

bornée. Enfin, nous avons illustré des méthodes utilisant la technique des intervalles pour la résolution des systèmes linéaires à intervalles.

Dans le chapitre suivant, nous allons aborder et traiter le problème d'estimation d'état et ceci en se basant sur l'outil de l'analyse par intervalles.

Chapitre 3 :

Estimation d'état

par l'arithmétique d'intervalles

Chapitre3 :

Estimation d'état par l'arithmétique d'intervalles

3.1 Introduction

Mettre en application une stratégie efficace de commande exige souvent la connaissance des variables d'état du système dynamique. Cependant, dans la plupart des applications, la totalité des états du système ne peut pas être mesurée. Les états doivent être estimés à partir de ses sorties par l'intermédiaire d'un observateur d'état. Les premiers observateurs d'état qui ont été développés, il y a plusieurs décennies, par Luenberger [Luenberger, 1964] et Kalman [Labarrer et al., 1993], peuvent reconstruire l'état d'un système dynamique en connaissance parfaite des paramètres du modèle. Cependant en pratique, le système peut comprendre des incertitudes (sur les paramètres, les conditions initiales du vecteur d'état, les perturbations et bruit de mesure). Ceci est une limitation sérieuse pour l'application de ces observateurs classiques puisque les solutions convergent vers une estimation d'état biaisée.

Dans la plupart des situations, les incertitudes sur le système sont mal connues mais pas totalement inconnues. En général, les limites supérieures et inférieures sont connues. Donc il serait intéressant de tirer profit de cette connaissance partielle pour obtenir une estimation fiable. Afin de remédier à ces incertitudes, deux grandes approches ont été développées :

- La première est l'estimation d'état par l'approche multi-modèle [Magill, 1965], [Athans et al., 1977], [Banerjee et al., 1977] et [Li et Bar-Shalom, 1966].
- Et la deuxième est l'approche par les observateurs intervalles [Gouzé et al., 2000] , [Rapaport et Gouzé, 2003] , [Moisan et al., 2009], [Raissi et al., 2005], [Mazenc et Bernard, 2010] , [Bernard et Gouzé, 2004] , [Moisan, 2007] , [Gouzé et al., 2000] , [Sauvage et al., 2007] et [Mazenc et Bernard, 2011] .

L'approche multi-modèle fournit une estimation d'état par combinaison des estimations obtenues par plusieurs systèmes avec différentes valeurs des paramètres [Magill, 1965]. Cette approche est en particulier adaptée aux procédés sujet aux multiples points de fonctionnement [Athans et al., 1977], [Banerjee et al., 1977] . Cependant, la détermination du nombre de modèles à considérer est un facteur difficile car avec un nombre limité de modèles, les informations fournies au sujet de la dynamique du procédé sont insuffisantes pour obtenir des estimations fiables. D'un autre côté, l'estimation pourrait être inefficace avec un nombre trop élevé de modèles. Cette inefficacité est due principalement à la présence de modèles nommés « modèles inutiles » [Li et Bar-Shalom, 1966].

L'approche par les observateurs intervalles a été introduite pour tenir compte des incertitudes caractérisant certaines classes des systèmes [Gouzé et al., 2000], [Rapaport et Gouzé, 2003]. Partant de la connaissance des bornes supérieures et inférieures des incertitudes des paramètres, des conditions initiales de l'état à estimer ou des signaux perturbateurs (perturbations, bruit), la technique consiste à concevoir deux observateurs :

- Le premier (observateur inférieur) est construit sur la base des bornes inférieures des incertitudes (paramètres, conditions initiales, perturbations, bruit).
- Le second (observateur supérieur) est construit sur la base des bornes supérieures des incertitudes (paramètres, conditions initiales, perturbations, bruit).

L'estimateur de l'état du système est ensuite calculé à partir des inégalités différentielles combinant les deux observateurs inférieur et supérieur.

Cette technique est largement utilisée dans plusieurs applications [Moisan et al., 2009], [Raissi et al., 2005], [Mazenc et Bernard, 2010] et elle a été particulièrement exploitée avec succès dans le domaine des procédés biologiques [Bernard et Gouzé, 2004] , [Moisan, 2007] et [Gouzé et al., 2000].

Les observateurs intervalles sont construits pour les systèmes vérifiant une condition de **coopérativité**. Un système est dit **coopératif** si les éléments hors diagonaux de la matrice d'état (pour les systèmes linéaires) ou la matrice Jacobienne (pour les systèmes non linéaires) sont non négatifs [Gouzé et al., 2000].

Ce troisième chapitre est structuré comme suit :

Nous allons introduire dans la section (3.2) les concepts de base de la théorie classique de l'estimation d'état des systèmes linéaires temps invariant (LTI) tels que la commandabilité, l'observabilité et un rappel sur les observateurs classiques.

Dans la section (3.3), nous allons redéfinir la commandabilité et l'observabilité pour la classe des systèmes LTI en présence d'incertitudes sur les paramètres des matrices des systèmes [Ahn et al., 2005] .

Dans la section (3.4.1), nous allons développer un observateur intervalle en se référant sur le travail de [Sauvage et al., 2007]. Cet observateur est synthétisé pour la classe des systèmes LTI à paramètres incertains. L'avantage qu'offre cet observateur par rapport aux observateurs cités précédemment est que l'observateur final pour le système LTI coopératif est obtenu par superposition pondérée des deux observateurs inférieur et supérieur de l'observateur intervalle. Un exemple d'application sera entrepris sur un modèle d'état d'un réacteur chimique continu.

Cependant, cette condition de coopérativité est une contrainte à la synthèse d'observateurs intervalles, car la plupart des systèmes physiques ne le sont pas.

Pour un système linéaire stable en boucle ouvert, on peut toujours le rendre coopératif par un changement de base temps variant et développer un observateur intervalle temps variant et convergent en boucle ouverte [Mazenc et Bernard, 2011]. Dans la section (3.4.2), nous allons illustrer les étapes de constructions de cet observateur, puis nous l'appliquerons à un exemple d'un système stable en boucle ouverte (d'ordre trois).

3.2 Revue sur les concepts de base de la théorie classique de l'estimation des systèmes linéaires temps invariant

Dans cette partie, nous allons revoir les concepts de base de la théorie classique de l'estimation tels que la commandabilité, l'observabilité des systèmes LTI. Un rappel sur les observateurs classiques est présenté.

3.2.1 Commandabilité

Soit le système linéaire temps invariant suivant :

$$\begin{cases} \dot{x} = Ax + Bu \\ y = Cx + Du \end{cases} \quad (3.1)$$

Où $x \in R^n$ est le vecteur d'état, $u \in R^r$ est l'entrée de commande du système, $y \in R^p$ est le vecteur des variables mesurées. Les matrices $A \in R^{n \times n}$, $B \in R^{n \times r}$ et $C \in R^{p \times n}$ sont des matrices constantes. L'état initial est à une valeur quelconque $x_0 = x(t_0)$ et nous supposons que la matrice directe est nulle ($D = 0$).

Le système LTI est dit commandable si pour toute instance x_1 du vecteur d'état, il existe un signal d'entrée $u(t)$ d'énergie finie qui permet au système de passer de l'état x_0 à l'état x_1 en un temps fini.

Il est possible que la commandabilité n'est pas vérifiée sur tout le vecteur d'état, mais elle puisse néanmoins l'être sur une partie de ses composantes, celles-ci sont dites variables d'état commandables du système.

La commandabilité peut être vue comme la possibilité de modifier les dynamiques d'un modèle en agissant sur ses entrées. Cette propriété ne se réfère qu'à l'état et à l'entrée du système. Il est donc clair qu'elle ne dépend que des matrices A et B . Selon le critère de Kalman, la paire (A, B) ou le système (3.1) est commandable si et seulement si :

$$\text{rang}(\mathcal{C}) = n \text{ où } \mathcal{C} = [B \ AB \ \dots \ A^{n-1}B] \quad (3.2)$$

La matrice $\mathcal{C}$ est dite matrice de commandabilité.

3.2.2 Observabilité

Le modèle d'état (3.1) est observable si quel que soit t_0 , il existe un intervalle de temps fini $[t_0, t_1]$ tel que la connaissance de l'entrée $u(t)$ et de la sortie $y(t)$ sur cet intervalle permet de reconstituer l'état $x(t_0)$.

Cette propriété peut aussi ne pas être vérifiée sur toutes les variables du vecteur d'état, mais sur une partie seulement, ces variables sont les états observables du système.

Selon le critère de Kalman, l'observabilité ne dépend en fait que des matrices A et C . La paire (A, C) ou le système (3.1) est observable si et seulement si :

$$\text{rang}(\mathcal{O}) = n \text{ où } \mathcal{O} = \begin{bmatrix} C \\ CA \\ \vdots \\ CA^{n-1} \end{bmatrix} \quad (3.3)$$

La matrice $\mathcal{O}$ est dite matrice de d'observabilité.

3.2.3 Définition d'un observateur

Afin de donner le concept général de l'objectif d'un observateur dans un système donné, nous considérons le cas général d'un système dynamique non linéaire donné par :

$$\begin{cases} \dot{x} = f(x(t), u(t)), x(t_0) = x_0 \\ y = h(x(t), u(t)) \end{cases} \quad (3.4)$$

Où $x \in R^n$ est le vecteur d'état, $f: R^n \times R^r \rightarrow R^n$ et $h: R^n \times R^r \rightarrow R^p$ sont des fonctions non linéaires, $y \in R^p$ est le vecteur des variables mesurées et $u \in R^r$ est l'entrée de commande du système.

Un observateur est un système auxiliaire dont l'objectif est de reconstruire les variables indisponibles à la mesure du vecteur d'état $x(t)$ du système dynamique, en utilisant les informations disponibles tels que le modèle mathématique (3.4) et les mesures y disponibles en ligne.

La définition la plus classique est la suivante :

Définition 3.1

Un observateur asymptotique du système (3.4) est un système dynamique de la forme suivante :

$$\begin{cases} \dot{z} = g(z(t), u(t), y(t)); z(t_0) = z_0 \\ \hat{x} = h(z(t), u(t), y(t)) \end{cases} \quad (3.5)$$

tel que

$$\lim_{t \rightarrow \infty} \|x(t) - \hat{x}(t)\| = 0 \quad (3.6)$$

Le principe de l'observateur est illustré par la figure 3.1. L'observateur converge asymptotiquement vers l'état du système.

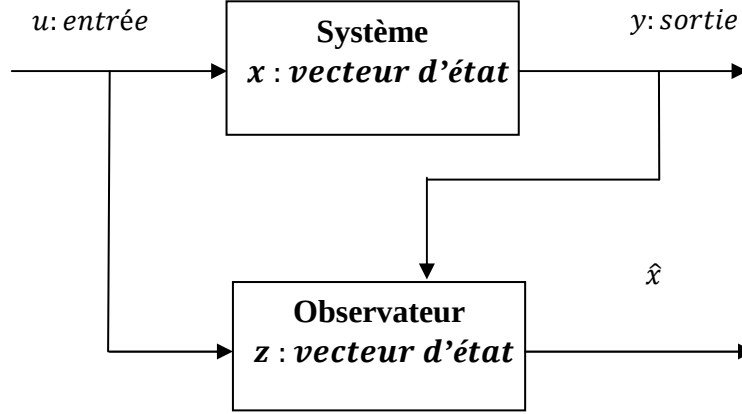


Figure 3.1 : Principe d'estimation d'état.

3.2.4 Méthodes classiques d'estimation

Nous allons rappeler quelques principaux observateurs classiques pour la classe des systèmes linéaires continus tels que l'observateur de Luenberger et le filtre de Kalman.

a) Observateur de Luenberger

Considérons le système (3.1) dans le cas d'un système linéaire que nous réécrivons :

$$\begin{cases} \dot{x} = Ax + Bu \\ y = Cx + Du \end{cases}$$

D'après [Luenberger, 1966], si le système est observable, alors il existe une matrice gain L telle que $(A - LC)$ est une matrice stable, le système suivant :

$$\dot{\hat{x}} = A \hat{x}(t) + Bu(t) + L(y - C\hat{x}(t)), \quad \hat{x}(t_0) = \hat{x}_0 \neq x_0 \quad (3.7)$$

est alors un observateur d'état du système (3.1). L'erreur dynamique est donnée par le système suivant :

$$\dot{e}(t) = (A - LC)e(t) \quad (3.8)$$

avec

$$e = x - \hat{x} \quad , \quad e(t_0) = x_0 - \hat{x}_0 \quad (3.9)$$

La convergence de l'état estimé vers l'état réel est garantie par le choix du gain L calculé au moyen de placement de pôles associés à la matrice $A - LC$ (vérifier que les valeurs propres de $(A - LC)$ sont à parties réelles strictement négatives).

b) Filtre de Kalman continu

Le filtre de Kalman continu [Labarrer et al., 1993] peut être vu comme une extension de l'observateur de Luenberger.

Considérons le système dynamique suivant :

$$\begin{cases} \dot{x}(t) = A(t)x(t) + u(t) + v(t), & x_0 = x(t_0) \\ y(t) = C(t) x(t) + w(t) \end{cases} \quad (3.10)$$

avec :

$$\begin{aligned} E[v(t)] &= 0 & E[v(t) v^T(\tau)] &= Q(t) \delta(t - \tau) \\ E[w(t)] &= 0 & E[w(t) w^T(\tau)] &= R(t) \delta(t - \tau) \\ E[x(t_0)] &= \hat{x}_0 & Cov[x(t_0)] &= P_0 \\ E[v(t) w^T(\tau)] &= E[v(t) x^T(t_0)] = E[w(t) x^T(t_0)] = 0 \end{aligned}$$

$w(t)$ et $v(t)$ sont deux bruit blancs centrés indépendant (perturbations gaussiennes) avec respectivement de covariance $Q(t)$ et $R(t)$. Le filtre de Kalman continu est un filtre linéaire non biaisé, à variance minimale.

Le filtre linéaire est sous la forme :

$$\dot{\hat{x}}(t) = F(t)\hat{x}(t) + B(t) u(t) + L(t)y(t) \quad \hat{x}(t_0) = \hat{x}_0 \quad (3.11)$$

On a comme hypothèse, l'estimé $\hat{x}$ est non biaisé, c'est-à-dire :

$$E[\hat{x}(t)] = E[x(t)] \quad E[\dot{\hat{x}}(t)] = E[\dot{x}(t)] \quad (3.12)$$

Prenons la moyenne des expression (3.10) , (3.11) et verifiant l'hypothèse (3.12) :

$$F(t) E[x(t)] + B(t) u(t) + L(t) C(t) E[x(t)] = A(t) E[x(t)] + u(t)$$

D'où les relations suivantes :

$$\begin{aligned} F(t) &= A(t) - L(t) C(t) \\ B(t) &= I \end{aligned}$$

et l'expression du filtre :

$$\dot{\hat{x}}(t) = A(t)\hat{x}(t) + u(t) + L(t)[y(t) - C(t)\hat{x}(t)]$$

Dont laquelle la matrice gain $L(t)$ doit minimiser la variance P de l'erreur $e = x - \hat{x}$. Cette erreur e vérifie l'équation différentielle dont l'expression est la suivante :

$$\dot{e}(t) = (A(t) - L(t) C(t)) e(t) + v(t) - L(t)w(t) \quad (3.12)$$

$e(t)$ est donc la sortie d'un système linéaire ayant pour entrée un bruit blanc $v(t) - L(t)w(t)$ qui a comme variance $Q + L R L^T$.

La variance P de l'erreur $e = x - \hat{x}$ est solution de l'équation différentielle (équation de Ricatti) suivante:

$$\dot{P}(t) = A(t)P(t) + P(t)A^T(t) - P(t)C^T(t)R^{-1}(t)C(t)P(t) + Q(t) \quad (3.13)$$

Pour minimiser la variance de l'erreur P , on minimise $\dot{P}(t)$. C'est à dire choisir la matrice gain comme suit :

$$L(t) = P(t) C^T R^{-1}(t)$$

D'où les équations du filtre :

$$\begin{aligned} \dot{\hat{x}} &= A(t)\hat{x}(t) + u(t) + L(t)(y(t) - C(t)\hat{x}(t)) \\ \dot{P}(t) &= A(t)P(t) + P(t)A^T(t) - P(t)C^T(t)R^{-1}(t)C(t)P(t) + Q(t) \\ L(t) &= P(t)C^T R^{-1}(t) \end{aligned}$$

avec, comme conditions initiales :

$$\begin{cases} \hat{x}(t_0) = \hat{x}_0 \\ P(t_0) = P_0 \end{cases}$$

- Régime permanent du filtre de Kalman

Dans le cas stationnaire :

$$\begin{cases} \dot{x}(t) = A x(t) + v(t) \\ y(t) = C x(t) + w(t) \end{cases}$$

avec les matrices A, C et les matrices Q, R constantes, si le système est complètement observable et commandable, le filtre de Kalman possède un régime permanent unique.

Les matrices de gain et de variance sont alors constantes et données par les équations suivantes :

$$\begin{aligned} L &= P C^T R^{-1} \\ \dot{P}(t) &= A P + P A^T - P C^T R^{-1} C P + Q = 0 \end{aligned} \quad (3.14)$$

La solution positive définie peut être obtenue en résolvant l'équation différentielle (3.14) et en testant le caractère positif défini des matrices solutions P .

3.3 Test de commandabilité et d'observabilité robuste des systèmes linéaires temps invariant à paramètres incertains par l'arithmétique d'intervalles

Des solutions ont été apportées au problème d'incertitudes des systèmes LTI dans le domaine de la commandabilité et de l'observabilité par l'utilisation de l'arithmétique d'intervalles. Cette nouvelle technique [Ahn et al., 2005] , fait appel à un nouveau concept qui est l'application de l'indépendance (dépendance) linéaire des vecteurs intervalles (pavés) pour le test de commandabilité (et de l'observabilité) robuste des systèmes LTI à paramètres incertains.

Sachant qu'un système LTI à paramètres incertains peut être représenté par des matrices intervalles, ces dernières aussi peuvent être considérées comme des ensembles de vecteurs intervalles colonnes. Nous allons voir que suggérer comme condition une indépendance linéaire de ces pavés contribue à la vérification du test de commandabilité (le test d'observabilité est un problème dual).

Nous allons en premier lieu, donner quelques définitions nécessaires afin d'illustrer cette technique. Puis, montrer l'utilisation de la propriété d'indépendance (dépendance) linéaire des vecteurs intervalles dans son application pour le test de commandabilité.

Deux méthodes sont illustrées, la première consiste à tester la commandabilité des systèmes LTI incertains, et la deuxième consiste à tester la non-commandabilité des systèmes LTI incertains.

Même stratégie adoptée pour l'observabilité.

Les méthodes utilisées sont développées à l'aide de la toolbox INTLAB (Annexe A) et testées sur des exemples de systèmes LTI à paramètres incertains.

Définition 3.2

Une matrice intervalle $X^I \in IR^{m \times n}$, comme elle a été déjà définie dans le chapitre 2, peut être considérée comme étant l'ensemble de n vecteurs intervalles de IR^m .

$$X^I = (x_1^I, x_2^I, \dots, x_n^I) \text{ avec } x_i^I \in IR^m \quad (3.15)$$

Cette matrice peut aussi être définie comme suit :

$$X^I = [\underline{X}^I \quad \overline{X}^I] \quad (3.16)$$

avec

$$\underline{X}^I = [\underline{x}_1^I, \underline{x}_2^I, \dots, \underline{x}_n^I] \quad (3.17)$$

et

$$\overline{X}^I = [\overline{x}_1^I, \overline{x}_2^I, \dots, \overline{x}_n^I] \quad (3.18)$$

Comme nous pouvons aussi définir X^I en utilisant sa matrice centrée et sa matrice rayon :

$$X^I = [X_0 - \Delta X, X_0 + \Delta X] \quad (3.19)$$

où X_0 représente la matrice centrée de X^I calculée comme suit :

$$X_0 = mid(X) = \frac{\underline{X}^I + \overline{X}^I}{2} \quad (3.20)$$

et ΔX est la matrice rayon de X^I calculée comme suit :

$$\Delta X = rad(X) = \frac{\overline{X}^I - \underline{X}^I}{2} \quad (3.21)$$

Définition 3.3

Soient deux vecteurs intervalles $x^I, y^I \in IR^m$, nous définissons le rapport entre ces deux vecteurs intervalles par μ_{xy} calculé comme suit :

$$\mu_{xy} = x^I \setminus y^I = \{x_1^I \oslash y_1^I, \dots, x_m^I \oslash y_m^I\} \quad (3.22)$$

3.3.1 Dépendance et indépendance linéaires des vecteurs intervalles

Définition 3.4

Soient n vecteurs intervalles $x_1^I, x_2^I, \dots, x_n^I$, ces vecteurs intervalles sont dit linéairement indépendants s'il existe la seule solution triviale ($a_1 = a_2 = \dots a_n = 0$) tel que :

$$a_1 x_1^I \oplus a_2 x_2^I \oplus \dots \oplus a_n x_n^I = 0^I \quad (3.23)$$

- **Cas de 2 vecteurs intervalles**

Soient deux vecteurs intervalles x_1^I, x_2^I , en se basant sur la définition 3.4, ces vecteurs intervalles sont linéairement indépendants s'il existe la seule solution triviale ($a_1 = a_2 = 0$) tel que :

$$a_1 x_1^I \oplus a_2 x_2^I = 0^I \quad (3.24)$$

La solution de l'équation (3.24) peut être facilement vérifiée à l'aide du rapport μ_{xy} définie précédemment et ceci en utilisant le théorème suivant.

Théorème 3.1 [Ahn et al., 2005]

Soient deux vecteurs intervalles $x^I, y^I \in IR^m$ tel que :

$$x^I = (x_1^I, x_2^I, \dots, x_m^I)^T \quad \text{et} \quad y^I = (y_1^I, y_2^I, \dots, y_m^I)^T \quad (3.25)$$

avec

$$0 \notin x_1^I \cap x_2^I \cap \dots \cap x_m^I \quad \text{et} \quad 0 \notin y_1^I \cap y_2^I \cap \dots \cap y_m^I \quad (3.26)$$

Les deux vecteurs intervalles sont dits linéairement indépendants si la relation suivante est vérifiée :

$$(\mu_{xy})_1 \cap (\mu_{xy})_2 \cap \dots \cap (\mu_{xy})_m = \emptyset \quad \text{où} \quad (\mu_{xy})_i = x_i^I \oslash y_i^I \quad (3.27)$$

La preuve du théorème est facilement vérifiée en utilisant la propriété de commutativité et d'associativité des intervalles scalaires.

- **Cas général**

Au delà de trois vecteurs intervalles, il est difficile d'étendre le théorème 3.1, celui-ci est valable que pour deux vecteurs intervalles. Dans le cas général ceci revient à vérifier la relation suivante :

$$a_1 x_1^I \oplus a_2 x_2^I \oplus \dots \oplus a_n x_n^I = 0^I \quad \text{ou} \quad x_i^I \in IR^m \quad (3.28)$$

La méthode utilisée dans ce cas est appliquée à la matrice intervalle X^I définie par (3.16).

La dimension de X^I est donnée par :

$$\dim(X^I) = (m \times n) \quad (3.29)$$

Trois cas sont considérés pour la dimension de la matrice X^I (1^{er} cas: $m > n$, 2^{ème} cas : $m = n$ et le 3^{ème} cas : $m < n$). Nous allons donner la méthode que pour le 1^{er} cas, pour les autres cas, nous pouvons utiliser les résultats du 1^{er} cas pour les retrouver. Cette technique consiste à construire toutes les sous-matrices carrées S^i de dimension $(n \times n)$ de la matrice X^I .

Le nombre des sous-matrices carrées S^i construites à partir de la matrice X^I est calculé par la relation suivante :

$$k = \frac{m(m-1)(m-2).....(m-n+1)}{n!} \quad (3.30)$$

Un ensemble comprenant toutes les sous-matrices carrées S^i est donné par l'ensemble suivant :

$$S_X = \{S^i, i = \overline{1, k}\} \quad (3.31)$$

et nous définissons un ensemble qui contient les indices des sous-matrices S^i par :

$$s^I = \{s^i, i = \overline{1, k}\} \quad (3.32)$$

Dans MATLAB, l'ensemble des indices s^I est donné par la commande suivante :

```
s=[nchoosek(1:k,n)]
si=s(i,:)
```

et les sous-matrices carrées sont obtenues par la commande suivante :

```
SX=combvec(X(:, si))
```

Nous définissons l'ensemble des matrices centrées des sous-matrices carrées S^i par l'ensemble suivant:

$$S_{Xc} = \left\{ S_0^i = \frac{\overline{s^i} + s^i}{2}, i = \overline{1, k} \right\} \quad (3.33)$$

L'ensemble des matrices rayon des sous-matrices de X^I est donné par la relation :

$$\Delta S_X = \left\{ \Delta S^i = \frac{\overline{s^i} - s^i}{2}, i = \overline{1, k} \right\} \quad (3.34)$$

Définition 3.5

Le rang de la matrice X^I peut être retrouvé en calculant le maximum des rangs des ces sous-matrices carrées S^i qui est donné par la relation suivante :

$$rang(X^I) = \max \{rang(S^i), i = \overline{1, k}\}$$

Le Lemme et les deux théorèmes suivants ont été proposé [Ahn et al., 2005] .

Lemme

Considérons une matrice intervalle carrée X^I et soit X_0 sa matrice centrée qui est non singulière (invertible). Supposons que le rayon spectral $\rho(|(X_0)^{-1}| \Delta X) < 1$ alors la matrice X^I est non singulière.

Théorème 3.2

S'il existe au moins une sous-matrice $S^i \in S_X$ telle que sa matrice centrée correspondante $S_0^i \in S_{Xc}$ est non singulière (invertible) et le rayon spectral

$$\rho(|(S_0^i)^{-1}| \Delta S^i) < 1 \quad (3.35)$$

alors les vecteurs intervalles $x_1^I, x_2^I, \dots, x_n^I$ de la matrice intervalle X^I sont **linéairement indépendants**.

Afin de tester maintenant la **dépendance** linéaire des vecteurs intervalles, le théorème suivant a été proposé.

Théorème 3.3

Pour toute sous-matrice carrée $S^i \in S_X$ et pour toute matrice centrée $S_0 \in S_{Xc}$ et pour toute matrice rayon $\Delta S \in \Delta S_X$ correspondante, s'il existe une matrice R et un nombre naturel p tel que :

$$(I + |I - S_0 R|)_p \leq (\Delta S |R|)_p \quad (3.36)$$

où $p \in \{1 \dots n\}$ et $(\cdot)_p$ représente la $p^{\text{ème}}$ colonne, alors les vecteurs intervalles $x_1^I, x_2^I, \dots, x_n^I$ de la matrice intervalle X^I sont **linéairement dépendants**.

3.3.2 Application au test de commandabilité robuste et d'observabilité robuste

Soit le système LTI incertain à intervalles suivant :

$$\dot{x} = Ax + Bu \quad (3.37)$$

$x \in IR^n$ est le vecteur d'état, $u \in IR^r$ est l'entrée du système, $A \in IR^{n \times n}$ et $B \in IR^{n \times r}$ sont des matrices intervalles.

Où

$$\text{rang}(B) = r \quad (3.38)$$

Les matrices A et B s'écrivent :

$$A \in A^I = [\underline{A}, \bar{A}] \quad \text{et} \quad B \in B^I = [\underline{B}, \bar{B}] \quad (3.39)$$

Le système (3.37) est dit commandable si le $\text{rang}(\mathcal{C}) = n$ pour tout $\mathcal{C} \in \mathcal{C}^I$, la matrice $\mathcal{C}^I$ est donnée par :

$$\mathcal{C}^I = [B^I, A^I \otimes B^I, A^I \otimes A^I \otimes B^I, \dots, \underbrace{A^I \otimes A^I \otimes \dots \otimes A^I}_{n-r} \otimes B^I] \quad (3.40)$$

avec la dimension de la matrice intervalle $\mathcal{C}^I$ est égale $(n \times m) = n \times \underbrace{(n - r + 1) \cdot r}_m$

- **Le cas sans intervalles**

Nous allons considérer en premier lieu le cas sans intervalles, le système est donné par :

$$\dot{x} = A_0 x + B_0 u \quad (3.41)$$

La matrice de commandabilité correspondante est la suivante :

$$\mathcal{C}_0 = [B_0, A_0 B_0, (A_0)^2 B_0, \dots, (A_0)^{n-r} B_0] \quad (3.42)$$

Si le système est commandable alors on a toujours :

$$\text{rang}(\mathcal{C}_0) = n \quad (3.43)$$

Afin de distinguer le cas sans intervalles du cas avec intervalles, considérons la sous-matrice $(\mathcal{C}_0)'$ construite à partir de $\mathcal{C}_0$:

$$(\mathcal{C}_0)' = [B_0, A_0 B_0, (A_0)^2 B_0, \dots, (A_0)^{n-r-q} B_0] \quad \text{où} \quad q \geq 1 \quad (3.44)$$

et supposons que :

$$\text{rang}(\mathcal{C}_0)' = n$$

alors, dans le cas sans intervalles on a toujours la relation suivante qui est vérifiée :

$$\text{rang}(\mathcal{C}_0)' = \text{rang}(\mathcal{C}_0) = n \quad (3.45)$$

- **Cas avec intervalles**

Dans ce cas, il faudra vérifier le rang de la matrice intervalle $(\mathcal{C}^I)$. Etant donné que la matrice $(\mathcal{C}^I)$ est de dimension $(n \times m)$, son rang est difficile à retrouver.

Alors, considérons la sous-matrice $(\mathcal{C}^I)'$ construite à partir de la matrice intervalle $\mathcal{C}^I$ définie par

$$(\mathcal{C}^I)' = [B^I, A^I \otimes B^I, A^I \otimes A^I \otimes B^I, \dots, \underbrace{A^I \otimes A^I \otimes \dots \otimes A^I}_{n-r-q} \otimes B^I] \quad (3.46)$$

Le Corollaire suivant a été proposé [Ahn et al., 2005] :

Corollaire

Si la sous-matrice $(\mathcal{C}^I)'$ de la matrice intervalle de commandabilité $\mathcal{C}^I$ vérifie la propriété d'indépendance linéaire des vecteurs intervalles définie par le théorème 3.2 alors le système LTI à paramètres incertains est commandable.

Preuve

Un système LTI est commandable si sa matrice de commandabilité est de rang plein. C'est équivalent à avoir la condition d'indépendance linéaire de ces vecteurs intervalles (preuve immédiate).

Le test d'observabilité est un problème dual.

Nous allons illustrer cette méthode sur deux exemples :

Exemple 3.1

Soit le système LTI incertain à intervalles suivant :

$$\dot{x} = Ax + Bu$$

Les matrices A et B sont données ci-dessous :

$$A \in A^I = \begin{pmatrix} [\begin{smallmatrix} 0.9499, & 1.0501 \end{smallmatrix} & [\begin{smallmatrix} 0.0000, & 0.0000 \end{smallmatrix} & [\begin{smallmatrix} 0.0000, & 0.0000 \end{smallmatrix} \\ [\begin{smallmatrix} 0.0000, & 0.0000 \end{smallmatrix} & [\begin{smallmatrix} 0.9599, & 1.0401 \end{smallmatrix} & [\begin{smallmatrix} 0.9699, & 1.0301 \end{smallmatrix} \\ [\begin{smallmatrix} 0.0000, & 0.0000 \end{smallmatrix} & [\begin{smallmatrix} -2.0801, & -1.9199 \end{smallmatrix} & [\begin{smallmatrix} 3.6000, & 4.4001 \end{smallmatrix} \end{pmatrix}$$

avec la dimension de la matrice A est donnée par :

$$\dim(A) = (n \times n) = (3 \times 3)$$

et B est une matrice ponctuelle donnée par :

$$B \in B^I = \begin{pmatrix} 1 & 0 \\ 0 & 0 \\ 0 & 1 \end{pmatrix}$$

la dimension de B est :

$$\dim(B) = (n \times r) = (3 \times 2)$$

La matrice de commandabilité calculée par l'arithmétique d'intervalles est la suivante :

$$C \in C^I = \begin{pmatrix} [1.0000, 1.0000] & [0.0000, 0.0000] & [0.9499, 1.0501] & [0.0000, 0.0000] \\ [0.0000, 0.0000] & [0.0000, 0.0000] & [0.0000, 0.0000] & [0.9699, 1.0301] \\ [0.0000, 0.0000] & [1.0000, 1.0000] & [0.0000, 0.0000] & [3.6000, 4.4001] \end{pmatrix}$$

avec

$$\dim(C) = (n \times (n - r + 1) \cdot r) = (n \times m) = (3 \times 4)$$

On a $n < m$ alors le nombre des sous-matrices carrées qui peuvent être construites à partir de la matrice intervalle C^I calculé à partir de (3.30) est $k = 4$.

- L'ensemble indice des sous-matrices carrées construites à partir de C^I est :

$$s^I = \{s^i, i = \overline{1, k}\} = \{s^1, s^2, s^3, s^4\}$$

- L'ensemble des sous-matrices carrées construites à partir de C^I est :

$$S_C = \{S^i, i = \overline{1, k}\} = \{S^1, S^2, S^3, S^4\}$$

Pour l'indice $s^1 = \{1 \quad 2 \quad 3\}$ on a la sous-matrice S^1 suivante :

$$S^1 = \begin{pmatrix} [1.0000, 1.0000] & [0.0000, 0.0000] & [0.9499, 1.0501] \\ [0.0000, 0.0000] & [0.0000, 0.0000] & [0.0000, 0.0000] \\ [0.0000, 0.0000] & [1.0000, 1.0000] & [0.0000, 0.0000] \end{pmatrix}$$

La matrice centrée de S^1 est donnée par :

$$S_0^1 = \begin{pmatrix} 1 & 0 & 1 \\ 0 & 0 & 0 \\ 0 & 1 & 0 \end{pmatrix}$$

et sa matrice rayon est:

$$\Delta S^1 = \begin{pmatrix} 0 & 0 & 0.05 \\ 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

Pour l'indice $s^2 = \{1 \quad 2 \quad 4\}$ on a la sous-matrice S^2 suivante :

$$S^2 = \begin{pmatrix} [\quad 1.0000, \quad 1.0000] & [\quad 0.0000, \quad 0.0000] & [\quad 0.0000, \quad 0.0000] \\ [\quad 0.0000, \quad 0.0000] & [\quad 0.0000, \quad 0.0000] & [\quad 0.9699, \quad 1.0301] \\ [\quad 0.0000, \quad 0.0000] & [\quad 1.0000, \quad 1.0000] & [\quad 3.6000, \quad 4.4001] \end{pmatrix}$$

La matrice centrée de S^2 est :

$$S_0^2 = \begin{pmatrix} 1 & 0 & 1 \\ 0 & 0 & 1 \\ 0 & 1 & 4 \end{pmatrix}$$

La matrice rayon de S^2 est :

$$\Delta S^2 = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0.0300 \\ 0 & 0 & 0.4000 \end{pmatrix}$$

Pour l'indice $s^3 = \{1 \quad 3 \quad 4\}$ on a la sous- matrice suivante :

$$S^3 = \begin{pmatrix} [\quad 1.0000, \quad 1.0000] & [\quad 0.9499, \quad 1.0501] & [\quad 0.0000, \quad 0.0000] \\ [\quad 0.0000, \quad 0.0000] & [\quad 0.0000, \quad 0.0000] & [\quad 0.9699, \quad 1.0301] \\ [\quad 0.0000, \quad 0.0000] & [\quad 0.0000, \quad 0.0000] & [\quad 3.6000, \quad 4.4001] \end{pmatrix}$$

La matrice centrée de S^3 est la suivante :

$$S_0^3 = \begin{pmatrix} 1 & 1 & 0 \\ 0 & 0 & 1 \\ 0 & 0 & 4 \end{pmatrix}$$

La matrice rayon de S^3 est :

$$\Delta S^3 = \begin{pmatrix} 0 & 0.0500 & 0 \\ 0 & 0 & 0.0300 \\ 0 & 0 & 0.4000 \end{pmatrix}$$

Pour l'indice $s^4 = \{2 \ 3 \ 4\}$ on a la sous matrice S^4 suivante :

$$S^4 = \begin{pmatrix} [0.0000, 0.0000] & [0.9499, 1.0501] & [0.0000, 0.0000] \\ [0.0000, 0.0000] & [0.0000, 0.0000] & [0.9699, 1.0301] \\ [1.0000, 1.0000] & [0.0000, 0.0000] & [3.6000, 4.4001] \end{pmatrix}$$

La matrice centrée de S^4 est :

$$S_0^4 = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 0 & 4 \end{pmatrix}$$

La matrice rayon de S^4 est :

$$\Delta S^4 = \begin{pmatrix} 0 & 0.0500 & 0 \\ 0 & 0 & 0.0300 \\ 0 & 0 & 0.4000 \end{pmatrix}$$

Test de commandabilité robuste

- Les matrices centrées S_0^1, S_0^3 des sous-matrices carrées S^1, S^3 respectivement sont singulières (non inversibles).
- Les matrices centrées S_0^2, S_0^4 des sous-matrices carrées S^2, S^4 respectivement sont non singulières.
- Le rayon spectrale $\rho(|(S_0^2)^{-1}| \Delta S^2) = 0.0300 < 1$ et $\rho(|(S_0^4)^{-1}| \Delta S^4) = 0.0500 < 1$.

Afin de vérifier le test de commandabilité robuste du système LTI incertain à intervalles, il suffit d'avoir la non singularité d'au moins une matrice centrée des sous-matrices intervalles de la matrice de commandabilité $\mathcal{C}^I$ (dans notre cas S_0^2 et S_0^4 sont non singulières). Et d'avoir un rayon spectral inférieur à 1 (dans notre cas $\rho(|(S_0^2)^{-1}| \Delta S^2) < 1$, $\rho(|(S_0^4)^{-1}| \Delta S^4) < 1$). Donc le système LTI incertain est commandable.

Exemple 3.2

Dans cet exemple l'incertitude sur les paramètres des matrices du système LTI à intervalles est représentée par la variable α . En faisant varier cette variable nous testons aux mêmes temps la commandabilité et la non-commandabilité du système LTI incertain.

Nous vérifions le test de la non-commandabilité du système LTI incertain par le théorème 3.3, en remplaçant dans la relation (3.36) la matrice R par la matrice $(S_0)^{-1}$ ce qui revient à vérifier l'inégalité suivante :

$$(I)_P \leq (\Delta S |(S_0)^{-1}|)_P \quad (3.47)$$

Les matrices A et B du système LTI incertain sont données par les matrices intervalles suivantes :

$$A \in A^I = \begin{pmatrix} [1 - \alpha, 1 + \alpha] & [2 - 2\alpha, 2 + 2\alpha] & [-1 - \alpha, -1 + \alpha] \\ [-2 - 2\alpha, -2 + 2\alpha] & [1 - \alpha, 1 + \alpha] & [1 - \alpha, 1 + \alpha] \\ [0.5 - 0.5\alpha, 0.5 + 0.5\alpha] & [-2 - 2\alpha, -2 + 2\alpha] & [4 - 4\alpha, 4 + 4\alpha] \end{pmatrix}$$

$$B \in B^I = \begin{pmatrix} [1 - \alpha, 1 + \alpha] & [-0.1 - 0.1\alpha, -0.1 + 0.1\alpha] \\ [0.1 - 0.1\alpha, 0.1 + 0.1\alpha] & [0.1 - 0.1\alpha, 0.1 + 0.1\alpha] \\ [-0.1 - 0.1\alpha, -0.1 + 0.1\alpha] & [1 - \alpha, 1 + \alpha] \end{pmatrix}$$

La matrice B dans cet exemple est prise incertaine (matrice intervalle pleine).

Nous donnons ci-dessous les résultats pour le test de commandabilité et de la non-commandabilité pour la valeur $\alpha=0.1$. Par la suite, on varie α , et nous dressons un tableau résumant les résultats des deux tests :

Pour la valeur $\alpha=0.1$ nous trouvons les matrices A et B suivantes:

$$A \in A^I = \begin{pmatrix} [0.9000, 1.1001] & [1.8000, 2.2001] & [-1.1001, -0.9000] \\ [-2.2001, -1.8000] & [0.9000, 1.1001] & [0.9000, 1.1001] \\ [0.4499, 0.5501] & [-2.2001, -1.8000] & [3.6000, 4.4001] \end{pmatrix}$$

$$\text{et } B \in B^I = \begin{pmatrix} [0.9000, 1.1001] & [-0.0900, -0.0899] \\ [0.0899, 0.1101] & [0.0899, 0.1101] \\ [-0.1101, -0.0899] & [0.9000, 1.1001] \end{pmatrix}$$

La matrice de commandabilité calculée par l'arithmétique d'intervalles est la suivante :

$$C \in C^I = \begin{pmatrix} [0.9000, 1.1001] & [-0.0900, -0.0899] & [1.0269, 1.5731] & [-1.1511, -0.6289] \\ [0.0899, 0.1101] & [0.0899, 0.1101] & [-2.4621, -1.5379] & [1.0309, 1.5291] \\ [-0.1101, -0.0899] & [0.9000, 1.1001] & [-0.3311, 0.1311] & [2.8684, 4.6416] \end{pmatrix}$$

Le nombre de sous-matrices qui peuvent être construites à partir de C^I est $k = 4$. Nous trouvons les sous-matrices carrées suivantes :

$$S^1 = \begin{pmatrix} [0.9000, 1.1001] & [-0.0900, -0.0899] & [1.0269, 1.5731] \\ [0.0899, 0.1101] & [0.0899, 0.1101] & [-2.4621, -1.5379] \\ [-0.1101, -0.0899] & [0.9000, 1.1001] & [-0.3311, 0.1311] \end{pmatrix}$$

$$S^2 = \begin{pmatrix} [0.9000, 1.1001] & [-0.0900, -0.0899] & [-1.1511, -0.6289] \\ [0.0899, 0.1101] & [0.0899, 0.1101] & [1.0309, 1.5291] \\ [-0.1101, -0.0899] & [0.9000, 1.1001] & [2.8684, 4.6416] \end{pmatrix}$$

$$S^3 = \begin{pmatrix} [0.9000, 1.1001] & [1.0269, 1.5731] & [-1.1511, -0.6289] \\ [0.0899, 0.1101] & [-2.4621, -1.5379] & [1.0309, 1.5291] \\ [-0.1101, -0.0899] & [-0.3311, 0.1311] & [2.8684, 4.6416] \end{pmatrix}$$

$$S^4 = \begin{pmatrix} [0.9000, 1.1001] & [1.0269, 1.5731] & [-1.1511, -0.6289] \\ [0.0899, 0.1101] & [-2.4621, -1.5379] & [1.0309, 1.5291] \\ [0.9000, 1.1001] & [-0.3311, 0.1311] & [2.8684, 4.6416] \end{pmatrix}$$

◇ Test de commandabilité

- Les matrices centrées des sous-matrices carrées S^1 , S^2 et S^3 sont toutes inversibles.
- Pour S^1 , S^2 et S^3 , le rayon spectral $\rho(|(S_0^i)^{-1}| \Delta S^i)$ calculé est 0.2903, 0.5558 et 0.4325 respectivement qui est inférieur à 1.
- La matrice centrée de la sous-matrice carrée S^4 est non singulière, le rayon spectral $\rho(|(S_0^4)^{-1}| \Delta S^4)$ calculé est 21.1779 > 1

Alors, le test de commandabilité du système LTI incertain à intervalles est vérifié. Car toutes les matrices centrées des sous-matrices carrées de la matrice de commandabilité $\mathcal{C}^I$ sont inversibles. Et nous avons un rayon spectral qui est inférieur à 1 pour $\rho(|(S_0^1)^{-1}| \Delta S^1)$, $\rho(|(S_0^2)^{-1}| \Delta S^2)$ et $\rho(|(S_0^3)^{-1}| \Delta S^3)$.

◇ Test de la non-commandabilité

Pour vérifier le test de la non-commandabilité, nous calculons $(\Delta S^i |(S_0^i)^{-1}|)$ à partir de S^1 , S^2 , S^3 et S^4 nous trouvons :

$$(\Delta S^1 |(S_0^1)^{-1}|) = \begin{pmatrix} 0.1083 & 0.1890 & 0.0164 \\ 0.0344 & 0.2228 & 0.0341 \\ 0.0314 & 0.1158 & 0.1129 \end{pmatrix}$$

$$(\Delta S^2 |(S_0^2)^{-1}|) = \begin{pmatrix} 0.1245 & 0.3279 & 0.0325 \\ 0.0433 & 0.3019 & 0.0429 \\ 0.1639 & 1.3067 & 0.2443 \end{pmatrix}$$

$$(\Delta S^3 |(S_0^3)^{-1}|) = \begin{pmatrix} 0.1185 & 0.1895 & 0.1199 \\ 0.0457 & 0.2228 & 0.1483 \\ 0.0479 & 0.1174 & 0.2810 \end{pmatrix}$$

$$(\Delta S^4 |(S_0^4)^{-1}|) = \begin{pmatrix} 11.0852 & 7.0543 & 0.2922 \\ 14.2805 & 9.0327 & 0.3720 \\ 39.2199 & 25.3326 & 0.8965 \end{pmatrix}$$

Afin de vérifier le test de la non-commandabilité, il faut que toutes les sous-matrices vérifient la condition suivante :

$$(I)_p \leq (\Delta S|(S_0)^{-1})_p$$

Cette condition n'est pas vérifiée pour les sous-matrices S^1 , S^2 et S^3 , donc nous ne pouvons pas conclure sur le test de la non-commandabilité pour la valeur de $\alpha = 0.1$.

Faisant varier la valeur de l'incertitude α du système LTI incertain et à chaque variation nous testons la commandabilité par le théorème 3.2 et la non-commandabilité par le théorème 3.3.

Les résultats des tests de commandabilité et de la non-commandabilité sont résumés dans le tableau 3.1 :

α	0.1	0.15	0.25	0.3	0.35	0.4	0.45	0.5	0.55
Test de commandabilité	vérifié	vérifié	vérifié	vérifié	Non vérifié	Non vérifié	Non vérifiée	Non vérifié	Non vérifié
Test de non-commandabilité	Non vérifié	Non vérifié	Non vérifié	Non vérifié	Non vérifié	vérifié	vérifié	vérifié	vérifié

Tableau 3.1 : Résultats du test de la commandabilité et de la non-commandabilité par variation de α

D'après les résultats du tableau 3.1, le système LTI incertain est commandable jusqu'à la valeur de $\alpha = 0.35$. Pour cette incertitude, ni le test de commandabilité ni le test de la non-commandabilité est vérifié. Au-delà de cette incertitude, le système LTI incertain vérifie le test de non-commandabilité.

Nous donnons ci-dessous les résultats des deux tests pour la valeur d'incertitude $\alpha = 0.35$:

- **Le test de commandabilité pour $\alpha = 0.35$**

Les matrices centrées des sous-matrices carrées S^1 , S^2 , S^3 et S^4 sont toutes inversibles. Le calcul du rayon spectral $\varphi\left(\left|(S_0^i)^{-1}\right| \Delta S^i\right)$ a donné 1.1078, 2.1472, 1.6629 et 41.8833 respectivement qui est supérieur à 1. Donc nous ne pouvons pas conclure sur le test de la commandabilité.

- **Le test de la non-commandabilité pour $\alpha = 0.35$**

Nous avons calculé les matrices $(\Delta S^i|(S_0^i)^{-1})$ qui sont données ci-dessous:

$$(\Delta S^1|(S_0^1)^{-1}) = \begin{pmatrix} 0.3841 & 0.7149 & 0.0537 \\ 0.1302 & 0.8701 & 0.1261 \\ 0.1145 & 0.4504 & 0.3972 \end{pmatrix}$$

$$(\Delta S^2 | (S^2_0)^{-1}) = \begin{pmatrix} 0.4473 & 1.3322 & 0.1199 \\ 0.1658 & 1.1951 & 0.1611 \\ 0.6398 & 5.1898 & 0.9063 \end{pmatrix}$$

$$(\Delta S^3 | (S^3_0)^{-1}) = \begin{pmatrix} 0.4239 & 0.7170 & 0.4574 \\ 0.1730 & 0.8703 & 0.5597 \\ 0.1827 & 0.4572 & 1.0886 \end{pmatrix}$$

$$(\Delta S^4 | (S^4_0)^{-1}) = \begin{pmatrix} 22.4154 & 15.1073 & 0.0537 \\ 28.3430 & 19.3305 & 0.1261 \\ 78.7393 & 51.6176 & 0.3972 \end{pmatrix}$$

et afin de vérifier le test de la non-commandabilité, il faut que toutes les sous-matrices vérifient la condition suivante:

$$(I)_P \leq (\Delta S | (S_0)^{-1})_P$$

or celle-ci ne l'est pas pour la sous-matrice S^1 . Donc, nous ne pouvons pas conclure aussi sur le test de la non-commandabilité.

Dans ce cas, avec la valeur de l'incertitude $\alpha = 0,35$, on ne peut conclure ni sur le test de commandabilité ni sur le test de la non commandabilité du système LTI incertain.

Cette limitation de la méthode est due principalement au choix de la matrice B . Dans cet exemple la matrice B considérée est une matrice intervalle incertaine pleine, c'est ce qui à causé un conservatisme, car avec la valeur de 35 pour cent d'incertitude, on ne peut déduire aucune conclusion sur le test de la commandabilité ou de la non-commandabilité du système LTI incertain.

3.4 Estimation d'état des systèmes linéaires temps invariant incertains par l'arithmétique des intervalles

Dans cette partie, nous allons développer deux observateurs d'état. Le premier observateur, est construit pour la classe des systèmes LTI à paramètres incertains, vérifiant la condition de coopérativité exigée pour la synthèse d'un observateur intervalle. L'estimation du vecteur d'état est obtenue par superposition pondérée des deux observateurs inférieur et supérieur de l'observateur intervalle, le facteur de pondération est calculé à partir de la différence entre la sortie mesurée et les bornes supérieure et inférieure retrouvées par l'observateur intervalle. Cet observateur n'est appliqué que pour les systèmes ayant une sortie mesurée. Comme exemple d'application, on procède à la synthèse d'un observateur intervalle pour un modèle d'état d'un réacteur chimique continu.

Le deuxième observateur est construit pour la classe des systèmes LTI présentant des incertitudes sur les conditions initiales du vecteur d'état et sur les perturbations additives, mais qui peuvent être bornées dans un vecteur intervalle. Le système considéré est stable en boucle ouverte mais ne vérifiant pas la condition de coopérativité. Alors, l'estimation d'état

par cet observateur est basée sur un changement de coordonnées temps-variant qui permet de rendre le système coopératif, l'observateur final est un observateur temps-variant asymptotiquement stable. Un exemple montrant les étapes de transformation du système non coopératif en un système coopératif par ce changement de coordonnées est illustré sur un modèle d'état d'un système LTI stable en boucle ouverte d'ordre trois.

3.4.1 Observateur intervalle pour les systèmes linéaires temps invariant à paramètres incertains

L'estimation d'état est effectuée grâce à un observateur intervalle qui donne une estimation supérieure (observateur supérieur) et une estimation inférieure (observateur inférieur) de l'état réel $x(t)$ du système. L'état estimé représenté par $\hat{x}(t)$ est ensuite calculé en utilisant un facteur poids α qui utilise la sortie mesurée du système et les limites inférieure et supérieure retrouvées par l'observateur intervalle.

Le système linéaire temps invariant considéré est donné par la structure suivante :

$$\begin{cases} \dot{x}(t) = Ax(t) + \theta Bu(t) + Dv(t) \\ y = Cx \end{cases} \quad (3.48)$$

avec l'état initial du système est $x(t_0) = x_0$, le vecteur d'état $x \in R^n$, l'entrée $u \in R^{r1}$, l'entrée $v \in R^{r2}$ et la sortie $y \in R$. Les matrices $A \in R^{n \times n}$, $B \in R^{n \times r1}$, $D \in R^{n \times r2}$ et $C \in R^{n \times 1}$ sont constantes.

Nous supposons que le système vérifie une condition de coopérativité qui est reliée à la matrice d'état A du système (3.48). Ceci est vérifié en imposant aux éléments hors diagonale de la matrice A d'être non négatifs, la relation suivante est donnée :

$$a_{i,j} \geq 0, \quad \forall i \neq j \quad (3.49)$$

avec, $a_{i,j}$ sont les éléments de la matrice A , et le vecteur d'entrée u évolue positivement :

$$\forall t > t_0: 0 < u_{min} \leq u(t) \quad (3.50)$$

Dans le système (3.48), nous supposons que θ est un paramètre constant inconnu, mais qui peut être borné dans un intervalle avec des limites connues. θ satisfait la condition donnée par l'inégalité suivante :

$$\theta^- Bu \leq \theta Bu \leq \theta^+ Bu \quad (3.51)$$

avec θ^- et θ^+ sont les bornes inférieure et supérieure du paramètre incertain θ respectivement.

Supposons aussi qu'il existe un gain L tel que les équations suivantes forment un observateur intervalle pour le système (3.48) :

$$\begin{cases} \dot{z}^-(t) = A z^- + \theta^- Bu + D v + L(y - C z^-) \\ \dot{z}^+(t) = A z^+ + \theta^+ Bu + D v + L(y - C z^+) \end{cases} \quad (3.52)$$

et l'état $x(t)$ est enfermé par les états estimés par l'observateur intervalle, comme le montre la relation suivante :

$$\forall t > t_0: \quad z^-(t) \leq x(t) \leq z^+(t) \quad (3.53)$$

L'état initial $x(t_0)$ doit vérifier l'inégalité suivante :

$$z^-(t_0) \leq x(t_0) \leq z^+(t_0) \quad (3.54)$$

Par la suite, l'estimation finale du vecteur d'état $\hat{x}(t)$ de l'état réel $x(t)$ du système est obtenue au moyen d'un facteur pondéré σ . Ce facteur est calculé en temps réel en utilisant la sortie mesurée y du système (3.48) et les estimations supérieure et inférieure $z^+(t)$ et $z^-(t)$ respectivement trouvées par l'observateur intervalle (3.52). σ est donné par la relation suivante :

$$\sigma = \frac{y - C z^+}{C(z^- - z^+)} \quad (3.55)$$

L'estimation d'état est alors décrite par l'équation suivante :

$$\hat{x}(t) = \sigma z^-(t) - (1 - \sigma) z^+(t) \quad (3.56)$$

NB : Cette méthode est valable que pour les systèmes à une sortie mesurée.

Le théorème suivant a été proposé [Sauvage et al., 2007].

Théorème 3.4

En supposant que la matrice gain L est choisie de telle manière à avoir la matrice

$A_e = A - L C$ stable (Hurwitz) et coopérative, alors l'estimation $\hat{x}(t)$ définie par (3.55) et (3.56) converge asymptotiquement vers l'état réel du système dynamique (3.48).

Nous allons illustrer cette méthode d'estimation à travers l'exemple suivant :

a) Application au modèle d'un réacteur chimique continu

Exemple 3.3

L'exemple considéré est un réacteur continu avec les réactions chimiques suivantes :



Le modèle dynamique du réacteur est décrit par les équations différentielles suivantes :

$$\begin{cases} \dot{C}_A = \frac{Q}{V} (C_{A,in} - C_A) - k_1 C_B \\ \dot{C}_B = -\frac{Q}{V} C_B + k_1 C_A - k_2 C_B \\ \dot{C}_C = -\frac{Q}{V} C_C + k_2 C_B \end{cases} \quad (3.58)$$

où C_A , C_B , C_C et $C_{A,in}$ sont les concentrations de A, B, C et la concentration d'admission de A respectivement en (mol/l) , Q est le débit d'alimentation en (m^3/s) , V est le volume du réacteur en (m^3) et k_1, k_2 sont les constantes cinétiques associées à A et B respectivement en (s^{-1}) .

Nous supposons que seulement la concentration C_C est disponible à la mesure et que $C_{A,in}$ est une variable incertaine qui est bornée entre deux valeurs connues comme suit :

$$C_{A,in,min} \leq C_{A,in} \leq C_{A,in,max} \quad (3.59)$$

Soient $\bar{Q}$ et $\bar{V}$ les conditions de fonctionnement nominales. Le vecteur d'état considéré comprend les variables d'état suivantes C_A , C_B et C_C .

Le modèle du réacteur chimique décrit par les équations différentielles (3.58) est linéarisé au tour du point d'équilibre, nous obtenons la représentation d'état suivante :

$$\begin{cases} \dot{x} = A x + \theta B u \\ y = C x \end{cases} \quad (3.60)$$

avec

$$A = \begin{pmatrix} -\frac{\bar{Q}}{\bar{V}} - k_1 & 0 & 0 \\ k_1 & -\frac{\bar{Q}}{\bar{V}} - k_2 & 0 \\ 0 & k_2 & -\frac{\bar{Q}}{\bar{V}} \end{pmatrix}$$

$$B = \begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}$$

$$C = (0 \ 0 \ 1)$$

où x, u, θ sont respectivement le vecteur d'état, l'entrée du système qui est le taux de dilution $\frac{Q}{V}$ et la concentration d'admission $C_{A,in}$ qui est le paramètre incertain.

D'après la relation (3.59), les bornes supérieure et inférieure $C_{A,in,max}$, $C_{A,in,min}$ respectivement du paramètre incertain $C_{A,in}$ représentent θ^+ et θ^- respectivement du paramètre incertain θ du système dynamique (3.60).

Alors, il faudra maintenant choisir un gain L donné par le vecteur colonne suivant:

$$L = [L_1, L_2, L_3]^T$$

tel que le système ci-dessous soit un observateur intervalle pour le système (3.60) :

$$\begin{cases} \dot{z}^-(t) = A z^- + \theta^- Bu + D v + L(y - C z^-) \\ \dot{z}^+(t) = A z^+ + \theta^+ Bu + D v + L(y - C z^+) \\ \forall t > 0: \quad z^-(t) \leq x(t) \leq z^+(t) \\ z^-(0) \leq x(0) \leq z^+(0) \end{cases} \quad (3.61)$$

Ceci est réalisé en imposant à la matrice $A_e = A - LC$ d'être **stable** et **coopérative** :

$$A_e = A - LC = \begin{pmatrix} -\frac{\bar{Q}}{V} - k_1 & 0 & -L_1 \\ k_1 & -\frac{\bar{Q}}{V} - k_2 & -L_2 \\ 0 & k_2 & -\frac{\bar{Q}}{V} - L_3 \end{pmatrix} \quad (3.62)$$

Afin de vérifier la condition de stabilité et de coopérativité A_e , le gain suivant peut être utilisé comme choix :

$$L = [0 \ 0 \ 0]^T \quad (3.63)$$

car en imposant ce choix, la matrice $A_e = A - LC$ vérifie la condition de coopérativité et de stabilité garantie par la matrice A (A est coopérative et stable).

La matrice A_e calculée est la suivante :

$$A_e = A = \begin{pmatrix} -12 & 0 & 0 \\ 10 & -7 & 0 \\ 0 & 5 & -2 \end{pmatrix}$$

On a les valeurs propres de A_e ($\lambda_1 = -2$, $\lambda_2 = -7$, $\lambda_3 = -12$) qui sont à parties réelles strictement négatives.

Les valeurs numériques des paramètres du système et des conditions initiales sont données par le tableau ci-dessous :

paramètres	valeurs	Conditions initiales	valeurs
$\overline{Q}/\overline{V}$	2 s^{-1}	$x(0)\text{ (mol/l)}$	$(0.83, 1.2, 3)^T$
$C_{A,in}$	5 (mol/l)	$z^-(0)\text{ (mol/l)}$	$(0, 1, 0)^T$
$C_{A,in,min}$	4 (mol/l)	$z^+(0)\text{ (mol/l)}$	$(1.1, 2, 5)^T$
$C_{A,in,max}$	8 (mol/l)		
k_1	$5\text{ (s}^{-1}\text{)}$		
k_2	$10\text{ (s}^{-1}\text{)}$		

Tableau 3.2 : Valeurs des paramètres du système pour les simulations

Les simulations numériques sont réalisées en utilisant les valeurs numériques dressées dans le tableau 3.2. Les simulations sont données par la figure 3.2 :

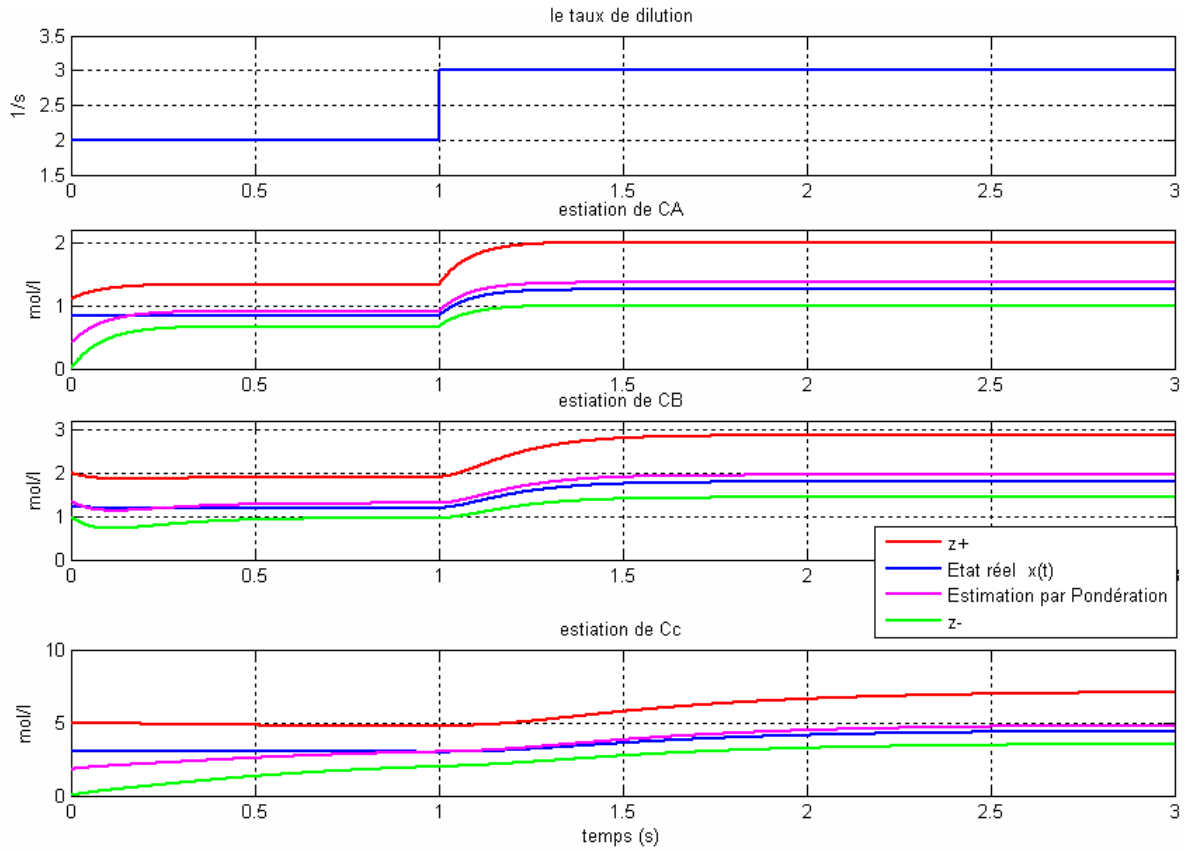


Figure 3.2 : Estimation de l'état réel $x(t)$ (en bleu) par l'observateur intervalle (en vert z^- et en rouge z^+) et par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour un choix de gain $L = [0 \ 0 \ 0]^T$.

Nous remarquons d'après la figure 3.2, que l'évolution de l'état réel $x(t)$ (en bleu) est enfermée par deux trajectoires distinctes qui sont les estimations z^- (en vert) et z^+ (en rouge) obtenues par l'observateur inférieur et supérieur respectivement de l'observateur intervalle (3.61). Et en vérifiant la condition de stabilité et de coopérativité de la matrice A_e , nous assurons la convergence asymptotique de l'estimation d'état $\hat{x}(t)$ en violet (estimation d'état obtenue par superposition pondérée de z^- et z^+) vers l'état réel $x(t)$ du système dynamique (3.60).

Cependant, la convergence de l'estimation d'état est régie par les valeurs propres de la matrice A_e et peut alors être améliorée par l'intermédiaire du choix du vecteur gain L . Il est à noter que ces valeurs propres ne peuvent pas être arbitrairement assignées même si le système est observable puisque la matrice A_e doit être coopérative. D'après les simulations numériques effectuées pour cet exemple d'application, la condition de coopérativité exige de placer des valeurs non-positives pour les paramètres L_1, L_2 du vecteur gain L .

Les résultats de simulation pour un choix du vecteur gain L suivant :

$$L = [0 \ 0 \ 10]^T$$

avec la matrice A_e calculée :

$$A_e = A - LC = \begin{pmatrix} -12 & 0 & 0 \\ 10 & -7 & 0 \\ 0 & 5 & -12 \end{pmatrix}$$

qui est une matrice stable (les valeurs propres de A_e ($\lambda_1 = -12$, $\lambda_2 = -7$, $\lambda_3 = -12$) sont à parties réelles strictement négatives), sont donnés par la figure 3.3 ci-dessous :

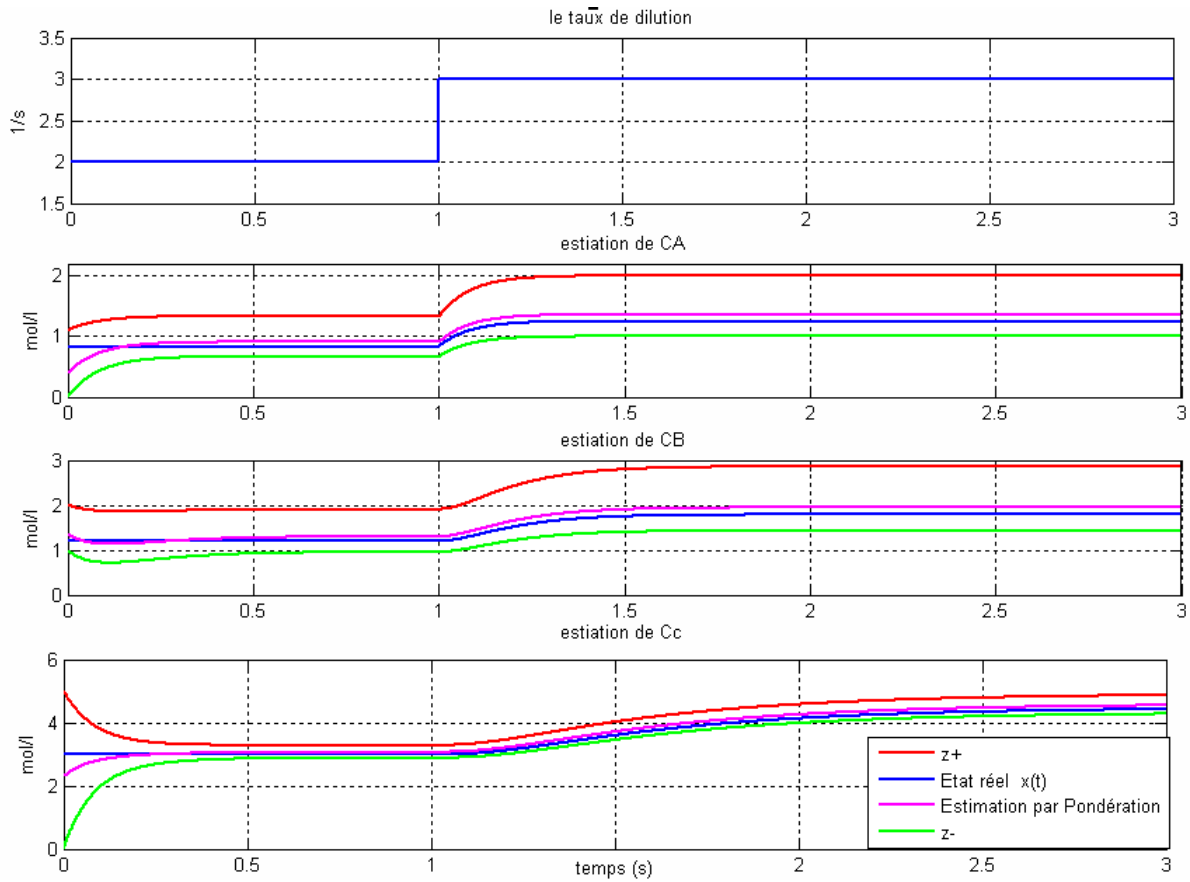


Figure 3.3 : Estimation de l'état réel $x(t)$ (en bleu) par l'observateur intervalle (en vert z^- et en rouge z^+) et par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour un choix de gain $L = [0 \ 0 \ 10]^T$.

Comme le montre la figure 3.3, les performances d'estimation peuvent être améliorées en plaçant L_1, L_2 égaux à zéro et en augmentant L_3 . Avec un choix du gain $L = [0 \ 0 \ 10]^T$, nous avons pu améliorer le taux de convergence de l'état estimé $\hat{x}(t)$ (en violet) vers l'état réel $x(t)$ (en bleu) du système dynamique.

Par la suite, nous effectuons une variation du paramètre incertain $C_{A,in}$, tout en gardant la relation (3.59) toujours vérifiée (les bornes $C_{A,in,min}$, $C_{A,in,max}$ restent inchangées). Et vue l'amélioration du taux de convergence en utilisant le gain $L = [0 \ 0 \ 10]^T$, celui-ci sera choisi pour les simulations.

Nous traçons les estimations d'état des trois variables C_A , C_B et C_C par l'observateur final (obtenu par superposition pondérée des limites supérieure et inférieure de l'observateur intervalle). Et nous traçons en même temps l'évolution de l'erreur $e = x - \hat{x}$ correspondante par variation du paramètre incertain $C_{A,in}$:

Pour la valeur du paramètre incertain $C_{A,in} = 4.5$, nous obtenons les tracés suivants :

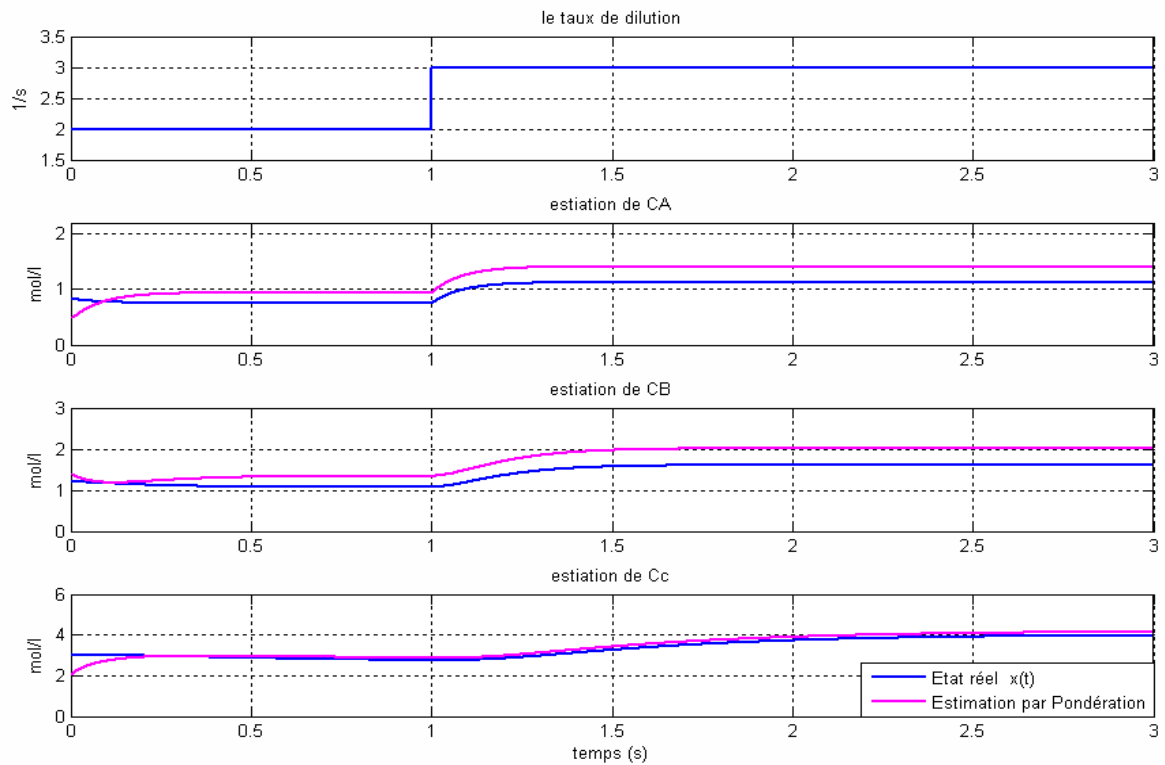


Figure 3.4 : Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 4.5$.

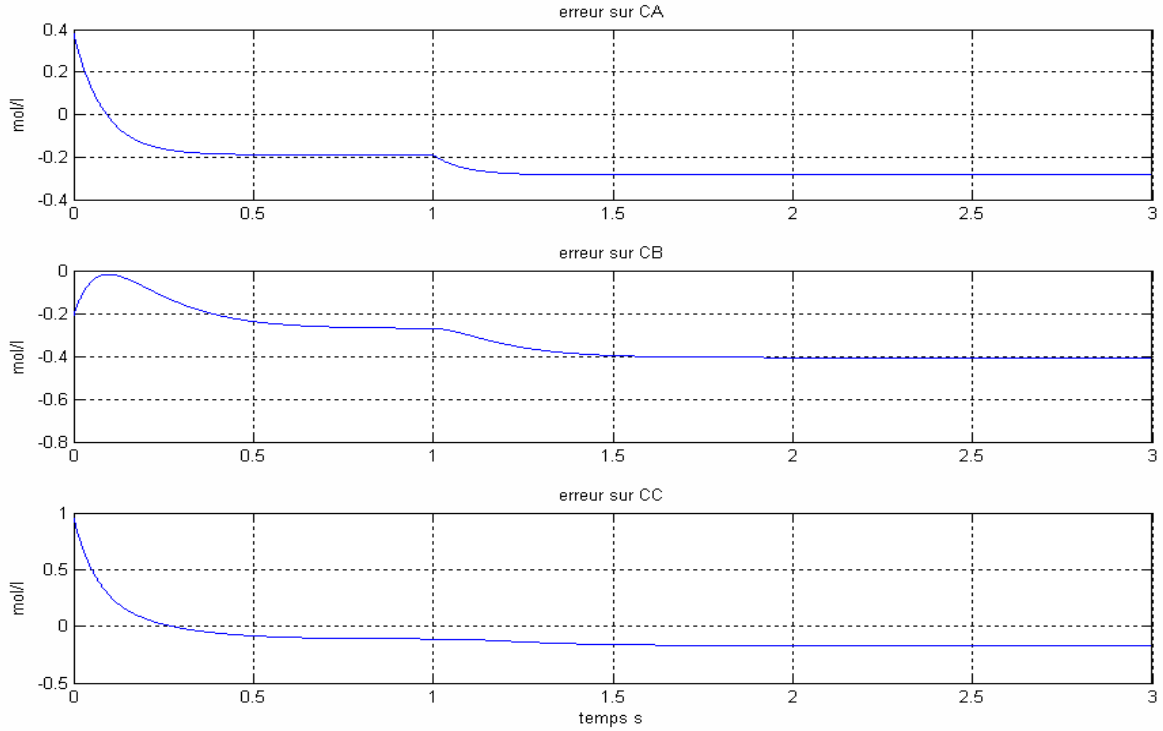


Figure 3.5 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 4.5$.

Pour la valeur du paramètre incertain $C_{A,in} = 5.5$, nous obtenons les tracés suivants :

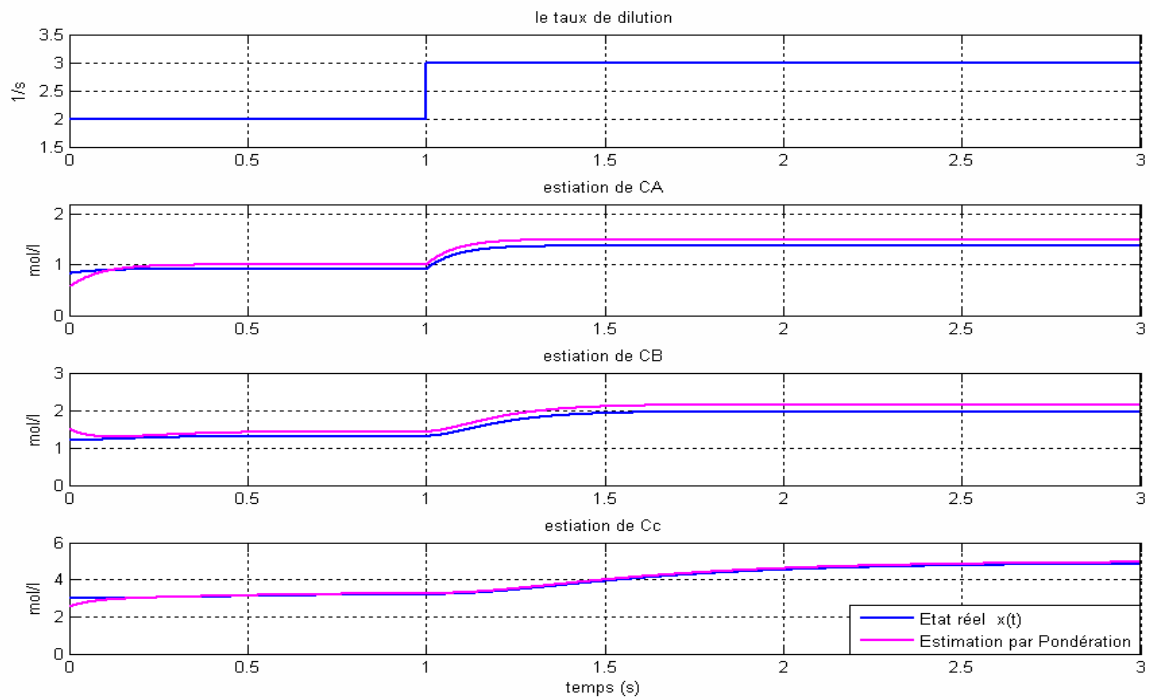


Figure 3.6 : Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 5.5$.

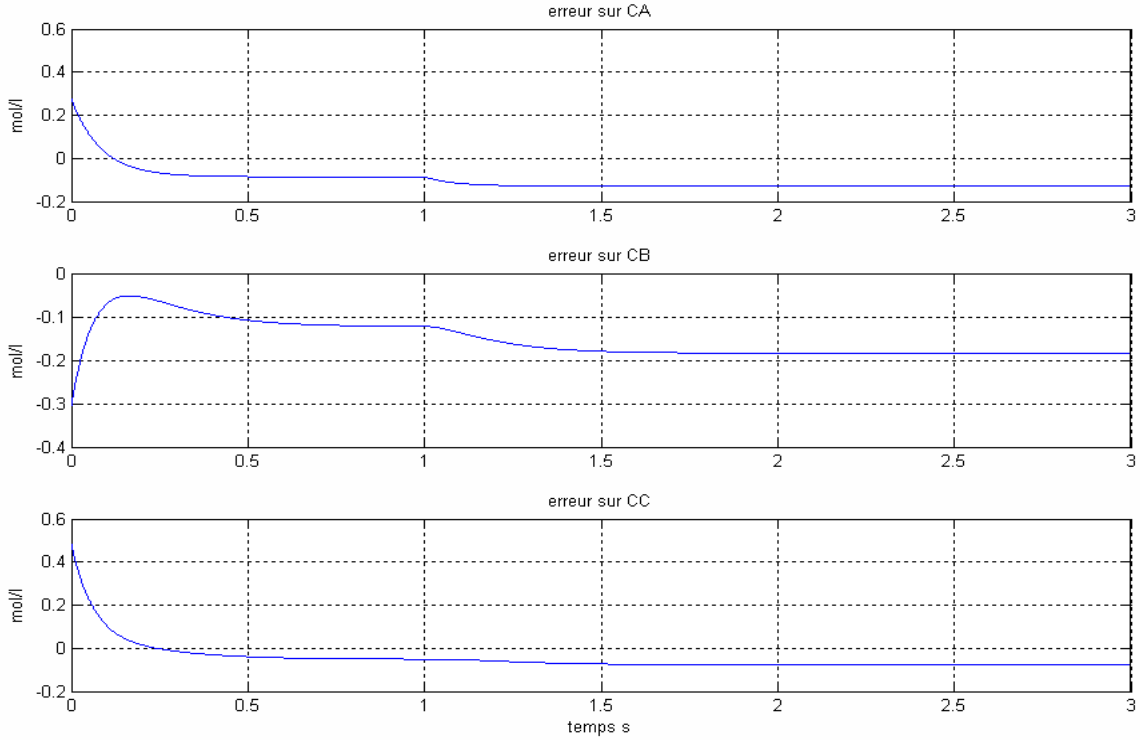


Figure 3.7 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 5.5$.

Pour la valeur du paramètre incertain $C_{A,in} = 6$, nous obtenons les tracés suivants :

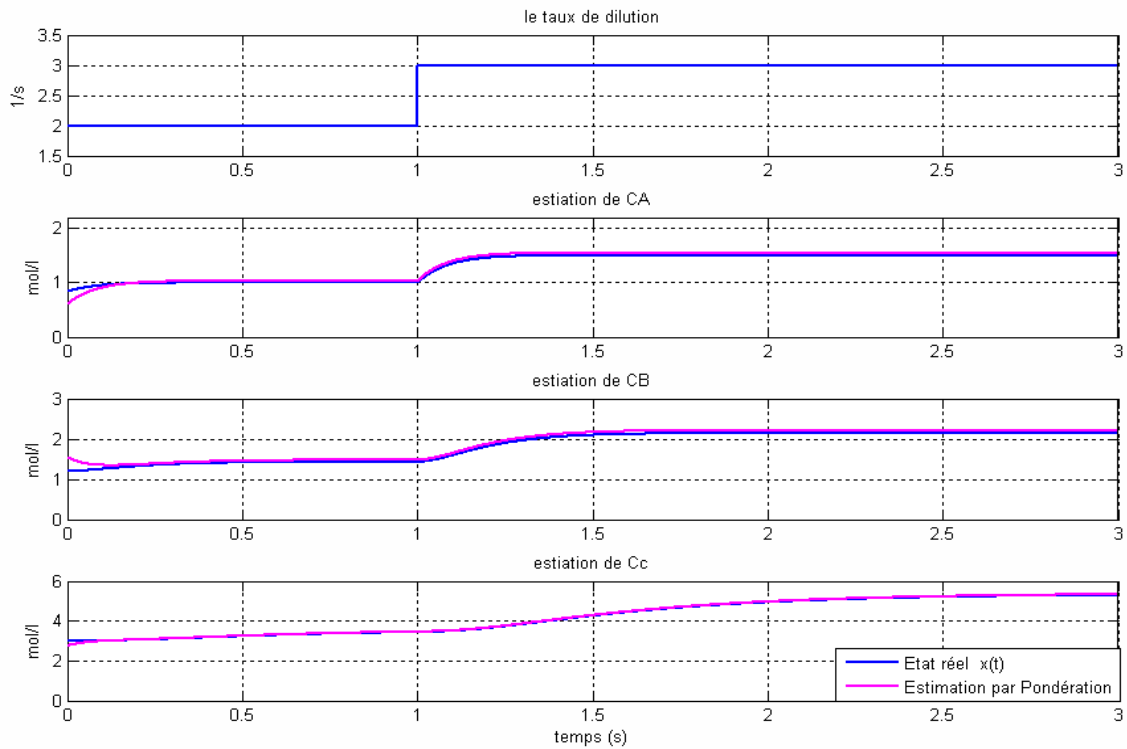


Figure 3.8 : Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 6$.

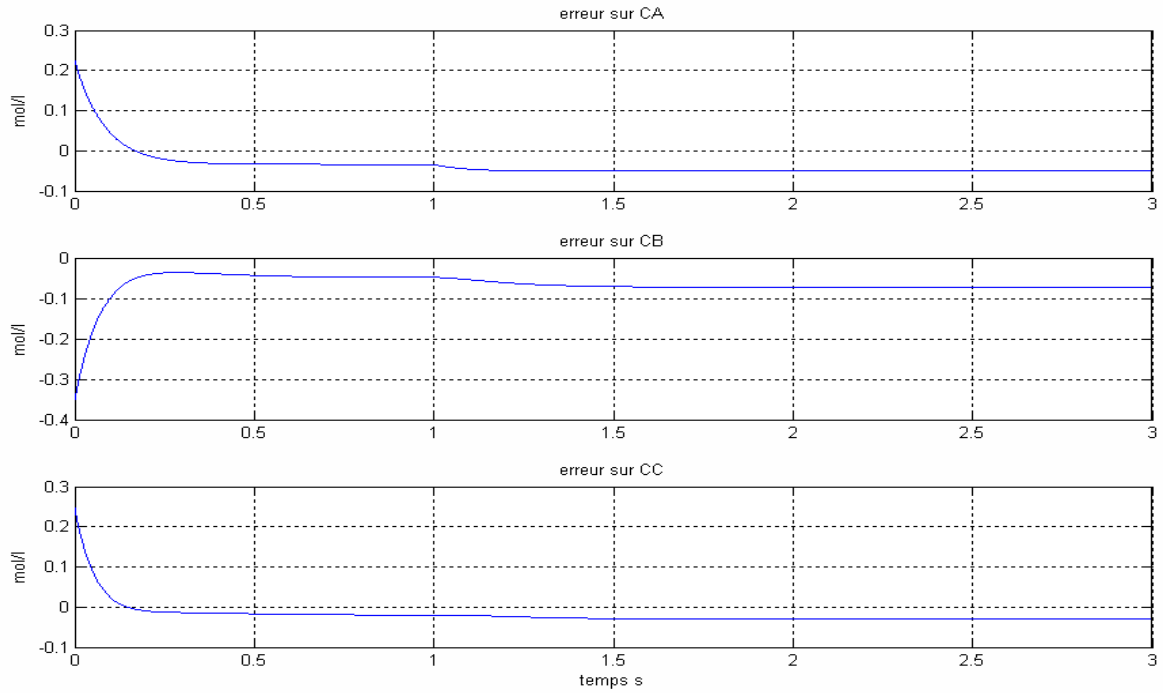


Figure 3.9 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 6$.

Pour la valeur du paramètre incertain $C_{A,in} = 7$, nous obtenons les tracés suivants :

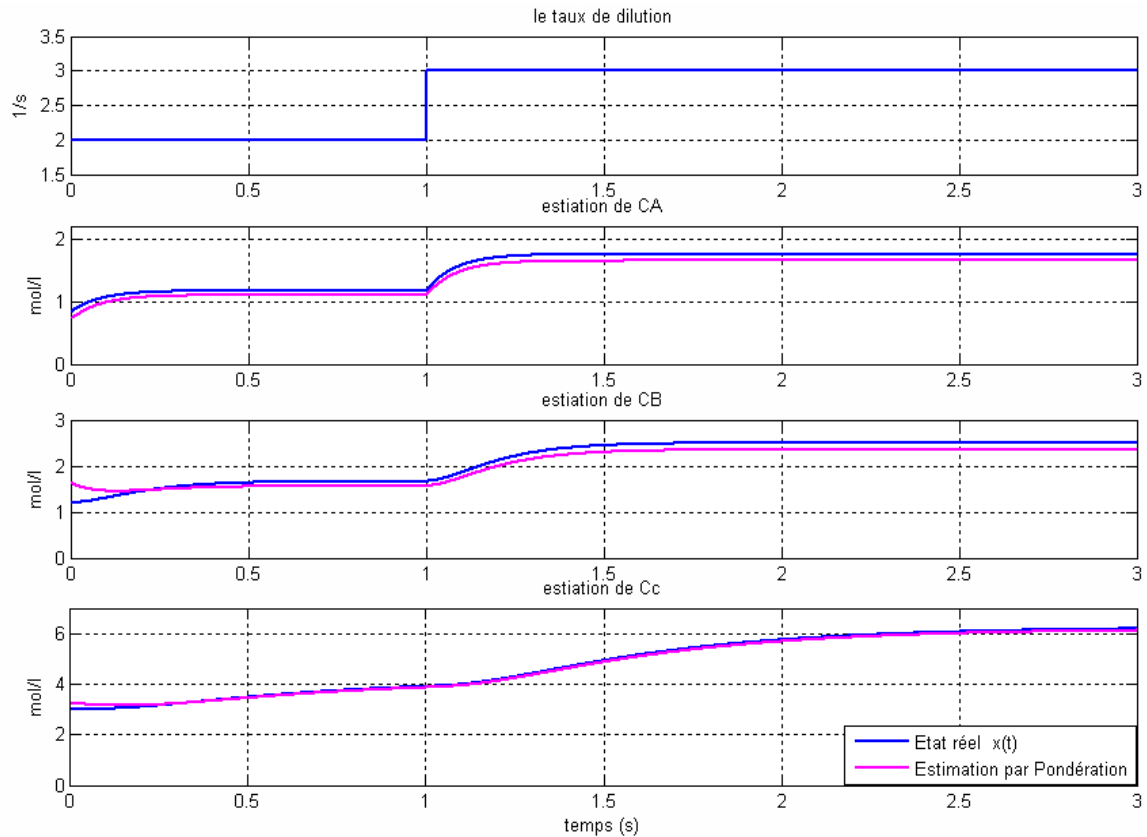


Figure 3.10: Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 7$.

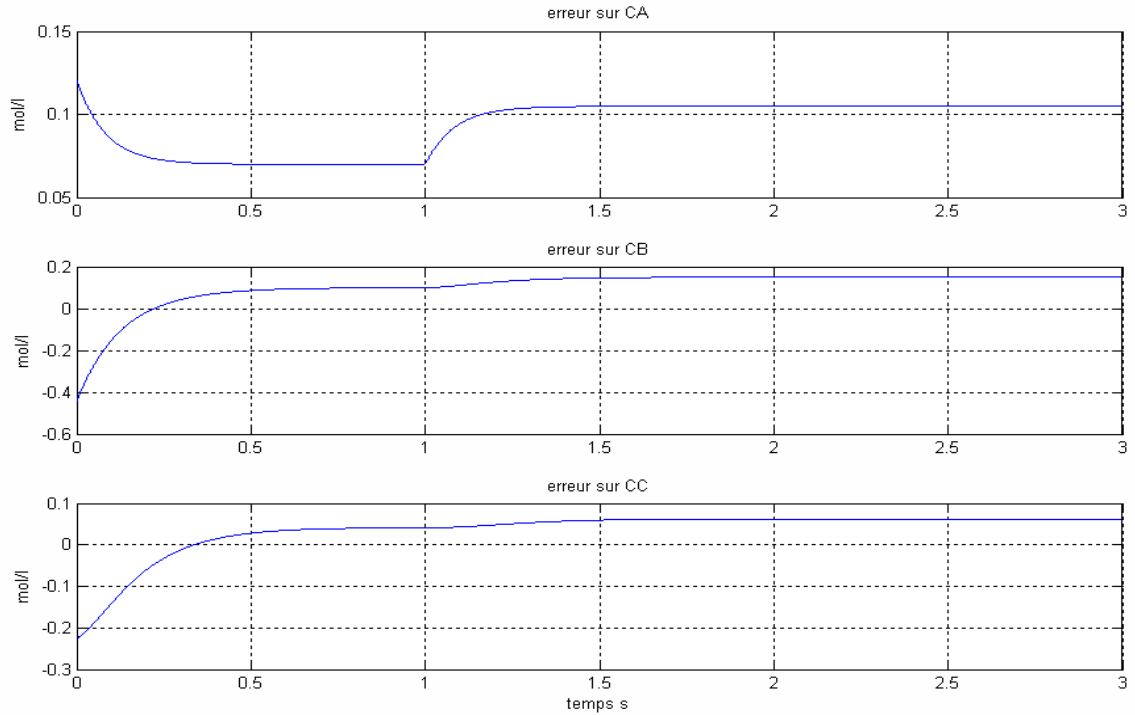


Figure 3.11 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 7$.

Pour la valeur du paramètre incertain $C_{A,in} = 7.5$, nous obtenons les tracés suivants :

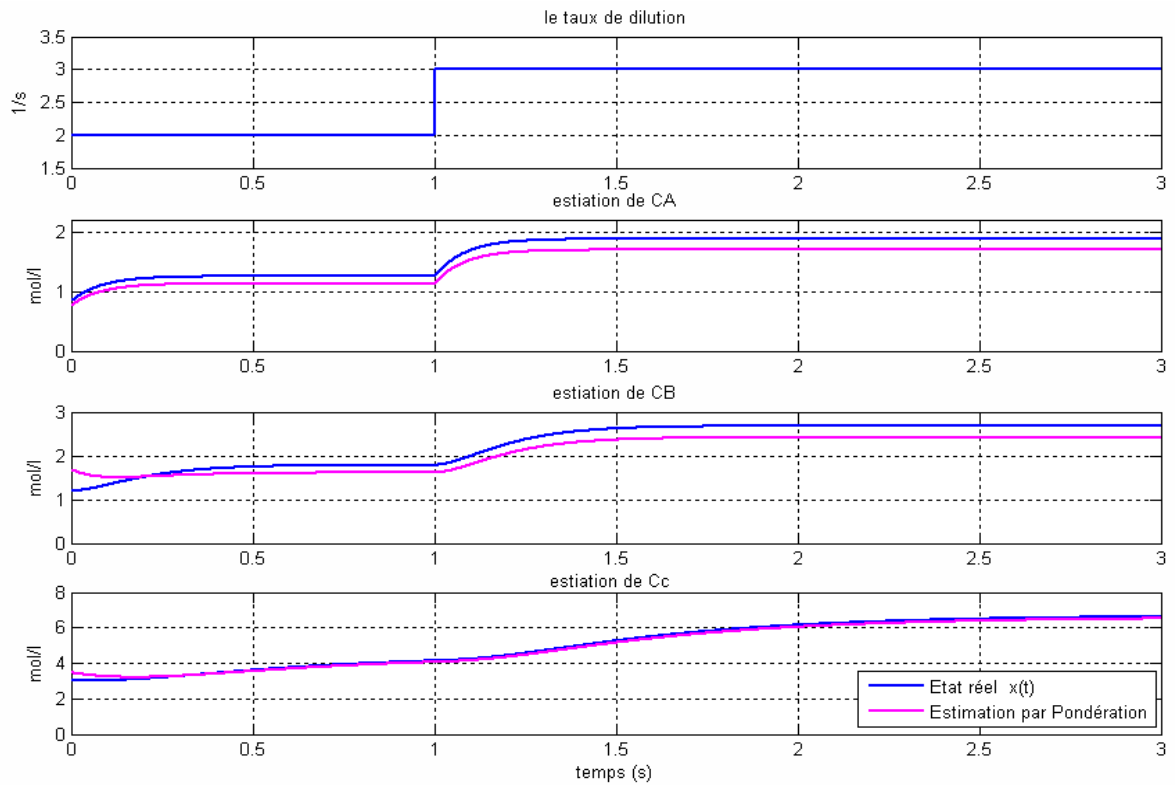


Figure 3.12 : Estimation de l'état réel $x(t)$ (en bleu) par l'estimation obtenue par pondération $\hat{x}(t)$ (en violet) pour une valeur du paramètre incertain $C_{A,in} = 7.5$.

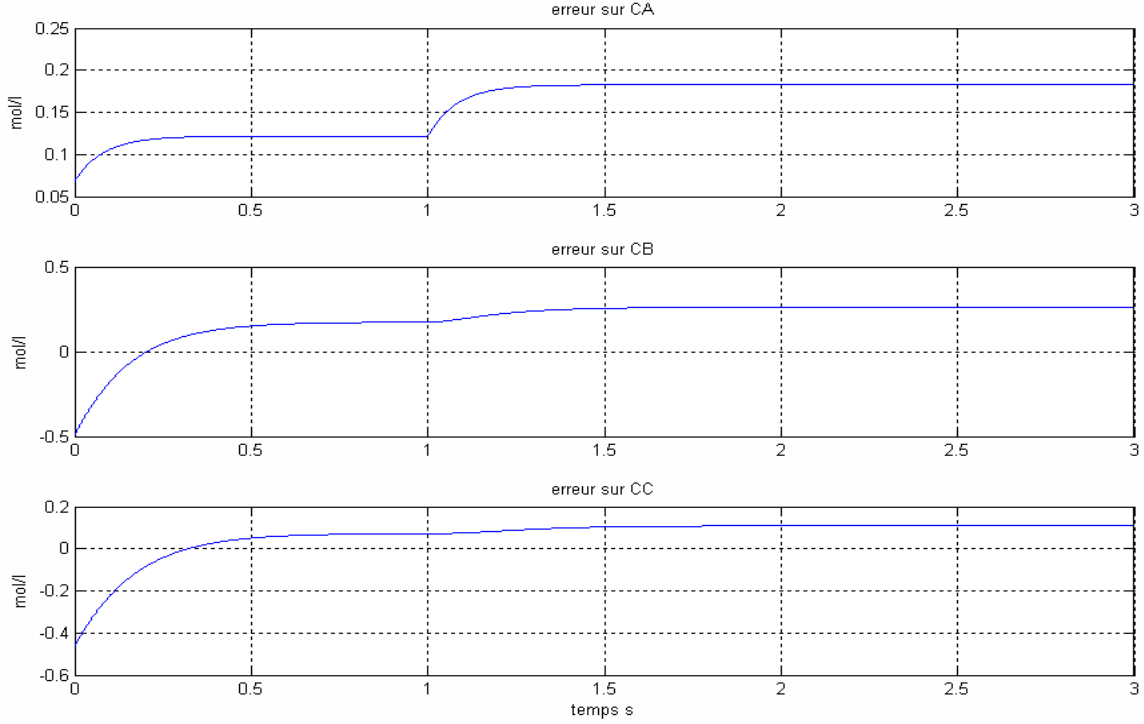


Figure 3.13 : Tracé de l'évolution de l'erreur $e = x - \hat{x}$ pour une valeur du paramètre incertain $C_{A,in} = 7.5$.

Nous remarquons d'après les simulations effectuées (tracés des estimations d'état et de l'erreur), qu'en variant le paramètre incertain $C_{A,in}$, en s'éloignant vers la limite inférieure $C_{A,in,min}$ (valeurs prises pour $C_{A,in}$ sont 4.5 et 5.5) et la limite supérieure $C_{A,in,max}$ (valeurs prises pour $C_{A,in}$ sont 7 et 7.5) que les performances de l'estimateur $\hat{x}(t)$ (obtenu par superposition pondérée des limites supérieure et inférieure de l'observateur intervalle) diminuent. Ici nous testons la convergence de l'état estimé vers l'état réel du système dynamique et sa rapidité.

Par contre, en prenant la valeur $C_{A,in} = 6$, nous avons pu améliorer la convergence asymptotique de l'état estimé $\hat{x}(t)$ vers l'état réel du système dynamique, et l'observateur est plus rapide. Nous avons l'évolution de l'erreur $e = x - \hat{x}$ qui tend vers zero plus rapidement.

Si nous considérons une autre écriture de la relation (3.59), que nous réécrivons :

$$C_{A,in,min} \leq C_{A,in} \leq C_{A,in,max}$$

comme suit :

$$[a] = [C_{A,in,min}, C_{A,in,max}] = [4,8].$$

alors, $C_{A,in}$ représente la variable incertaine appartenant au support $[a]$. La valeur $C_{A,in} = 6$ représente en réalité l'estimé de la variable incertaine $C_{A,in}$. Car cette valeur représente le milieu de l'intervalle $[a]$ ($mid[a] = 6$). Pour cette valeur, nous avons pu améliorer la convergence et la rapidité de l'estimation d'état $\hat{x}(t)$ vers l'état réel $x(t)$ du système dynamique.

3.4.2 Observateur intervalle pour les systèmes linéaires temps invariant en présence des perturbations

Sachant que tous les systèmes ne vérifient pas la condition de coopérativité définie précédemment, l'observateur intervalle développé dans cette section permet de palier cette difficulté, en montrant que tout système linéaire temps invariant exponentiellement stable non coopératif avec des perturbations additives peut être transformé en un système coopératif par un changement de coordonnées « temps variant ». L'observateur issu de cette transformation est un observateur intervalle « temps variant » exponentiellement stable. La technique s'appuie sur une transformation sous la forme canonique de Jordan.

Le théorème suivant [Mazenc et Bernard, 2011] explicite cette nouvelle technique. Un exemple d'application est illustré sur un modèle d'état d'un système LTI stable en boucle ouverte (d'ordre 3) présentant des perturbations additives mal connues.

Théorème 3.5 [Mazenc et Bernard, 2011]

Soit le système suivant

$$\dot{x} = Ax + w(t) \quad (3.68)$$

avec

$x \in R^n$, A est une matrice stable et constante, $w(t) \in R^n$ est une fonction Lipschitz continue inconnue bornée par deux fonctions Lipschitz continues connues comme suit :

$$\forall t \geq 0, w^-(t) \leq w(t) \leq w^+(t)$$

alors, il existe $P: R \rightarrow R^{n \times n}$ de classe C^∞ de norme finie tel que : $\forall t \in R^+$, $P(t)$ est inversible coopérative et stable. Et une matrice constante $N \in R^{n \times n}$ tel que :

$\forall t \in R$, $P(t)$ est une solution de $\dot{P}(t) = N P(t) - P(t)A$. L'observateur suivant est un observateur intervalle exponentiellement stable pour le système (3.68) :

$$\begin{cases} \dot{z}^+ = N z^+ + P^+(t) w^+(t) - P^-(t) w^-(t) \\ \dot{z}^- = N z^- + P^+(t) w^-(t) - P^-(t) w^+(t) \\ z^+(t_0) = P^+(t_0)x_0^+ - P^-(t_0)x_0^- \\ z^-(t_0) = P^+(t_0)x_0^- - P^-(t_0)x_0^+ \\ x^+ = M^+(t)z^+ - M^-(t)z^- \\ x^- = M^+(t)z^- - M^-(t)z^+ \end{cases} \quad (3.69)$$

où $P^+(t) = \max(0, P(t))$, $P^-(t) = P^+(t) - P(t)$ et la $M(t)$ est l'inverse de la matrice $P(t)$, $M^+(t) = \max(0, M(t))$ et $M^-(t) = M^+(t) - M(t)$.

a) Exemple d'application sur un système d'ordre trois

Exemple 3.4 : Nous allons montrer à travers cet exemple, les différentes étapes de construction de l'observateur intervalle.

Considérons le système suivant :

$$\dot{x} = Ax + w(t)$$

où la matrice d'état A est donnée par :

$$A = \begin{pmatrix} -1/3 & 0 & -4/3 \\ 6/3 & -9/3 & 0 \\ 10/3 & 0 & -5/3 \end{pmatrix}$$

La matrice A est non coopérative. Donc, afin d'appliquer le théorème précédent pour la synthèse d'un observateur intervalle, il faudra vérifier la stabilité de A .

Les valeurs propres de A ($\lambda_1 = -3$, $\lambda_2 = -1 + 2i$ et $\lambda_3 = -1 - 2i$) sont à parties réelles strictement négatives, donc la matrice A est Hurwitz (stable).

L'incertitude sur le système se traduit dans le fait que l'entrée $w(t)$ est une perturbation inconnue mais qui est bornée par deux fonctions Lipschitz continus connues w^+ et w^- comme suit:

$$\begin{cases} w^-(t) \leq w(t) \leq w^+(t) \\ w^+(t) = w(t) + \Gamma \\ w^-(t) = w(t) - \Gamma \end{cases} \quad (3.70)$$

Γ : est un paramètre variable permettant de considérer plusieurs bornes inférieure et supérieure pour w^- et w^+ respectivement.

Nous prenons comme exemple pour $w(t)$, un bruit uniformément distribué dans l'intervalle $[-2,1]$.

- Dans MATLAB, $w(t)$ est obtenu par la commande « rand » et par une interpolation linéaire de $w(t_i)$ entre deux vecteurs temps (tp, t) de période d'échantillonnage différentes

Par exemple :

```
t=linspace(0,20,100);
tp=linspace(0,20,200);
w = -2 + 3.*rand(size(t));
wti= interp1(t,w,tp) ;
```

$w^+(t)$ et $w^-(t)$ peuvent être considérés comme des vecteurs intervalles dont la dimension est dépendante de la longueur du vecteur temps.

Considérons le vecteur d'état initial donné par le vecteur intervalle suivant :

$$x_0 \in ([-0.05, 0.05], [-0.05, 0.05], [-0.05, 0.05])^T \quad (3.71)$$

x_0 peut être réécrit de la manière suivante :

$$\begin{aligned} x_0^- &\leq x_0 \leq x_0^+ \\ x_0^+ &= (0.05, 0.05, 0.05)^T \\ x_0^- &= (-0.05, -0.05, -0.05)^T \end{aligned} \quad (3.71)'$$

b) Etapes de construction de l'observateur intervalle temps-variant

Nous développons ci-dessous les différentes étapes pour la construction de l'observateur intervalle du système incertain

b.1) Changement de coordonnées temps-invariant par la forme canonique de Jordan

La matrice A peut être transformée en une matrice de forme de Jordan [Perko, 1991] en considérant deux entiers $r \in \{0, 1, \dots, n\}$ et $s \in \{0, 1, \dots, n-1\}$ et un changement de base « temps-invariant » comme suit

$$Y = Px$$

qui transforme le système

$$\dot{x} = Ax$$

en

$$\dot{Y} = JY$$

avec

$$J = \begin{pmatrix} J_1 & 0 & \dots & 0 \\ 0 & J_2 & \dots & 0 \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & J_s \end{pmatrix} \in R^{n \times n} \quad (3.72)$$

où les matrices J_i sont partitionnées en deux groupes, les r premières matrices sont associées aux r valeurs propres réelles de multiplicité n_i de A et les restantes aux valeurs propres complexes de multiplicité m_i de A . Avec $n = \sum_{i=1}^r n_i + \sum_{r+1}^s 2m_i$

Pour $i = \overline{1, r}$, nous avons :

$$J_i = \begin{pmatrix} -\mu_i & 1 & \dots & 0 \\ 0 & -\mu_i & \dots & 0 \\ \vdots & \ddots & \ddots & 1 \\ 0 & \dots & 0 & -\mu_i \end{pmatrix} \in R^{n_i \times n_i} \quad (3.73)$$

où les μ_i sont des nombres réels positifs.

Et pour $i = \overline{r+1, s}$, nous avons :

$$J_i = \begin{pmatrix} \Lambda_i & I_2 & 0 & \cdots & 0 \\ 0 & \Lambda_i & I_2 & \ddots & \vdots \\ \vdots & \ddots & \ddots & \ddots & 0 \\ \vdots & & & \ddots & I_2 \\ 0 & \dots & \dots & 0 & \Lambda_i \end{pmatrix} \in R^{2m_i \times 2m_i} \quad (3.74)$$

avec $\Lambda_i = \begin{pmatrix} -\beta_i & \gamma_i \\ -\gamma_i & -\beta_i \end{pmatrix}$ et $I_2 = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix}$

où les β_i sont des nombres réels positifs et γ_i sont des nombres réels différents de zéro (β_i et γ_i représentent les parties réelles et imaginaires respectivement des valeurs propres complexes de la matrice A).

Dans notre exemple la matrice P calculée est la suivante :

$$P = \frac{1}{3} \begin{pmatrix} 1 & -3 & 1 \\ 1 & 0 & 1 \\ 2 & 0 & -1 \end{pmatrix}$$

Pour un nombre de valeurs propres réelles $r = 1$ de multiplicité $n_i = 1$ on a :

$$J_1 = (-3) \in R^{1 \times 1}$$

et pour $i = \overline{r+1, 3}$, $m_i = 1$ est la multiplicité des valeurs propres complexes :

$$J_2 = \begin{pmatrix} -1 & 2 \\ -2 & -1 \end{pmatrix} \in R^{2 \times 2}$$

La forme canonique de Jordon calculée est alors donnée par :

$$J = \begin{pmatrix} -3 & 0 & 0 \\ 0 & -1 & 2 \\ 0 & -2 & -1 \end{pmatrix}$$

tel que l'égalité suivante:

$$P A P^{-1} = J \quad (3.75)$$

est vérifiée. La matrice J obtenue n'est pas coopérative, donc il faudra effectuer un changement de base « temps-variant ».

b.2) Le changement de coordonnées temps-variant

Le changement de coordonnées considéré est le suivant :

$$P(t) = n(t) P \quad (3.76)$$

La matrice $n(t)$ est donnée par :

$$n(t) = \begin{pmatrix} I_r & 0 & \dots & 0 \\ 0 & \mathcal{N}_{r+1}(t) & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \mathcal{N}_s(t) \end{pmatrix} \quad (3.77)$$

avec, la matrice $I_r \in R^{r \times r}$, correspond à la matrice identité associée aux r valeurs propres réelles de la matrice A . Les matrices $\mathcal{N}_i(t)$ avec $i = \overline{r+1, s}$ associées aux valeurs propres complexes sont calculées de la manière suivante :

$$\mathcal{N}_i(t) = \begin{pmatrix} \mathcal{W}_i(t) & 0 & \dots & 0 \\ 0 & \mathcal{W}_i(t) & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \mathcal{W}_i(t) \end{pmatrix} \in R^{2m_i \times 2m_i} \quad (3.78)$$

avec

$$\mathcal{W}_i(t) = \begin{pmatrix} \cos(\gamma_i t) & -\sin(\gamma_i t) \\ \sin(\gamma_i t) & \cos(\gamma_i t) \end{pmatrix} \quad (3.79)$$

et γ_i est la partie imaginaire des valeurs propres complexes de A , dans cet exemple γ_i est égale à 2 ($\lambda_{2,3} = -1 \pm 2i$).

Alors la matrice $n(t)$, d'après (3.77) est égale à :

$$n(t) = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos(2t) & -\sin(2t) \\ 0 & \sin(2t) & \cos(2t) \end{pmatrix}$$

La matrice $P(t)$ calculée par (3.76) est :

$$P(t) = \frac{1}{3} \begin{pmatrix} 1 & -3 & 1 \\ \cos(2t) - 2\sin(2t) & 0 & \cos(2t) - \sin(2t) \\ \sin(2t) + 2\cos(2t) & 0 & \sin(2t) - \cos(2t) \end{pmatrix}$$

La matrice $P(t)$ est la solution de l'équation différentielle :

$$\dot{P}(t) = N P(t) - P(t)A .$$

où la matrice N obtenue est une matrice constante et qui peut être calculée de la manière suivante :

$$N = \begin{pmatrix} G & 0 & 0 & 0 \\ 0 & \mathcal{M}_{r+1} & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ 0 & \dots & 0 & \mathcal{M}_s \end{pmatrix} \quad (3.80)$$

avec G représente les r matrices J_i (3.73)

Et pour $i = \overline{r+1, s}$, les matrices $\mathcal{M}_i$ sont données par :

$$\mathcal{M}_i = \begin{bmatrix} -\beta_i I_2 & I_2 & 0 & \dots & 0 \\ 0 & -\beta_i I_2 & I_2 & \ddots & \vdots \\ \vdots & \ddots & \ddots & \ddots & 0 \\ \vdots & & \ddots & \ddots & I_2 \\ 0 & \dots & \dots & 0 & -\beta_i I_2 \end{bmatrix} \quad (3.81)$$

Nous obtenons après calcul, la matrice N suivante :

$$N = \begin{pmatrix} -3 & 0 & 0 \\ 0 & -1 & 0 \\ 0 & 0 & -1 \end{pmatrix}$$

La matrice $M(t) = P(t)^{-1}$ calculée est :

$$M(t) = P(t)^{-1} = \begin{pmatrix} 0 & \cos(2t) - \sin(2t) & \cos(t) + \sin(t) \\ -1 & \cos(2t) & \sin(2t) \\ 0 & 2\cos(2t) + \sin(t) & 2\sin(2t) - \cos(2t) \end{pmatrix}$$

Cette dernière sera utilisée dans l'étape suivante pour la construction de l'observateur intervalle :

b.3) L'observateur intervalle

Selon le théorème 3.5, le système admet l'observateur intervalle temps-variant suivant:

$$\begin{cases} \dot{z}^+ = N z^+ + P^+(t) w^+(t) - P^-(t) w^-(t) \\ \dot{z}^- = N z^- + P^+(t) w^-(t) - P^-(t) w^+(t) \\ z^+(t_0) = P^+(t_0)x_0^+ - P^-(t_0)x_0^- \\ z^-(t_0) = P^+(t_0)x_0^- - P^-(t_0)x_0^+ \end{cases} \quad (3.82)$$

avec

$$P^+(t) = \max(0, P(t)) = \frac{1}{3} \begin{pmatrix} 1 & 0 & 1 \\ \mathcal{L}(\cos(2t) - 2\sin(2t)) & 0 & \mathcal{L}(\cos(2t) - \sin(2t)) \\ \mathcal{L}(\sin(2t) + 2\cos(2t)) & 0 & \mathcal{L}(\sin(2t) - \cos(2t)) \end{pmatrix}$$

où $\mathcal{L}(x(t))$ représente le $\max_t(0, x(t))$.

$$P^-(t) = P^+(t) - P(t)$$

À l'instant $t_0 = 0$ la matrice P^+ et P^- calculées sont données ci-dessous :

$$P^+ = \begin{pmatrix} 0.3333 & 0 & 0.3333 \\ 0.3333 & 0 & 0.3333 \\ 0.6667 & 0 & 0 \end{pmatrix}$$

et

$$P^- = P^+ - P = \begin{pmatrix} 0 & 1 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & 0.3333 \end{pmatrix}$$

En prenant comme conditions initiales pour le vecteur d'état x_0 (3.71)' à l'instant $t_0 = 0$, nous obtenons comme conditions initiales pour l'observateur intervalle (3.82) les vecteurs colonnes suivants :

$$\begin{aligned} z^+(t_0) &= (0.0833, 0.0333, 0.0500)^T \\ z^-(t_0) &= (-0.0833, -0.0333, -0.0500)^T \end{aligned}$$

$z^+(t)$ et $z^-(t)$ qui représentent respectivement l'estimation d'état par l'observateur supérieur et l'observateur inférieur de l'observateur intervalle, sont obtenues en résolvant le système d'équations différentielles (3.82) .

Enfin, l'estimation d'état finale est donnée par les équations suivantes :

$$\begin{aligned} x^+ &= M^+(t)z^+ - M^-(t)z^- \\ x^- &= M^+(t)z^- - M^-(t)z^+ \end{aligned}$$

avec $M^+(t) = \max(0, M(t))$ est donnée par:

$$M^+(t) = \begin{pmatrix} 0 & \mathcal{L}(\cos(2t) - \sin(2t)) & \mathcal{L}(\cos(t) + \sin(t)) \\ 0 & \mathcal{L}(\cos(2t)) & \mathcal{L}(\sin(2t)) \\ 0 & \mathcal{L}(2\cos(2t) + \sin(t)) & \mathcal{L}(2\sin(2t) - \cos(2t)) \end{pmatrix}$$

et $M^-(t) = M^+(t) - M(t)$

Les figures suivantes représentent l'estimation d'état du système par l'observateur intervalle, en présence d'incertitudes sur l'entrée $w(t)$. Nous faisons varier à chaque fois la valeur de w^+ et w^- par le paramètre variable Γ (3.70) :

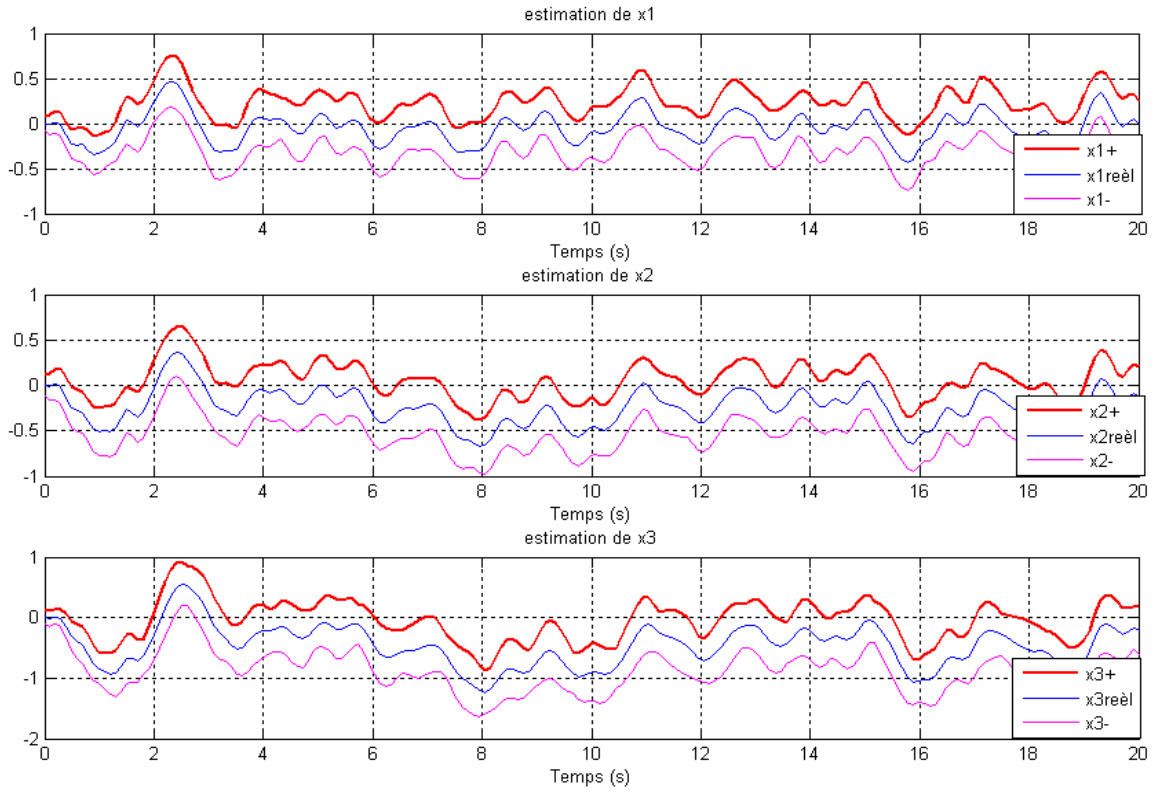


Figure 3.14 : Estimation de l'état réel (en bleu) par un observateur intervalle (l'observateur supérieur en rouge et l'observateur inférieure en violet) et pour $w^+ = w(t) + 0.2$, $w^- = w(t) - 0.2$.

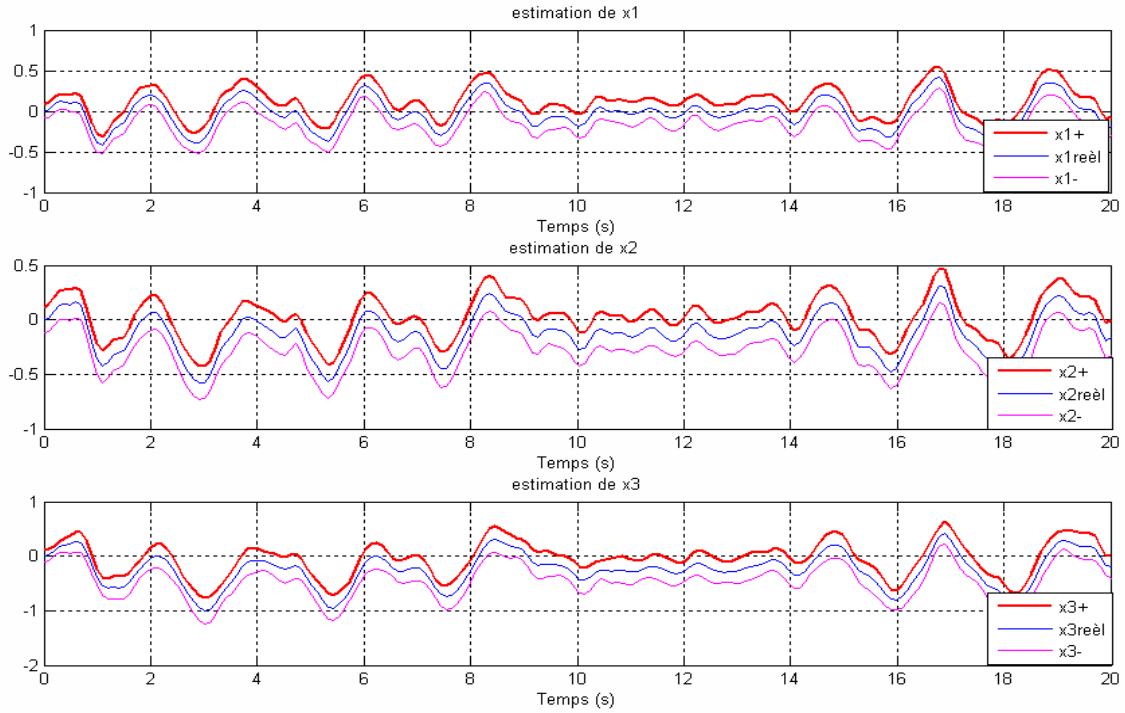


Figure 3.15 : Estimation de l'état réel (en bleu) par un observateur intervalle (l'observateur supérieur en rouge et l'observateur inferieur en violet) et pour $w^+ = w(t) + 0.1$, $w^- = w(t) - 0.1$.

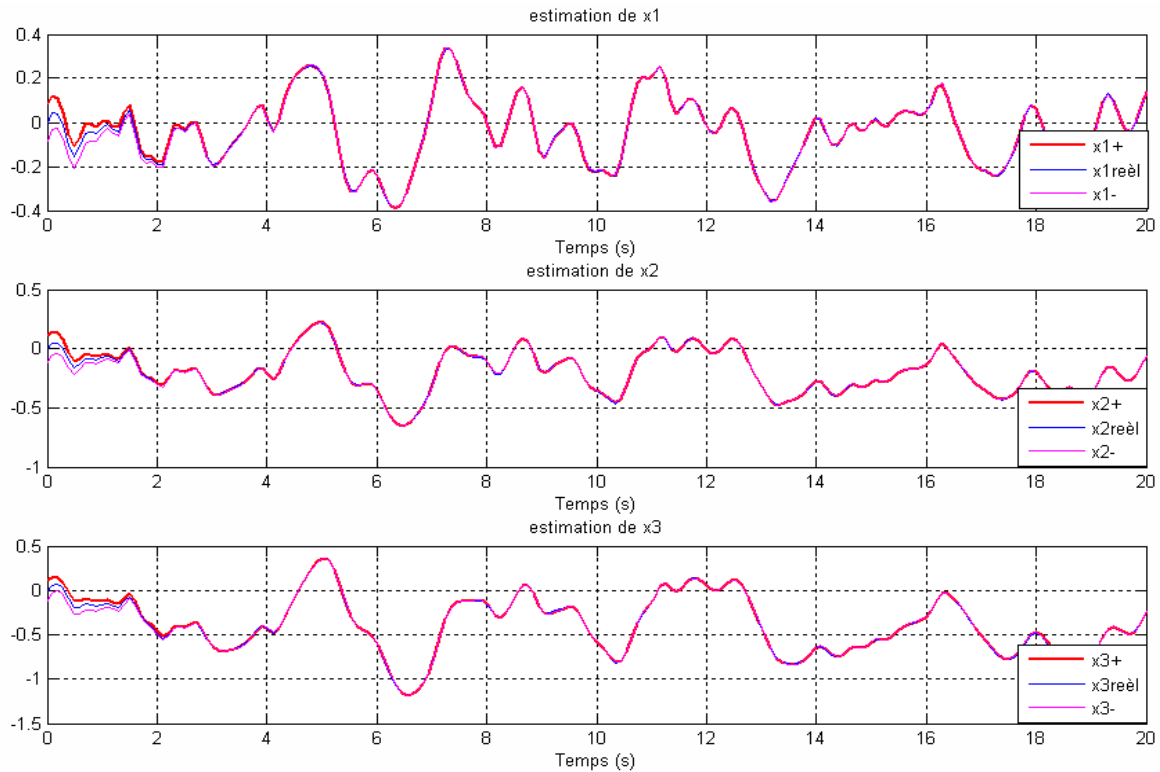


Figure 3.16 : Estimation de l'état réel (en bleu) par un observateur intervalle (l'observateur supérieur en rouge et l'observateur inferieur en violet), pour $w^+ = w(t) + 0$, $w^- = w(t) - 0$.

D'après les simulations réalisées, nous avons l'évolution de l'état réel $x(t)$ (en bleu) du système incertain (conditions initiales inconnues mais qui sont bornées et des perturbations additives bornées) qui est enfermée par les estimations supérieure et inférieure obtenues par l'observateur intervalle. Donc, cet observateur intervalle nous garantit d'envelopper l'évolution de l'état réel par les deux trajectoires supérieure et inférieure.

Nous remarquons qu'en variant le paramètre Γ pour les deux fonctions $w^+(t)$ et $w^-(t)$ dans l'équation (3.70) en le réduisant, nous obtenons la convergence de l'état estimé par l'observateur supérieur $x^+(t)$ (en rouge) et l'état estimé par l'observateur inférieur $x^-(t)$ (en violet) vers l'état réel du système $x(t)$ en bleu. Et pour la valeur de $\Gamma = 0$, l'estimation supérieure et l'estimation inférieure sont confondues avec l'état réel du système. Donc l'observateur intervalle obtenu est asymptotiquement stable.

3.5 Conclusion

Dans ce chapitre, nous avons introduit l'outil de l'arithmétique des intervalles pour redéfinir les propriétés structurelles de commandabilité et l'observabilité des systèmes LTI à paramètres incertains et ceci en se basant sur la propriété de l'indépendance (dépendance) linéaire des vecteurs intervalles.

L'utilisation de cette arithmétique dans le domaine d'estimation d'état a permis de prendre en charge les incertitudes influant sur un système donné. Deux observateurs intervalles ont été développés, le premier a été conçu pour la classe des systèmes LTI coopératifs à paramètres incertains. Cet observateur offre une estimation d'état convergente vers l'état réel et ceci au moyen d'une estimation d'état pondérée calculée en utilisant la sortie mesurée du système et les estimations supérieure et inférieure retrouvées par l'observateur supérieur et inférieur respectivement de l'observateur intervalle.

Cependant, la limitation de cette méthode d'estimation par cet observateur intervalle est qu'il est appliqué à une classe des systèmes exigeant de vérifier la condition de coopérativité. Le deuxième observateur intervalle développé dans ce chapitre permet de palier cette difficulté par la construction d'un observateur intervalle temps-variant exponentiellement stable pour un système linéaire stable en boucle ouverte. Et ceci au moyen d'un changement de base « temps-variant » qui transforme le système en un système coopératif.

Pour un système instable en boucle ouverte, on peut toujours associer « observateur classique et observateur intervalle » pour une estimation d'état convergente. On procède en deux étapes, on élabore un observateur classique de type Luenberger par exemple (le système doit être observable) en boucle fermée. Puis, on procède par un changement de base temps variant pour rendre coopératif le système d'erreur $e(t)$.

Ce dernier observateur sera appliqué en association avec un observateur Luenberger classique pour l'estimation d'état d'un pendule inversé dans le chapitre suivant.

Chapitre 4 :

Application au pendule inversé

Chapitre4 :

Application au pendule inversé

4.1 Introduction

En prenant référence de l'observateur intervalle développé dans le troisième chapitre (conçu pour la classe des systèmes LTI non coopératifs exponentiellement stable avec des perturbations additives), nous allons construire un observateur intervalle au modèle d'état d'un pendule inversé.

Le système présente des perturbations sur les sorties mesurées qui sont mal connues, et les conditions initiales du vecteur d'état sont inconnues mais qui peuvent être bornées par des limites supérieures et inférieures connues à priori.

A cet effet, la méthode d'estimation que nous allons adopter, utilise simultanément un observateur intervalle [Mazenc et Bernard, 2011] et un observateur classique [Luenberger, 1966] pour l'estimation d'état du système incertain.

La technique peut être expliquée de la manière suivante:

En premier lieu, nous allons construire un observateur classique de type Luenberger [Luenberger, 1966] au système perturbé sans tenir compte des incertitudes. En se basant sur l'erreur dynamique associée à cette estimation d'état classique, nous allons construire un observateur intervalle pour l'erreur dynamique.

Cette erreur dynamique vérifie la condition de synthèse de l'observateur intervalle [Mazenc et Bernard, 2011] qui est la stabilité de la matrice d'état de l'erreur notée A_e . Cette condition permet de concevoir un observateur intervalle pour l'erreur dynamique dont l'expression comprend un terme relié aux perturbations additives qui sont incertaines.

Autrement dit, nous allons concevoir un observateur intervalle temps-variant exponentiellement stable pour l'erreur dynamique, qui est un système linéaire temps invariant exponentiellement stable présentant un terme relié aux perturbations additives au système dans son expression.

Par la suite, la connaissance des limites supérieure et inférieure (obtenues par l'observateur intervalle de l'erreur) de la distance entre l'état réel et son estimation par l'observateur Luenberger classique, contribue à finaliser l'estimation d'état du système.

Pour la commande du pendule inversé, nous allons développer une commande par retour d'état classique à la quelle nous allons introduire cet aspect d'incertitudes influant sur le système par l'approche intervalle. Ceci est réalisé par l'utilisation simultanée « observateur classique-observateur intervalle » pour reconstruire et corriger l'état du système incertain.

Ce chapitre est organisé de la manière suivante :

Dans la section (4.2), une modélisation du pendule inversé permet d'avoir un modèle d'état indispensable pour l'estimation d'état et la commande. Dans la section (4.3), nous allons construire un observateur « Luenberger classique » pour le pendule inversé en absence des perturbations sur le système.

Dans la section (4.4), un observateur intervalle pour le système est conçu en trois étapes :

- La première étape est consacrée à la conception d'un observateur « Luenberger classique » pour le système en tenant compte des perturbations en sortie et en absence des incertitudes sur le système (section (4.4.1)).
- Dans la deuxième étape (section (4.4.2)), en prenant en considération les incertitudes sur le système perturbé, un observateur intervalle est conçu pour l'erreur dynamique associée à l'estimation par l'observateur Luenberger classique.
- Dans la troisième étape, à partir de la différence entre les bornes supérieure et inférieure de l'observateur intervalle de l'erreur et l'état estimé par l'observateur classique dans la section (4.4.1), nous allons obtenir un observateur intervalle du système qui est illustré dans la section (4.4.3).

Dans la section (4.5), nous allons réaliser une commande par retour d'état du pendule inversé, basée sur une reconstruction d'état par l'observateur Luenberger et l'observateur intervalle dont les étapes sont citées dans la section (4.5.1).

Les programmes de calcul des différents observateurs sont implantés sous l'environnement MATLAB.

4.2 Modélisation du pendule inversé

Dans cette section, une description du pendule inversé nous permet d'établir un modèle d'état pour celui-ci, la modélisation se base sur le Lagrangien appliqué au système.

4.2.1 Equations du mouvement du système

Afin de modéliser le pendule inversé, nous allons utiliser le Lagrangien $\mathcal{L}_g$ qui est défini comme la différence entre l'énergie cinétique E_c et l'énergie potentielle E_p du système.

$\mathcal{L}_g$ est donné par l'équation suivante :

$$\mathcal{L}_g = E_c - E_p \quad (4.1)$$

L'ensemble chariot-pendule peut être représenté par le schéma ci-dessous :

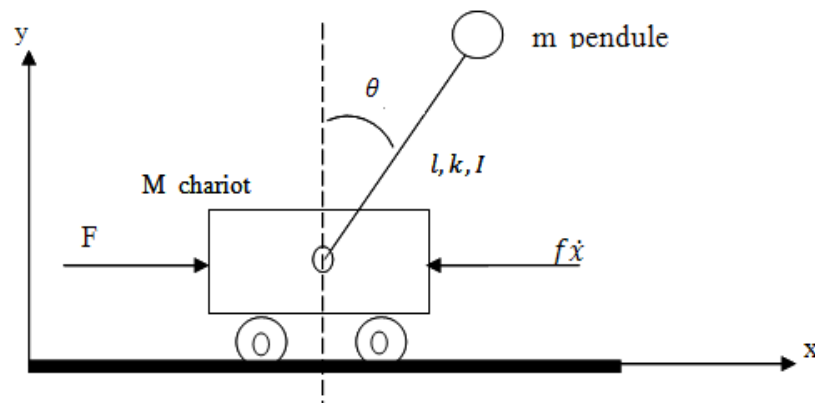


Figure 4.1: Schéma de l'ensemble chariot-pendule

Sur le schéma précédent, apparaissent les différentes grandeurs physiques définies comme suit :

$x(t)$: Position et direction du chariot (en m)

$\theta(t)$: Position angulaire du pendule (en rad)

$F(t)$: Force d'entraînement transmise par la courroie et due au moteur (en N)

M : Masse du chariot (en kg)

m : Masse du pendule (en kg)

g : Accélération de pesanteur = $9,81 \text{ m s}^{-2}$

l : Longueur du pendule (en m)

I : Moment d'inertie du pendule par rapport au centre de masse (en $N \text{ m}^2$)

f : Coefficient de frottement visqueux entre le chariot et le rail et s'opposant à l'action $F(t)$ (en $N/m/s$)

k : Coefficient de frottement visqueux du pendule sur son axe (en $N/rad/s$)

Par la suite, nous adapterons les notations suivantes :

$\frac{d\theta}{dt} = \dot{\theta}$ est la vitesse angulaire du centre de masse du pendule.

$\frac{d^2\theta}{dt^2} = \ddot{\theta}$ est l'accélération angulaire du centre de masse du pendule.

$\frac{dx}{dt} = \dot{x}$ est la vitesse du centre de masse du chariot.

$\frac{d^2x}{dt^2} = \ddot{x}$ est l'accélération du centre de masse du chariot.

Pour une position x du chariot, la position et la vitesse du centre de masse du pendule en projetant sur les axes x et y sont données par le vecteur position $\vec{P}$ et le vecteur vitesse $\vec{V}$ suivants:

$$\vec{P} = \begin{pmatrix} x - l \sin(\theta) \\ -l \cos(\theta) \end{pmatrix} \text{ et } \vec{V} = \begin{pmatrix} \dot{x} - \dot{\theta} l \cos(\theta) \\ l \dot{\theta} \sin(\theta) \end{pmatrix} \quad (4.2)$$

L'énergie cinétique du chariot E_{CM} et celle du pendule E_{cm} sont données par les deux équations suivantes respectivement:

$$E_{CM} = \frac{1}{2} M \dot{x}^2 \quad (4.3)$$

$$\begin{aligned} E_{cm} &= \frac{1}{2} m \dot{V}^2 + \frac{1}{2} I \dot{\theta}^2 \\ &= \frac{1}{2} m \left(\dot{x}^2 - 2 \dot{x} \dot{\theta} l \cos(\theta) + \dot{\theta}^2 l^2 \cos^2(\theta) + \dot{\theta}^2 l^2 \sin^2(\theta) \right) + \frac{1}{2} I \dot{\theta}^2 \end{aligned} \quad (4.4)$$

L'énergie cinétique totale de l'ensemble pendule-chariot a pour expression :

$$\begin{aligned} E_C &= E_{CM} + E_{cm} \\ &= \frac{1}{2} M \dot{x}^2 + \frac{1}{2} m \left(\dot{x}^2 - 2 \dot{x} \dot{\theta} l \cos(\theta) + \dot{\theta}^2 l^2 \cos^2(\theta) + \dot{\theta}^2 l^2 \sin^2(\theta) \right) + \frac{1}{2} I \dot{\theta}^2 \end{aligned} \quad (4.5)$$

Lorsque le chariot est en mouvement, l'énergie potentielle du système, fonction de l'angle du pendule est donnée par la relation suivante :

$$E_p = -mgl(1 - \cos(\theta)) \quad (4.6)$$

En remplaçant (4.5) et (4.6) dans (4.1), nous obtenons le Lagrangien $\mathcal{L}_g$ du système qui est donné par l'expression suivante:

$$\begin{aligned} \mathcal{L}_g &= \frac{1}{2} M \dot{x}^2 + \frac{1}{2} m \left(\dot{x}^2 - 2 \dot{x} \dot{\theta} l \cos(\theta) + \dot{\theta}^2 l^2 \cos^2(\theta) + \dot{\theta}^2 l^2 \sin^2(\theta) \right) + \frac{1}{2} I \dot{\theta}^2 \\ &\quad + mgl(1 - \cos(\theta)) \end{aligned} \quad (4.7)$$

Les équations non-linéaires du mouvement vont être déterminées à partir de l'équation de Lagrange suivante :

$$\frac{d}{dt} \left(\frac{\partial \mathcal{L}_g}{\partial \dot{\xi}} \right) - \left(\frac{\partial \mathcal{L}_g}{\partial \xi} \right) = Q_\xi \quad (4.8)$$

où $\xi = (x, \theta)$ est l'ensemble des cordonnées généralisées du système à savoir le déplacement horizontal du chariot x et l'angle du pendule θ .

et $Q_\xi = (Q_x, Q_\theta)$ représente l'ensemble des forces non conservatrices.

Q_x : correspond à la force d'entraînement F du chariot qui s'oppose à la force du frottement visqueux dont le coefficient est noté f qui est donné par :

$$Q_x = F - f \dot{x} \quad (4.9)$$

Q_θ : correspond à la force de frottement visqueux autour de l'axe dont le coefficient est noté k , donné par :

$$Q_\theta = -k\dot{\theta} \quad (4.10)$$

Alors, l'équation de Lagrange par rapport au chariot est :

$$\frac{d}{dt} \left(\frac{\partial \mathcal{L}_g}{\partial \dot{x}} \right) - \left(\frac{\partial \mathcal{L}_g}{\partial x} \right) = Q_x \quad (4.11)$$

En remplaçant (4.7) dans (4.11) nous obtenons :

$$\begin{cases} \left(\frac{\partial \mathcal{L}_g}{\partial \dot{x}} \right) = (M + m)\dot{x} - ml \dot{\theta} \cos(\theta) \\ \frac{d}{dt} \left(\frac{\partial \mathcal{L}_g}{\partial \dot{x}} \right) = (M + m)\ddot{x} - ml \ddot{\theta} \cos(\theta) + ml \dot{\theta}^2 \sin(\theta) \\ \left(\frac{\partial \mathcal{L}_g}{\partial x} \right) = 0 \end{cases}$$

Nous trouvons :

$$(M + m)\ddot{x} - ml \ddot{\theta} \cos(\theta) + ml \dot{\theta}^2 \sin(\theta) = F - f \dot{x} \quad (4.12)$$

L'équation de Lagrange par rapport au pendule s'écrit :

$$\frac{d}{dt} \left(\frac{\partial \mathcal{L}_g}{\partial \dot{\theta}} \right) - \left(\frac{\partial \mathcal{L}_g}{\partial \theta} \right) = Q_\theta \quad (4.13)$$

En remplaçant (4.7) dans (4.13) nous obtenons :

$$\begin{cases} \left(\frac{\partial \mathcal{L}_g}{\partial \dot{\theta}} \right) = -ml \dot{x} \cos(\theta) + ml^2 \dot{\theta} + I\dot{\theta} \\ \frac{d}{dt} \left(\frac{\partial \mathcal{L}_g}{\partial \dot{\theta}} \right) = -ml \ddot{x} \cos(\theta) + ml \dot{x} \sin(\theta) + (ml^2 + I)\ddot{\theta} \\ \left(\frac{\partial \mathcal{L}_g}{\partial \theta} \right) = ml \dot{x} \sin(\theta) - mgl \sin(\theta) \end{cases}$$

Et nous trouvons comme équation du mouvement :

$$(ml^2 + I)\ddot{\theta} - ml\ddot{x} \cos(\theta) + mgl \sin(\theta) = -k\dot{\theta} \quad (4.14)$$

Les équations du mouvement du système sont alors données par les équations suivantes:

$$\begin{cases} (M + m)\ddot{x} + f\dot{x} - ml\ddot{\theta}\cos(\theta) + ml\dot{\theta}^2 \sin(\theta) = F \\ -ml\ddot{x} \cos(\theta) + (ml^2 + I)\ddot{\theta} + k\dot{\theta} + mgl \sin(\theta) = 0 \end{cases} \quad (4.15)$$

4.2.2 Linéarisation des équations

A partir des équations du mouvement non-linéaires du système (4.15), nous allons effectuer une linéarisation autour des deux positions d'équilibre, l'une pour laquelle $\theta_e = 0$ (point d'équilibre stable) et l'autre pour laquelle $\theta_e = \pi$ (point d'équilibre instable).

En considérant des petites variations $\Delta\theta$ du pendule autour du point d'équilibre θ_e , la position peut être écrite par la relation suivante :

$$\theta = \theta_e + \Delta\theta \quad (4.16)$$

Le développement de Taylor d'ordre 1 d'une fonction non linéaire $f(\theta)$ autour de θ_e est :

$$f(\theta) \approx f(\theta_e) + \Delta\theta \left. \frac{df}{d\theta} \right|_{\theta=\theta_e} \quad (4.17)$$

- **Pendule en position $\theta_e = 0$**

Pour $\theta_e = 0$, nous avons la linéarisation suivante :

$$\begin{aligned} \cos(\theta) &\approx \cos(0) + (\theta - 0)(-\sin(0)) = 1 \\ \sin(\theta) &\approx \sin(0) + (\theta - 0)(\cos(0)) = \theta \end{aligned}$$

Les équations du mouvement (4.15) deviennent alors :

$$\begin{cases} (M + m)\ddot{x} + f\dot{x} - ml\ddot{\theta} = F \\ -ml\ddot{x} + (ml^2 + I)\ddot{\theta} + k\dot{\theta} + mgl\theta = 0 \end{cases} \quad (4.18)$$

- **Pendule en position $\theta_e = \pi$**

Pour $\theta_e = \pi$, nous avons la linéarisation suivante :

$$\begin{aligned} \cos(\theta) &\approx \cos(\pi) + (\theta - \pi)(-\sin(\pi)) = -1 \\ \sin(\theta) &\approx \sin(\pi) + (\theta - \pi)(\cos(\pi)) = (\pi - \theta) \end{aligned}$$

Afin de simplifier, nous allons effectuer le changement de variable sur la position θ comme suit:

$$\theta' = \theta - \pi$$

Alors, nous aurons :

$$\sin(\theta) = \sin(\pi + \theta') = -\theta'$$

Notons aussi que

$$\dot{\theta} = \dot{\theta}' \text{ et } \ddot{\theta} = \ddot{\theta}'$$

Alors les équations du mouvement (4.15) pour $\theta_e = \pi$ deviennent:

$$\begin{cases} (M + m)\ddot{x} + f\dot{x} + ml\ddot{\theta} = F \\ ml\ddot{x} + (ml^2 + I)\ddot{\theta} + k\dot{\theta} - mgl\theta = 0 \end{cases} \quad (4.19)$$

4.2.3 Représentation dans l'espace d'état

Afin de synthétiser un observateur classique ou intervalle au système, une représentation d'état des équations décrivant le mouvement du système est indispensable.

Pour cela, nous définissons le vecteur d'état $x(t)$ de la manière suivante :

$$x(t) = \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} = \begin{pmatrix} x(t) \\ \dot{x}(t) \\ \theta(t) \\ \dot{\theta}(t) \end{pmatrix} \quad (4.20)$$

- **Position stable**

En se référant aux équations relatives au pendule en position stable (4.18) et en négligeant le moment d'inertie I , nous obtenons les équations du mouvement suivantes :

$$\begin{cases} (M + m)\dot{x}_2 + f x_2 - ml \dot{x}_4 = F \\ -ml \dot{x}_2 + ml^2 \dot{x}_4 + kx_4 + mglx_3 = 0 \end{cases} \quad (4.21)$$

Posons :

$$\dot{x}_1 = x_2 = \dot{x}(t) \quad (4.22)$$

$$\dot{x}_3 = x_4 = \dot{\theta}(t) \quad (4.23)$$

Par le changement de variables (4.22) et (4.23) et par substitution de ces dernières dans (4.21) et par élimination successive des termes $\dot{x}_2$ et $\dot{x}_4$, nous obtenons les équations suivantes :

$$\dot{x}_2 = -\frac{f}{M} x_2 - \frac{mg}{M} x_3 - \frac{k}{ml} x_4 + \frac{1}{M} F \quad (4.24)$$

$$\dot{x}_4 = -\frac{f}{Ml} x_2 - \frac{(M+m)g}{Ml} x_3 - \frac{(M+m)k}{Mml^2} x_4 + \frac{1}{Ml} F \quad (4.25)$$

La représentation d'état est alors décrite comme suit :

$$\dot{x}(t) = \begin{pmatrix} \dot{x}_1(t) \\ \dot{x}_2(t) \\ \dot{x}_3(t) \\ \dot{x}_4(t) \end{pmatrix} = \underbrace{\begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & -\frac{f}{M} & -\frac{mg}{M} & -\frac{k}{ml} \\ 0 & 0 & 0 & 1 \\ 0 & -\frac{f}{Ml} & -\frac{(M+m)g}{Ml} & -\frac{(M+m)k}{Mml^2} \end{pmatrix}}_A \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} + \underbrace{\begin{pmatrix} 0 \\ \frac{1}{M} \\ 0 \\ \frac{1}{Ml} \end{pmatrix}}_B F \quad (4.26)$$

- **Position instable**

En suivant le même raisonnement pour le pendule en position instable, nous obtenons la représentation d'état suivante :

$$\begin{pmatrix} \dot{x}_1(t) \\ \dot{x}_2(t) \\ \dot{x}_3(t) \\ \dot{x}_4(t) \end{pmatrix} = \underbrace{\begin{pmatrix} 0 & 1 & 0 & 0 \\ 0 & -\frac{f}{M} & -\frac{mg}{M} & \frac{k}{ml} \\ 0 & 0 & 0 & 1 \\ 0 & \frac{f}{Ml} & \frac{(M+m)g}{Ml} & -\frac{(M+m)k}{Mml^2} \end{pmatrix}}_A \begin{pmatrix} x_1 \\ x_2 \\ x_3 \\ x_4 \end{pmatrix} + \underbrace{\begin{pmatrix} 0 \\ \frac{1}{M} \\ 0 \\ -\frac{1}{Ml} \end{pmatrix}}_B F \quad (4.27)$$

4.3 Synthèse d'un observateur Luenberger classique pour le pendule inversé en absence des perturbations en sortie

Le modèle d'état pris pour l'étude du système est sa position d'équilibre stable (4.26).

- **Paramètres du modèle**

Les différentes constantes du système pendule-chariot sont prises de [Moussouami et Aupoix, 2007], elles sont indiquées ci- dessous :

- La masse M du chariot est calculée comme suit : C'est la masse du chariot qui est de 1.6 Kg à la quelle est rajoutée la masse inertielle du moteur (J/R^2), ce qui fait une masse totale M de 6 Kg .
- La masse totale du pendule est de 0.31 Kg .
- La longueur du pendule est de 0.5 m .
- La grandeur de commande du système est une tension et non pas une force, la tension est reliée à la force par un coefficient de force ks qui est déterminé expérimentalement, il est de 225 N/V . Alors, la tension à appliquée est égale à $u = F / ks$.

Nous avons développé sous SIMULINK le modèle et l'observateur du pendule inversé. La grandeur d'entrée est une impulsion d'amplitude qui vaut ± 1 volts, pendant une durée de 0.1s. Par la suite l'impulsion u est multipliée par le coefficient ks afin d'obtenir la force F qui est l'entrée du système (4.26).

- Le frottement f représentant la valeur du frottement du chariot sur le rail et l'inertie du moteur à courant continu est choisi très grand afin d'avoir une dynamique très rapide $f = 210 \text{ N/m/s}$.
- Le frottement du pendule est pris très petit $k = 0.003 \text{ N/rad/s}$.

• **Conditions posées pour la synthèse de l'observateur Luenberger classique**

Avant d'entreprendre la synthèse de l'observateur, nous mettons comme cahier de charge les conditions suivantes :

- 1) Supposons que seulement la position x et l'angle θ sont disponibles à la mesure, alors la matrice d'observabilité est définie comme suit:

$$C = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 \end{pmatrix} \quad (4.28)$$

de telle manière à ce que nous aurons la paire (A, C) observable.

- 2) Prenons comme conditions initiales des variables du vecteur d'état du système réel et celles de l'observateur classique respectivement comme suit :

$$x_0(t_0) = (0, 0, 0, 0)^T \quad (4.29)$$

$$\hat{x}(t_0) = (0.01, 0.1, 0.025, 0.1)^T \quad (4.30)$$

- 3) Le système possède deux pôles réels dont un est à l'origine et deux pôles complexes donnés ci-dessous:

$$\begin{aligned} p1 &= 0 \\ p2 &= -34.9734 \\ p3 &= -0.0336 + 4.431 i \\ p4 &= -0.0336 - 4.431 i \end{aligned} \quad (4.31)$$

$p1$ et $p2$ sont les pôles relatifs au chariot, $p3$ et $p4$ sont les pôles relatifs au pendule. Les deux derniers pôles complexes conjugués traduisent l'aspect oscillatoire du pendule.

Nous choisissons le gain L de l'observateur qui est calculé par placement de pôles de telle manière à avoir la matrice $(A - LC)$ Hurwitz (C'est-à-dire les valeurs propres de $(A - LC)$ sont à parties réelles strictement négatives).

Nous choisissons de changer le pôle à l'origine et le mettre à -1, nous gardons $p2$ tel quel et afin d'amortir cet aspect oscillatoire des deux pôles complexes, nous effectuons une réduction de la pulsation γ_i des deux pôles conjugués, nous prenons la valeur $\gamma_i = 0.5$. Les pôles choisis sont alors les suivants :

$$\begin{aligned} p1 &= -1 \\ p2 &= -34.9734 \\ p3 &= -1.5 + 0.5 i \\ p4 &= -1.5 - 0.5 i \end{aligned} \quad (4.32)$$

L'observateur Luenberger classique est donné par l'équation suivante :

$$\hat{\dot{x}} = A\hat{x} + Bu + L(y - \hat{y}) \quad (4.33)$$

et qui peut être réécrit de la manière suivante :

$$\hat{\dot{x}} = A_e \hat{x} + [B \quad L] \begin{bmatrix} u \\ y \end{bmatrix} \quad (4.34)$$

avec $A_e = A - LC$

Nous avons les dimensions suivantes pour le système (4.34) : $A_e \in R^{4 \times 4}$, $B \in R^{4 \times 1}$ et $L \in R^{4 \times 2}$.

Comme nous pouvons aussi avoir l'écriture suivante pour (4.33) :

$$\hat{\dot{x}} = A \hat{x} + [B \quad L] \begin{bmatrix} u \\ (y - \hat{y}) \end{bmatrix} \quad (4.35)$$

L est obtenu par placement de pôles, en utilisant la commande « place » sur MATLAB. Nous obtenons la matrice suivante pour les pôles désirés (4.32):

$$L = \begin{pmatrix} 1.1084 & 0.0877 \\ 1.0004 & -0.5247 \\ -0.3435 & 2.8243 \\ 2.2924 & -18.5602 \end{pmatrix} \quad (4.36)$$

L'observateur Luenberger (4.35) peut être représenté par le schéma de simulation suivant (sous SIMULINK) :

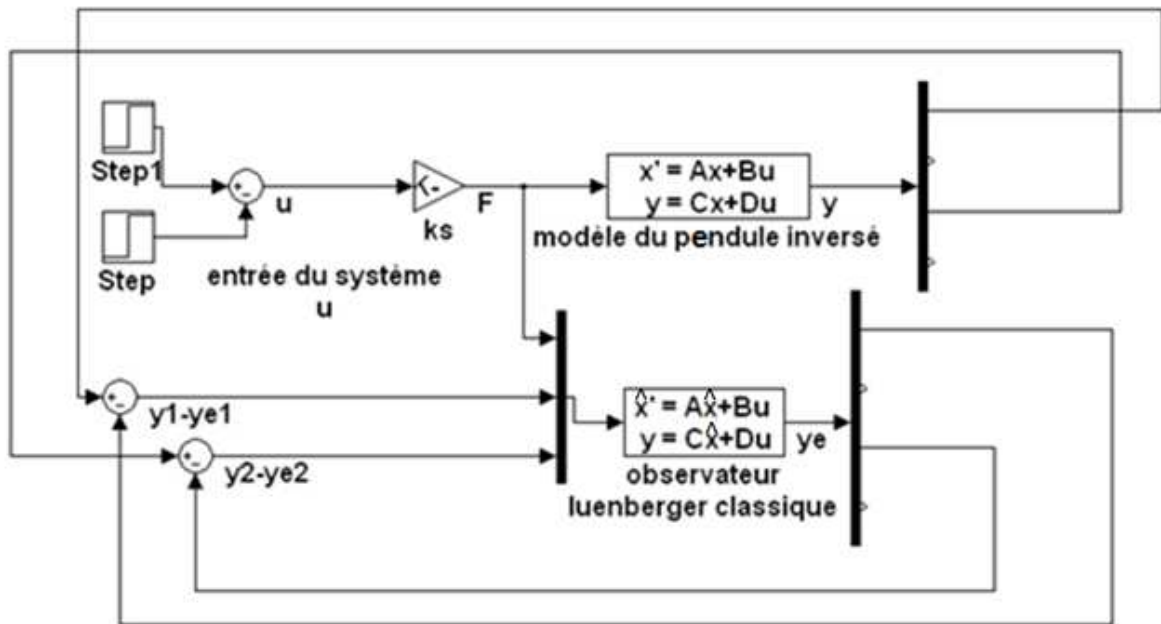


Figure 4.2 : Schéma de simulation sous SIMULINK du système réel doté d'un observateur Luenberger classique.

La réponse à une impulsion d'amplitude 1V et d'une durée 0.1s, permet de tracer les états du système estimés par l'observateur Luenberger classiques, présentés par les figures suivantes :

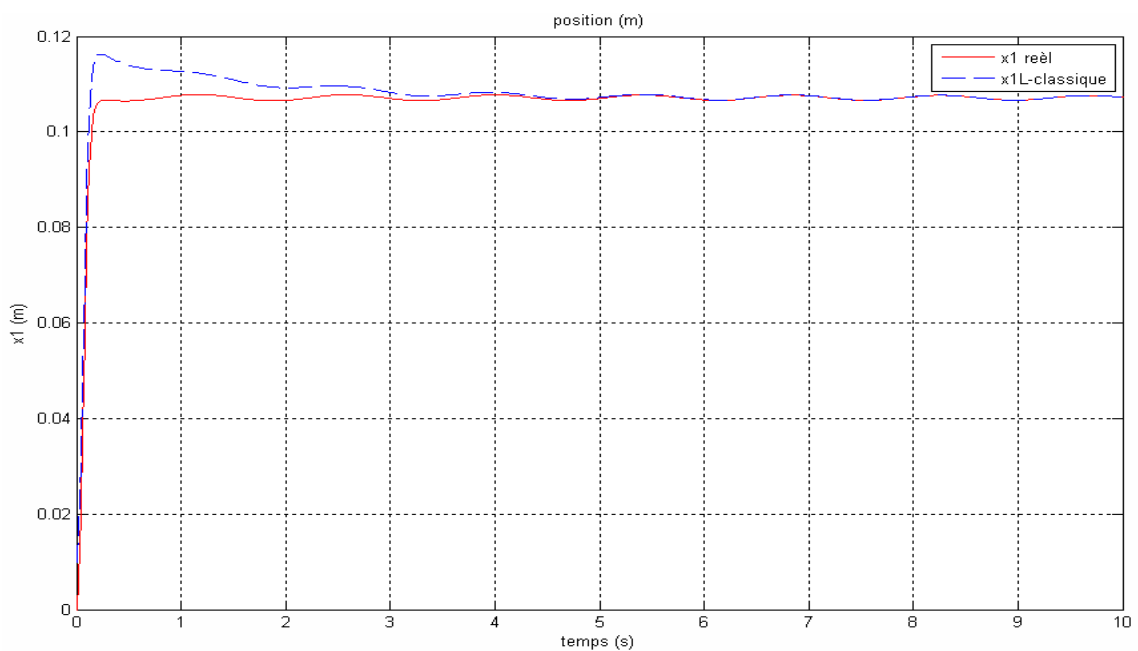


Figure 4.3 : Réponse en position du chariot, état réel (en rouge) et état estimé par l'observateur classique de Luenberger (en bleu).

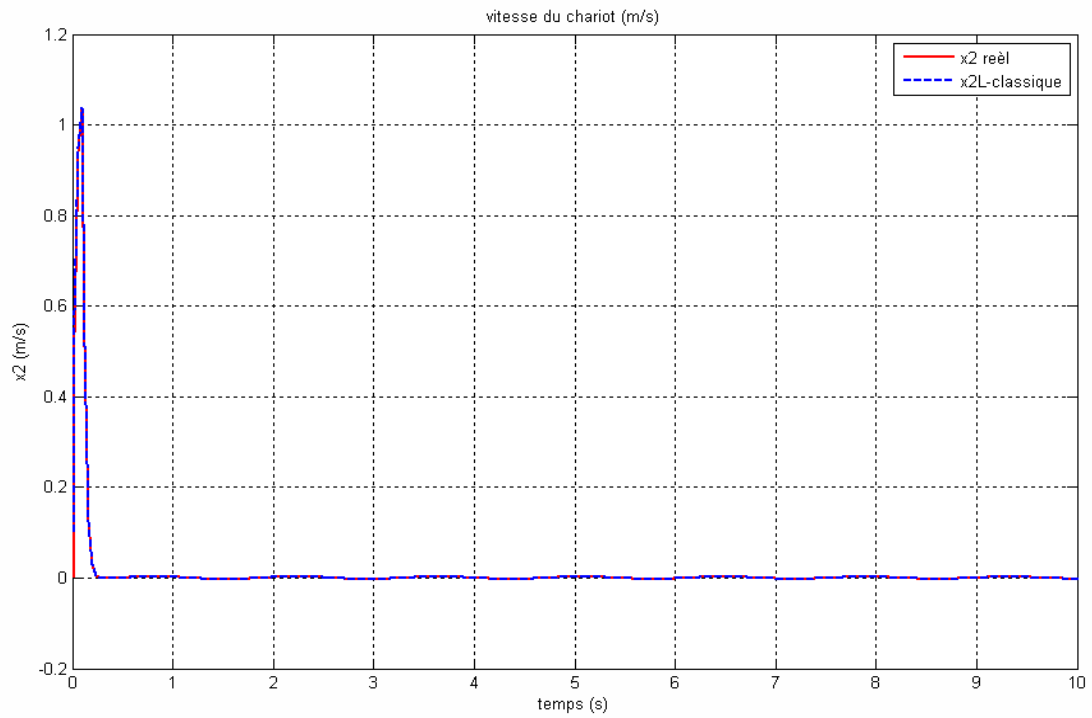


Figure 4.4 : Réponse en vitesse du chariot, état réel (en rouge) et état estimé (en bleu) par l'observateur classique de Luenberger.

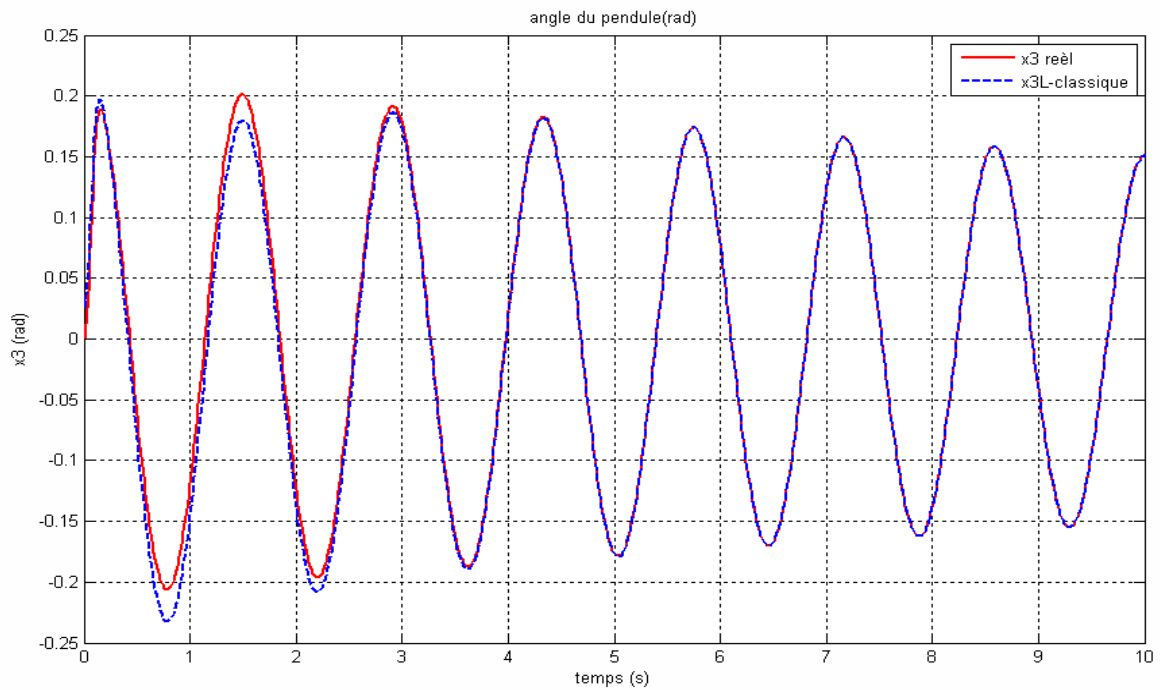


Figure 4.5 : Réponse angulaire du pendule, état réel (en rouge) et état estimé (en bleu) par l'observateur classique de Luenberger.

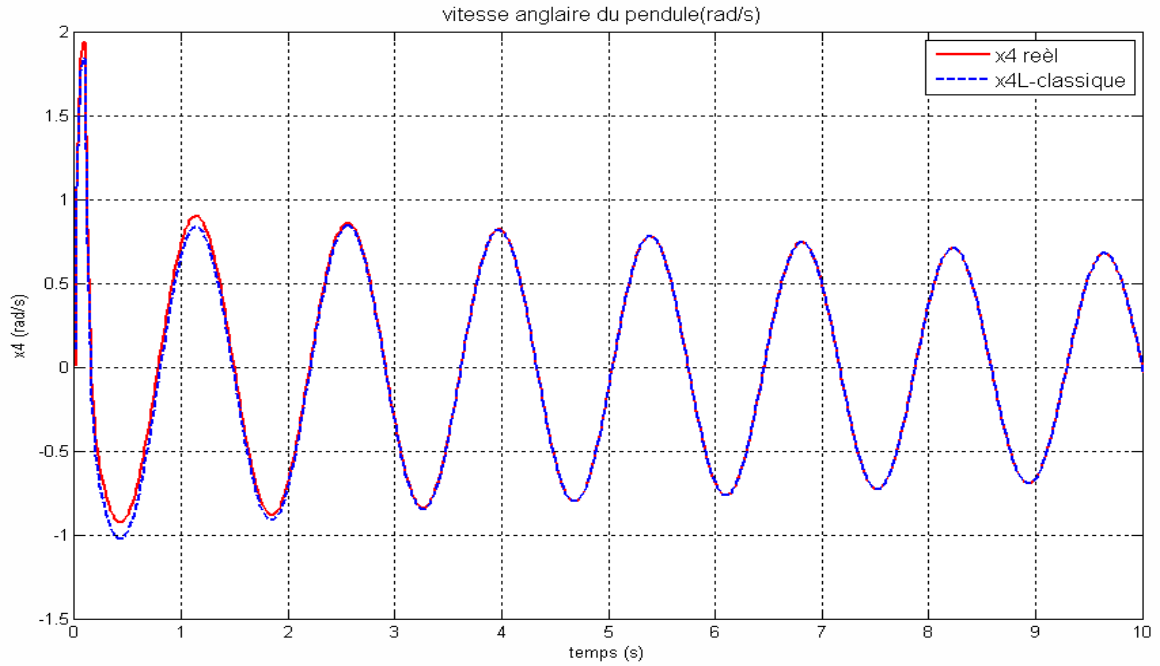


Figure 4.6 : Réponse en vitesse angulaire du pendule, état réel (en rouge) et état estimé (en bleu) par l'observateur classique de Luenberger.

Nous remarquons d'après les réponses du système pour l'entrée considérée (impulsion de 1V, durée 0.1s), que l'observateur classique choisi suit bien la dynamique du système réel. L'évolution de l'état estimé par l'observateur de Luenberger est confondue à celle de l'état réel.

L'erreur dynamique associée est donnée par le système suivant:

$$\dot{e} = A_e e \quad (4.37)$$

avec $e = (x - \hat{x})$ est la distance entre l'état réel x du système et l'état estimé par l'observateur classique de Luenberger $\hat{x}$.

Le vecteur initial de l'erreur dynamique calculé à partir de (4.29) et (4.30) est donné par :

$$e_0 = (-0.01, -0.1, -0.025, -0.1)^T \quad (4.38)$$

Nous traçons ci-dessous l'évolution de l'erreur dynamique associée, illustrée par la figure 4.7 suivante :

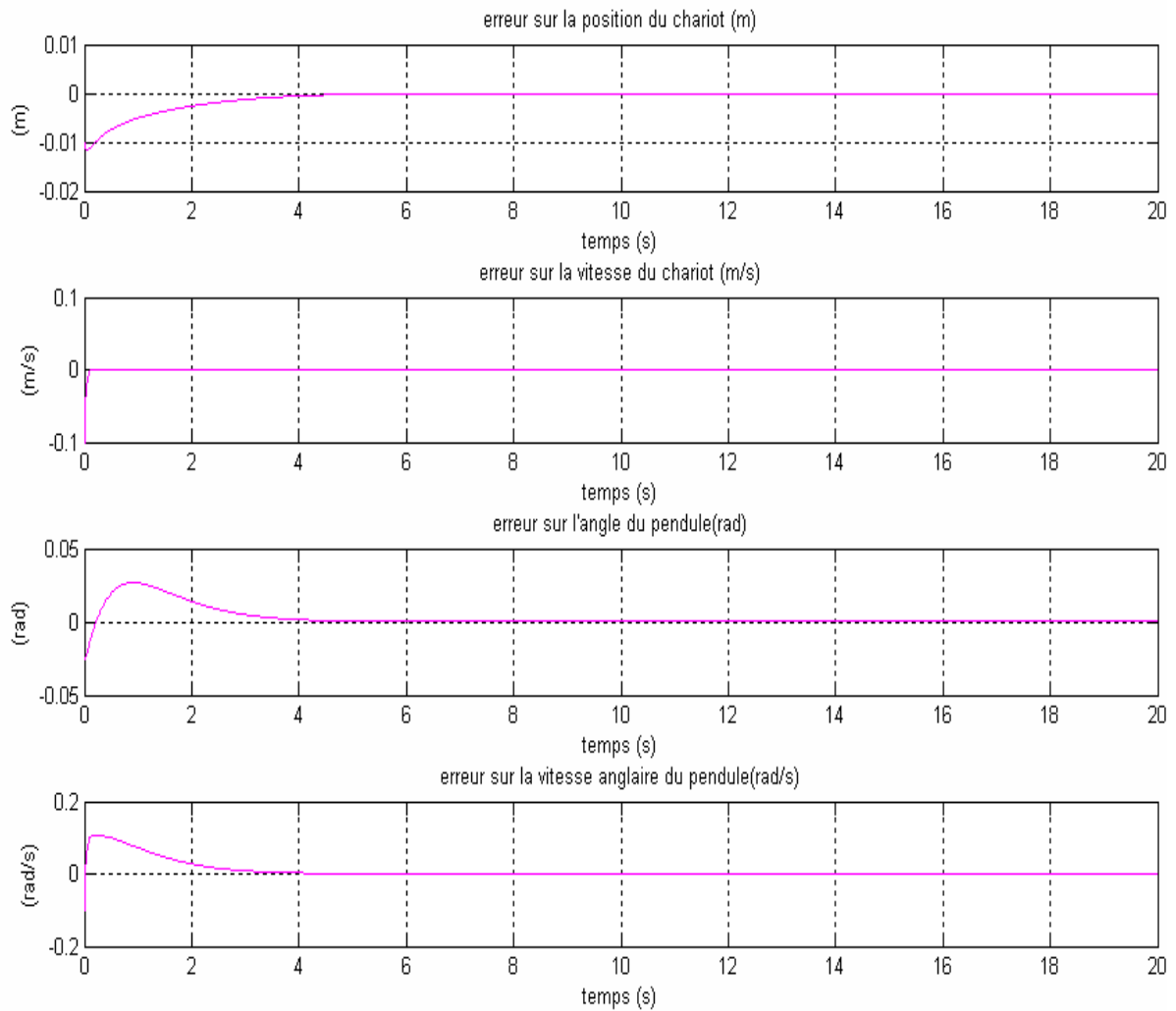


Figure 4.7 : Evolution de l'erreur dynamique pour les quatre variables d'état du pendule inversé.

Nous remarquons d'après le tracé de l'évolution de l'erreur $e(t)$ (figure 4.7), que celle-ci s'annule rapidement. Ce qui fait que l'état estimé par l'observateur de Luenberger $\hat{x}$ converge rapidement vers l'état réel x du système. Donc, le choix du gain L considéré est un bon choix et le gain L nous garantit cette convergence de l'état estimé vers l'état réel.

4.4 Estimation d'état en présence d'incertitudes sur le système

Dans cette partie, nous allons démontrer comment un observateur intervalle et un observateur classique, peuvent être utilisés simultanément pour obtenir des informations complémentaires sur les solutions d'un système LTI incertain observable, présentant des perturbations additives en sortie.

Nous considérons le modèle d'état précédent du pendule inversé en présence d'incertitudes sur la sortie du système (perturbations en sortie). Et en même temps, nous avons les conditions initiales du vecteur d'état qui sont inconnues mais qui peuvent être bornées dans un vecteur intervalle dont les limites sont connues à priori.

Le système est alors décrit par la représentation d'état suivante :

$$\begin{cases} \dot{x} = Ax + Bu \\ y = Cx + \varphi(t) \end{cases} \quad (4.39)$$

Où les variables du vecteur d'état initial x_0 à l'instant $t = 0s$ sont supposées comprises dans un vecteur intervalle de bornes connues à priori. Et $\varphi(t)$ est un bruit de mesure qui est une fonction Lipschitz continue inconnue.

En présence des perturbations en sortie sur le système, l'observateur Luenberger décrit précédemment (4.33), admet l'erreur dynamique suivante :

$$\dot{e} = A_e e - L\varphi(t) \quad (4.40)$$

Donc, afin de synthétiser un observateur intervalle pour le système incertain, nous allons concevoir un observateur intervalle pour l'erreur dynamique (4.40).

Posons :

$$w(t) = -L\varphi(t) \quad (4.41)$$

tel que $w(t)$ soit une fonction Lipschitz continue incertaine, comprise entre deux fonctions Lipschitz continues connues $w^+(t)$ et $w^-(t)$. Et sous la condition du choix du gain L de l'observateur classique, nous avons la matrice ($A_e = A - LC$) qui est stable.

Alors, le théorème 3.5 [Mazenc et Bernard, 2011] peut être appliqué à (4.40) et ceci afin de construire un observateur intervalle pour l'erreur dynamique. Et à partir de cette dernière, nous obtenons un observateur intervalle pour le système incertain.

L'erreur dynamique (4.40) peut être réécrite de la manière suivante :

$$\dot{e} = A_e e + w(t) \quad (4.42)$$

avec

$$\begin{aligned} w^-(t) &\leq w(t) \leq w^+(t) \\ w^+(t) &= w(t) + \Gamma \\ w^-(t) &= w(t) - \Gamma \end{aligned} \quad (4.43)$$

Les conditions initiales de l'erreur dynamique sont aussi incertaines, le vecteur initial de l'erreur dynamique $e_0(t_0 = 0s)$ est défini comme suit :

$$e_0^- \leq e_0 \leq e_0^+ \quad (4.44)$$

Alors, en appliquant le théorème 3.5 à (4.42), nous obtenons l'observateur intervalle suivant :

$$\begin{cases} \dot{z}^+ = N z^+ + P^+(t) w^+(t) - P^-(t) w^-(t) \\ \dot{z}^- = N z^- + P^+(t) w^-(t) - P^-(t) w^+(t) \\ z^+(t_0) = P^+(t_0) e_0^+ - P^-(t_0) e_0^- \\ z^-(t_0) = P^+(t_0) e_0^- - P^-(t_0) e_0^+ \end{cases} \quad (4.45)$$

Avec, les conditions initiales $z^+(t_0)$ et $z^-(t_0)$ sont calculées à partir des bornes inférieure et supérieure de l'état initiale de l'erreur (4.44).

L'observateur intervalle final de l'erreur est alors donné par les équations suivantes:

$$\begin{cases} e^+ = M^+(t) z^+ - M^-(t) z^- \\ e^- = M^+(t) z^- - M^-(t) z^+ \end{cases} \quad (4.46)$$

Où

$P^+(t) = \max(0, P(t))$, $P^-(t) = P^+(t) - P(t)$ et la matrice $M(t)$ et l'inverse de $P(t)$,
 $M^+(t) = \max(0, M(t))$ et $M^-(t) = M^+(t) - M(t)$

Par la suite, nous pouvons écrire :

$$\begin{aligned} e^- &\leq e \leq e^+ \\ e^- &\leq x - \hat{x} \leq e^+ \end{aligned} \quad (4.47)$$

Enfin, l'observateur intervalle du système est déduit à partir de l'inégalité (4.47). Il est donné comme suit:

$$e^- + \hat{x} \leq x \leq e^+ + \hat{x} \quad (4.48)$$

Procédons étape par étape pour synthétiser un observateur intervalle pour le pendule inversé à sa position stable, en présence d'incertitudes sur le système:

4.4.1 Synthèse d'un observateur Luenberger classique en tenant compte des perturbations en sortie et en absence d'incertitudes

Supposons que $\varphi(t)$ est un bruit de mesure uniformément distribué sur l'intervalle $[-0.02, 0.01]$.

$\varphi(t)$ est obtenu par une interpolation linéaire entre deux vecteur temps de périodes d'échantillonnages différentes.

L'observateur Luenberger classique en tenant compte des perturbations en sortie, peut être réécrit de la manière suivante :

$$\hat{\dot{x}} = A\hat{x} + Bu + L((Cx + \varphi(t)) - \hat{y}) \quad (4.49)$$

Prenons comme conditions initiales des variables du vecteur d'état, respectivement pour le système réel et le système doté de l'observateur de Luenberger classique en présence des perturbations en sortie, les vecteurs suivants :

$$\begin{aligned} x_0 &= (0.0062, 0.0062, 0.0062, 0.0062)^T \\ \hat{x}_0 &= (0.0162, 0.1062, 0.0312, 0.1062)^T \end{aligned} \quad (4.50)$$

Nous traçons ci-dessous, l'évolution des états du système réel et des états estimés par l'observateur Luenberger classique en présence des perturbations en sortie :

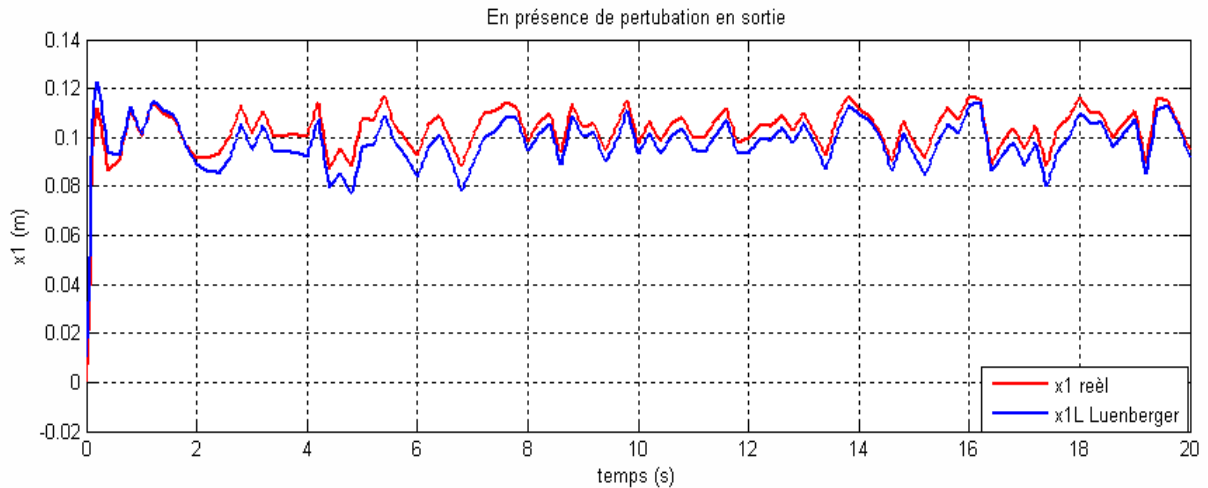


Figure 4.8 : Estimation de la position du chariot par l'observateur Luenberger classique en présence des perturbations sur la mesure.

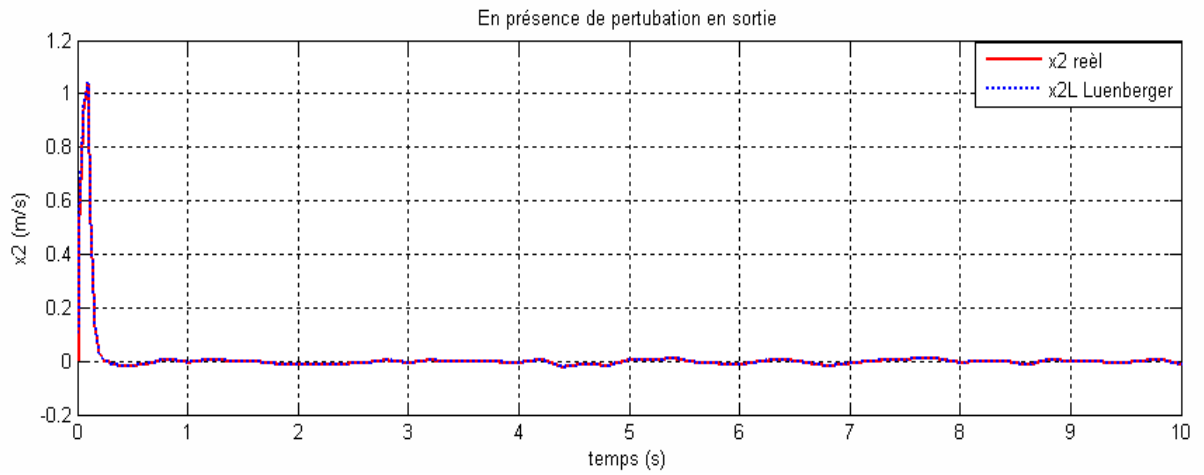


Figure 4.9 : Estimation de la vitesse du chariot par l'observateur Luenberger classique en présence des perturbations sur la mesure.

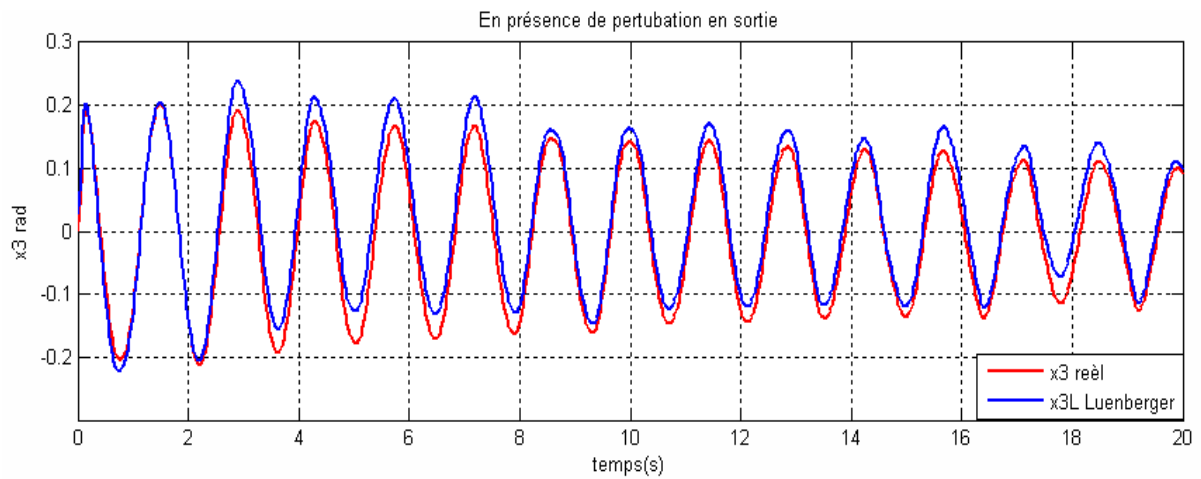


Figure 4.10 : Estimation de l'angle du pendule par l'observateur Luenberger classique en présence des perturbations sur la mesure.

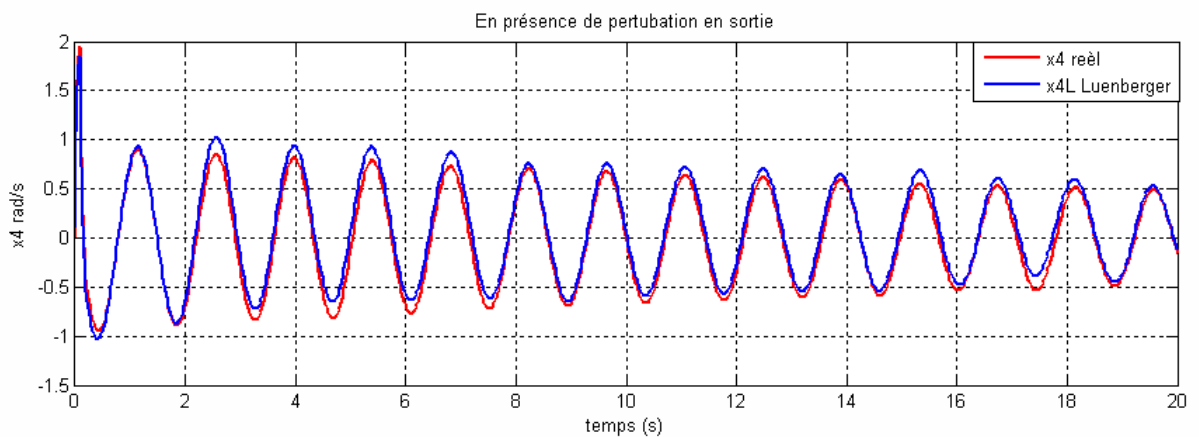


Figure 4.11 : Estimation de la vitesse angulaire du pendule par l'observateur Luenberger classique en présence des perturbations sur la mesure.

En présence des perturbations en sortie du système, nous remarquons d'après les tracés des réponses pour l'entrée considérée (impulsion de 1 V, durée 0.1s), qu'en gardant le même gain L pour l'observateur classique choisi, que les états estimés suivent les dynamiques des états réels à une différence près.

Traçons par la suite, l'erreur dynamique associée :

Les conditions initiales pour l'erreur dynamique (4.42) sont calculées par la différence entre x_0 et $\hat{x}_0$ de (4.50). Nous trouvons le vecteur colonne suivant :

$$e_0 = (-0.01, -0.1, -0.025, -0.1)^T \quad (4.51)$$

L'évolution de l'erreur dynamique des quatre états est donnée par la figure suivante :

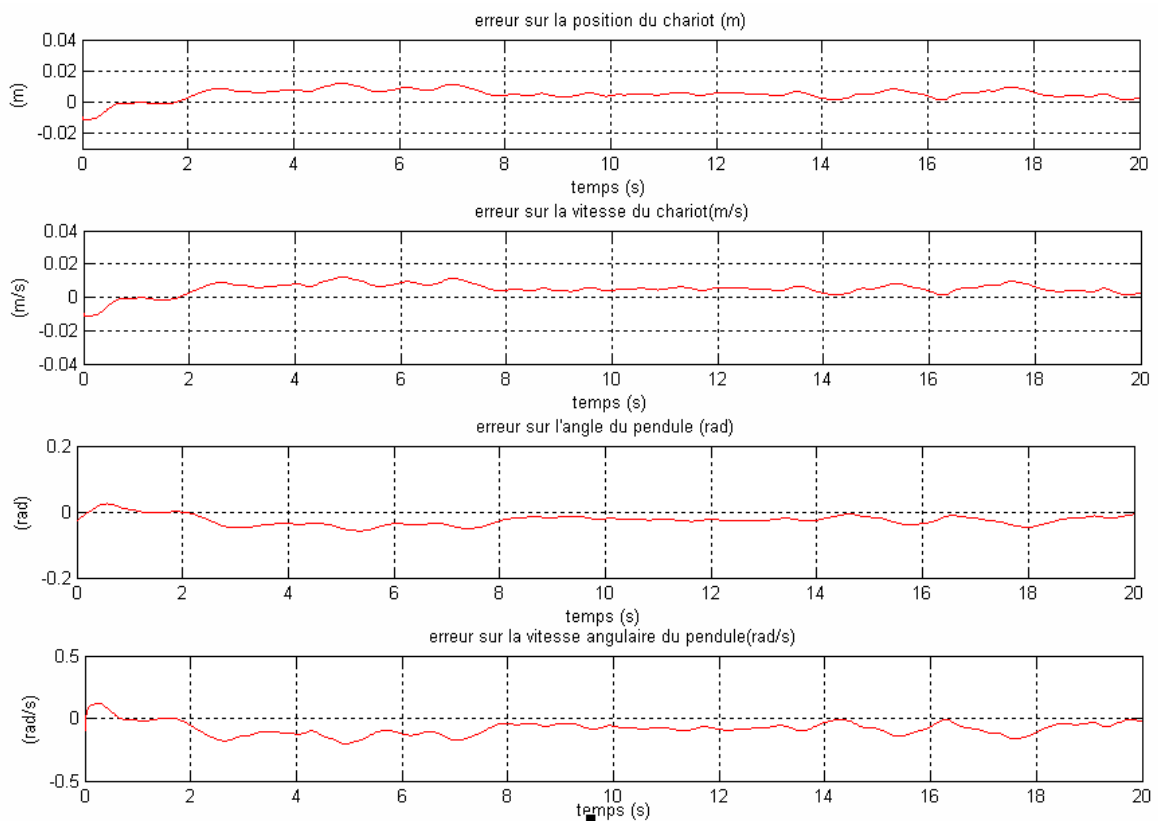


Figure 4.12 : Evolution de l'erreur dynamique du système perturbé.

Nous remarquons d'après les tracés de l'erreur dynamique des quatre états considérés, que celle-ci tend vers zéro. Donc, les états estimés par l'observateur classique de Luenberger convergent vers les états réels du système.

4.4.2 Synthèse d'un observateur intervalle pour l'erreur dynamique du système

Nous considérons l'erreur dynamique donnée par (4.42), telle que les perturbations en sortie sont inconnues mais bornées entre deux fonctions Lipschitz continues connues (4.43). Et les conditions initiales (4.51) de l'erreur dynamique sont considérées incertaines appartenant à un vecteur intervalle de bornes connues à priori.

Nous supposons que le vecteur initial de l'erreur dynamique e_0 à l'instant $t = 0$ est donné par le vecteur intervalle suivant :

$$e_0 = ([-0.06, 0.04] \quad , [-0.15, -0.05] \quad , [-0.075, 0.025] \quad , [-0.15, -0.05])^T \quad (4.52)$$

et qui peut être aussi réécrit par l'inégalité (4.44) que nous réécrivons:

$$e_0^- \leq e_0 \leq e_0^+$$

avec e_0^- et e_0^+ sont donnés par les deux vecteur colonnes suivants respectivement :

$$\begin{aligned} e_0^- &= (-0.06, -0.15, -0.075, -0.15)^T \\ e_0^+ &= (0.04, -0.05, 0.025, -0.05)^T \end{aligned} \quad (4.53)$$

La matrice $A_e = A - LC$ calculée est la suivante :

$$A_e = \begin{pmatrix} -1.1084 & 1 & -0.0877 & 0 \\ -1.0004 & -35 & 0.0178 & -0.001 \\ 0.3435 & 0 & -2.8243 & 1.0000 \\ -2.2924 & -70 & -2.0735 & -0.0407 \end{pmatrix} \quad (4.54)$$

La matrice A_e (qui est stable) est non coopérative. Les valeurs propres de A_e qui sont aussi les pôles désirés (4.32) sont:

$$\lambda_1 = -1, \quad \lambda_2 = -34.9734 \text{ et } \lambda_{3,4} = -1.5 \pm 0.5 i \quad (4.55)$$

Nous suivons les étapes ci-dessous pour l'obtention de l'observateur intervalle pour l'erreur dynamique.

a) Transformation sous forme canonique de Jordan

La transformation de la matrice A_e en forme canonique de Jordan [Perko, 1991] est effectuée de la manière suivante :

Pour la matrice :

$$P = \begin{pmatrix} -2.6592 & -1.4186 & -0.6260 & 0.6510 \\ 0.0661 & 2.2373 & -0.0011 & 0.0001 \\ -3.4666 & 4.2326 & 2.0259 & -2.1070 \\ -0.1320 & 4.4597 & 4.1311 & -2.1059 \end{pmatrix} \quad (4.56)$$

l'égalité

$$P A P^{-1} = J \quad (4.57)$$

avec

$$J = \begin{pmatrix} -1 & 0 & 0 & 0 \\ 0 & -34.9734 & 0 & 0 \\ 0 & 0 & -1.5 & 0.5 \\ 0 & 0 & -0.5 & -1.5 \end{pmatrix} \quad (4.58)$$

est vérifiée, d'après la forme canonique de Jordan J de A_e celle-ci n'est pas coopérative, donc il faudra effectuer un changement de base « temps -variant ».

b) Changement de coordonnées temps-variant

Le changement de coordonnées considéré est « temps-variant » $P(t) = n(t) P$. $n(t)$ calculée par (3.77) est donnée par la matrice suivante :

$$n(t) = \begin{pmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & \cos(\gamma_i t) & -\sin(\gamma_i t) \\ 0 & 0 & \sin(\gamma_i t) & \cos(\gamma_i t) \end{pmatrix} \quad (4.59)$$

avec, $\gamma_i = 0.5$ est la partie imaginaire des valeurs propres complexes (4.55) de A_e .

La matrice $P(t)$ calculée est la suivante :

$$P(t) = \begin{pmatrix} -2.659 & -1.419 & -0.626 & 0.651 \\ 0.06611 & 2.237 & -0.001055 & 9.426 \cdot 10^{-5} \\ 0.132 S(t) - 3.467 C(t) & 4.233 C(t) - 4.46 S(t) & 2.026 C(t) - 4.131 S(t) & 2.106 S(t) - 2.107 C(t) \\ -0.132 C(t) - 3.467 S(t) & 4.46 C(t) + 4.233 S(t) & 4.131 C(t) + 2.026 S(t) & -2.106 C(t) - 2.107 S(t) \end{pmatrix} \quad (4.60)$$

avec $C(t) = \cos(\gamma_i t)$ et $S(t) = \sin(\gamma_i t)$

La matrice $P(t)$ est la solution de $\dot{P}(t) = N P(t) - P(t)A$. La matrice N calculée par (3.80) est donnée comme suit:

$$N = \begin{pmatrix} -1 & 0 & 0 & 0 \\ 0 & -34.9734 & 0 & 0 \\ 0 & 0 & -1.5 & 0 \\ 0 & 0 & 0 & -1.5 \end{pmatrix} \quad (4.61)$$

La matrice $M(t) = P(t)^{-1}$ qui sera utilisée dans l'étape suivante lors du calcul de l'observateur intervalle de l'erreur dynamique est :

$$M = \begin{pmatrix} -0.2683 & , & -0.01328 & , & (-0.0829 C(t) - 9.074 \cdot 10^{-17} S(t)) & , & (9.074 \cdot 10^{-17} C(t) - 0.0829 S(t)) \\ 8.092 \cdot 10^{-3} & , & 0.4473 & , & (2.316 \cdot 10^{-3} C(t) - 2.047 \cdot 10^{-4} S(t)) & , & (2.047 \cdot 10^{-4} C(t) + 2.316 \cdot 10^{-3} S(t)) \\ 0.4237 & , & -0.02766 & , & (-0.3436 C(t) - 0.4748 S(t)) & , & (0.4748 C(t) - 0.3436 S(t)) \\ 0.8651 & , & 0.8938 & , & (-0.664 C(t) - 0.4569 S(t)) & , & (0.4569 C(t) - 0.664 S(t)) \end{pmatrix} \quad (4.62)$$

c) Observateur intervalle pour l'erreur dynamique

L'observateur intervalle temps-variant correspondant est d'ordre huit, qui est donné par le système d'équations différentielles défini précédemment par (4.45).
avec

$$P^+(t) = \begin{pmatrix} 0 & 0 & 0 & 0.651 \\ 0.06611 & 2.237 & 0 & 9.426 \cdot 10^{-5} \\ \mathcal{L}(0.132 S(t) - 3.467 C(t)) & \mathcal{L}(4.233 C(t) - 4.46 S(t)) & \mathcal{L}(2.026 C(t) - 4.131 S(t)) & \mathcal{L}(2.106 S(t) - 2.107 C(t)) \\ \mathcal{L}(-0.132 C(t) - 3.467 S(t)) & \mathcal{L}(4.46 C(t) + 4.233 S(t)) & \mathcal{L}(4.131 C(t) + 2.026 S(t)) & \mathcal{L}(-2.106 C(t) - 2.107 S(t)) \end{pmatrix} \quad (4.63)$$

et $P^-(t) = P^+(t) - P(t)$

En prenant comme conditions initiales à l'instant ($t_0 = 0s$) e_0^+ et e_0^- définies précédemment par (4.53), nous obtenons comme conditions initiales $z^+(t_0)$ et $z^-(t_0)$ du système différentiel (4.45) les vecteurs colonnes suivants :

$$\begin{aligned} z^+(t_0) &= (0.3868, -0.1091, 0.3631, 0.2041)^T \\ z^-(t_0) &= (-0.1487, -0.3396, -0.8202, -0.8788)^T \end{aligned} \quad (4.64)$$

avec, les matrices P^+ et P^- calculées à $(t_0 = 0s)$ sont les suivantes :

$$P^+ = \begin{pmatrix} 0 & 0 & 0 & 0.6510 \\ 0.0661 & 2.2370 & 0 & 0.0001 \\ 0 & 4.2330 & 2.0260 & 0 \\ 0 & 4.4600 & 4.1310 & 0 \end{pmatrix} \quad (4.65)$$

et

$$P^- = P^+ - P = \begin{pmatrix} 2.6590 & 1.4190 & 0.6260 & 0 \\ 0 & 0 & 0.0011 & 0 \\ 3.4670 & 0 & 0 & 2.107 \\ 0.1320 & 0 & 0 & 2.1060 \end{pmatrix} \quad (4.66)$$

$z^+(t)$ et $z^-(t)$ sont obtenues en résolvant le système d'équations différentielles (4.45).

Enfin, l'estimation d'état de l'erreur dynamique est calculée par les équations (4.46) et la matrice $M^+(t)$ calculée à partir de la matrice $M(t)$ donnée par (4.62) est :

$$M^+(t) = \max(0, M(t)) = \begin{pmatrix} 0, 0, \mathcal{L}(-0.0829 C(t) - 9.074 \cdot 10^{-17} S(t)), \mathcal{L}(9.074 \cdot 10^{-17} C(t) - 0.0829 S(t)) \\ 8.092 \cdot 10^{-3}, 0.4473, \mathcal{L}(2.316 \cdot 10^{-3} C(t) - 2.047 \cdot 10^{-4} S(t)), \mathcal{L}(2.047 \cdot 10^{-4} C(t) + 2.316 \cdot 10^{-3} S(t)) \\ 0.4237, 0, \mathcal{L}(-0.3436 C(t) - 0.4748 S(t)), \mathcal{L}(0.4748 C(t) - 0.3436 S(t)) \\ 0.8651, 0.8938, \mathcal{L}(-0.664 C(t) - 0.4569 S(t)), \mathcal{L}(0.4569 C(t) - 0.664 S(t)) \end{pmatrix} \quad (4.67)$$

et

$$M^-(t) = M^+(t) - M(t) \quad (4.68)$$

Nous obtenons l'observateur intervalle suivant pour l'erreur dynamique :

$$\begin{aligned} e^+ &= M^+(t)z^+ - M^-(t)z^- \\ e^- &= M^+(t)z^- - M^-(t)z^+ \end{aligned} \quad (4.69)$$

Les figures suivantes présentent les limites supérieure et inférieure de l'erreur dynamique obtenues par l'observateur intervalles conçu pour l'erreur $e(t)$. On varie le paramètre Γ des bornes supérieure w^+ et inférieure w^- des perturbations $w(t)$:

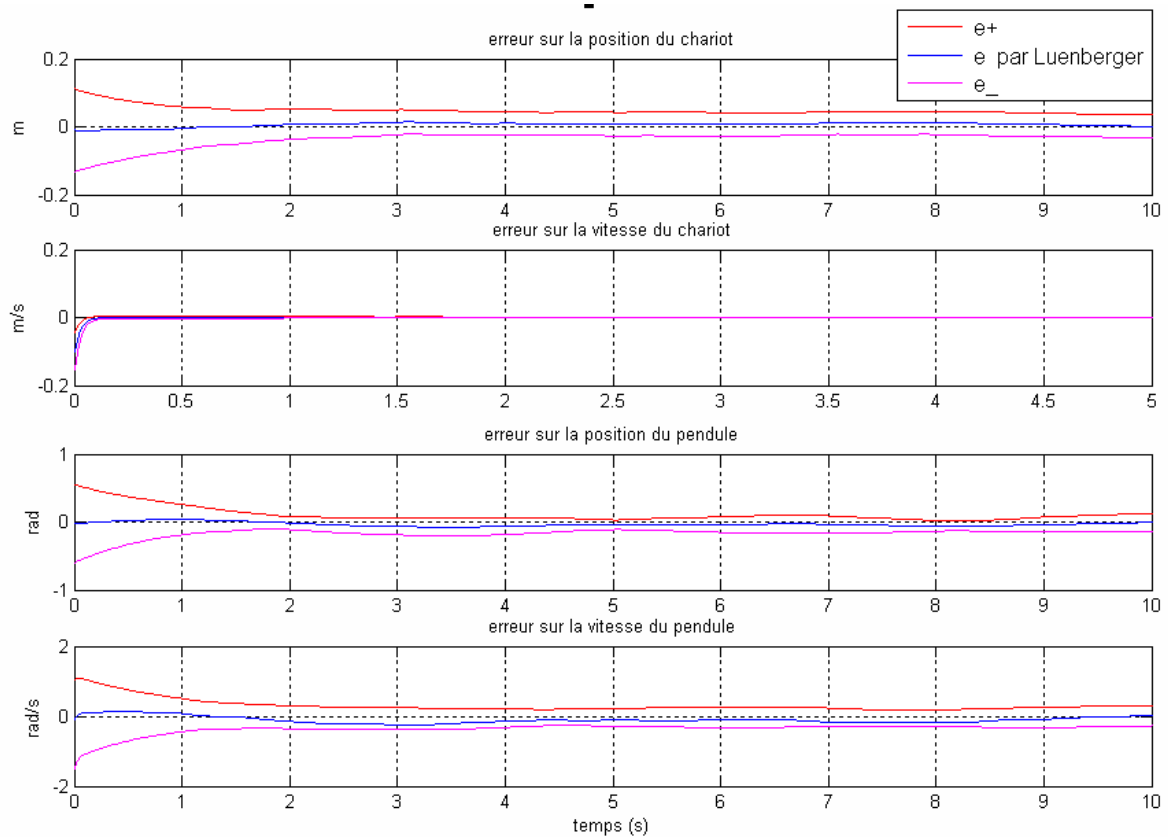


Figure 4.13: Limites supérieure et inférieure de l'erreur dynamique obtenues par un observateur intervalle pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$.

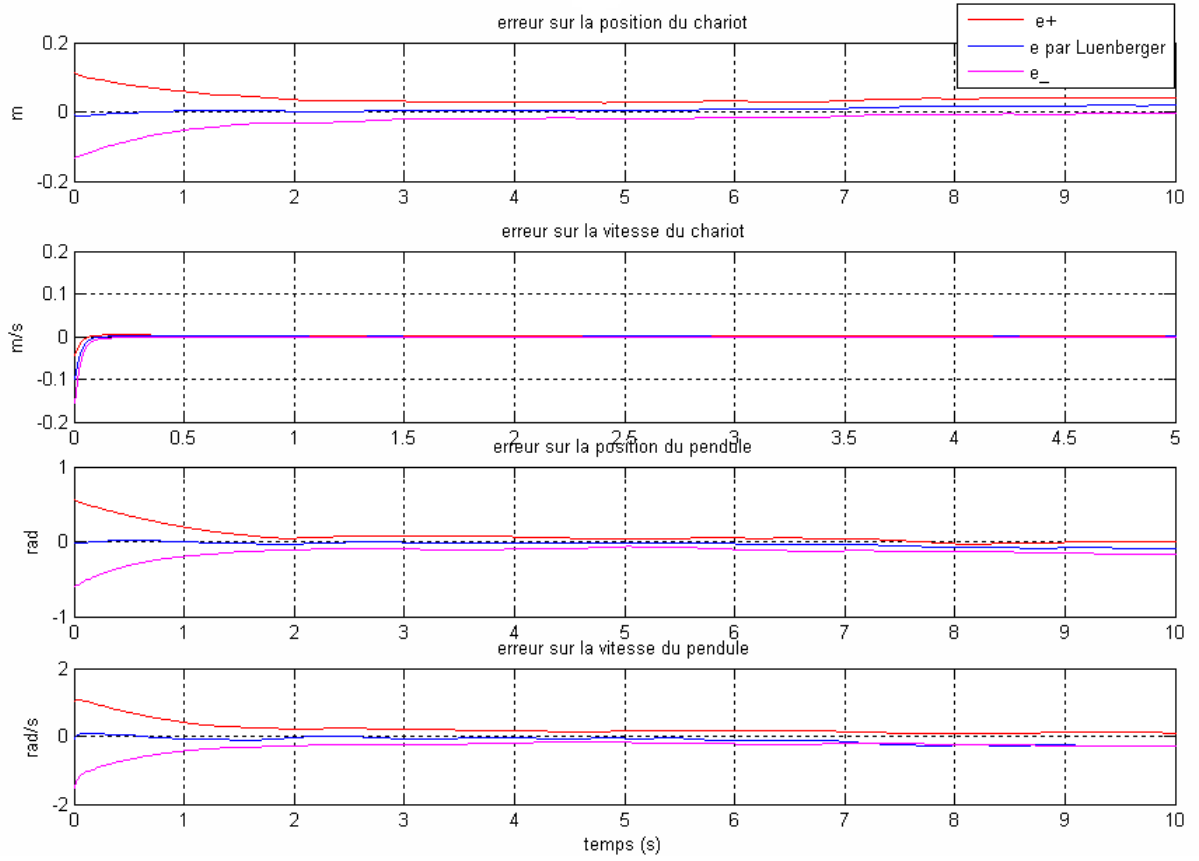


Figure 4.14: Limites supérieure et inférieure de l'erreur dynamique obtenues par un observateur intervalle pour $w^+ = w(t) + 0.01$ et $w^- = w(t) - 0.01$

Nous avons construit un observateur intervalle pour l'erreur $e = x - \hat{x}$ associée à l'observateur classique de Luenberger. Cet observateur intervalle permet de garantir de contenir l'évolution de l'erreur dynamique par les deux limites supérieure et inférieure, qui sont obtenues en considérant les incertitudes sur les bornes supérieure et inférieure des perturbations $w(t)$ et des conditions initiales de l'erreur e_0 .

4.4.3 Observateur intervalle du système

En utilisant simultanément l'observateur de Luenberger classique $\hat{x}(t)$ (4.49) et l'observateur intervalle construit pour l'erreur dynamique décrit par les équations (4.69), nous pouvons déduire l'observateur intervalle du système :

$$e^- + \hat{x}(t) \leq x(t) \leq e^+ + \hat{x}(t) \quad (4.70)$$

Et les conditions initiales du vecteur d'état à $t_0 = 0s$ vérifient:

$$e_0^- + \hat{x}_0 \leq x_0 \leq e_0^+ + \hat{x}_0 \quad (4.71)$$

Les figures suivantes représentent l'estimation d'état du système par un observateur intervalle, en présence d'incertitudes sur les perturbations en sortie et les conditions initiales du vecteur d'état qui sont inconnues mais bornées dans un vecteur intervalle. Nous faisons varier à chaque fois la valeur de w^+ et w^- dans (4.43) par le paramètre Γ :

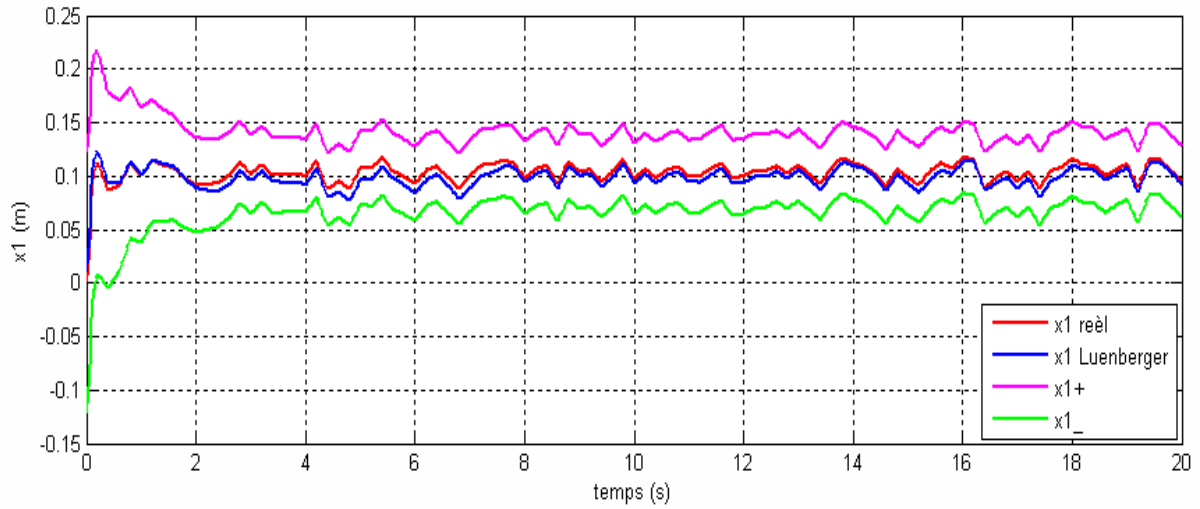


Figure 4.15 : Estimation d'état de la position du chariot (x_1), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par l'observateur intervalle en vert et en violet pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$.

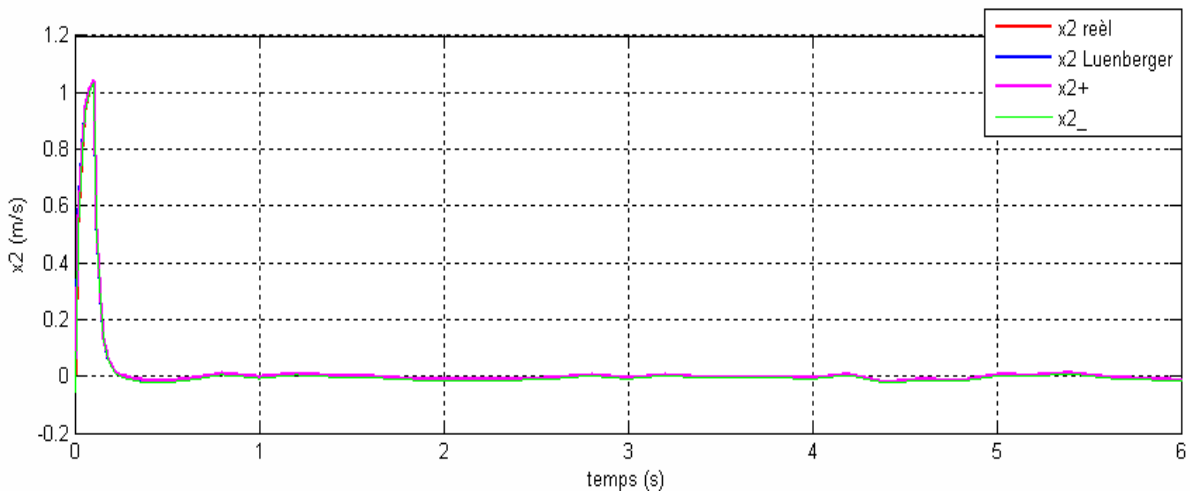


Figure 4.16 : Estimation d'état de la vitesse du chariot (x_2), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet et pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$.

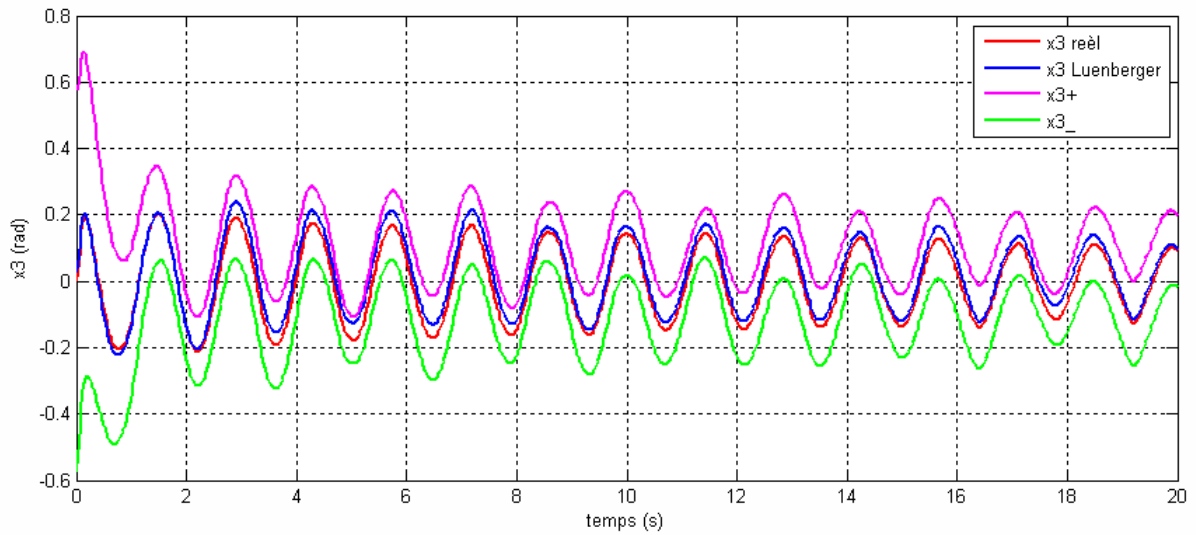


Figure 4.17 : Estimation d'état de l'angle du pendule (x_3), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et violet, pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$.

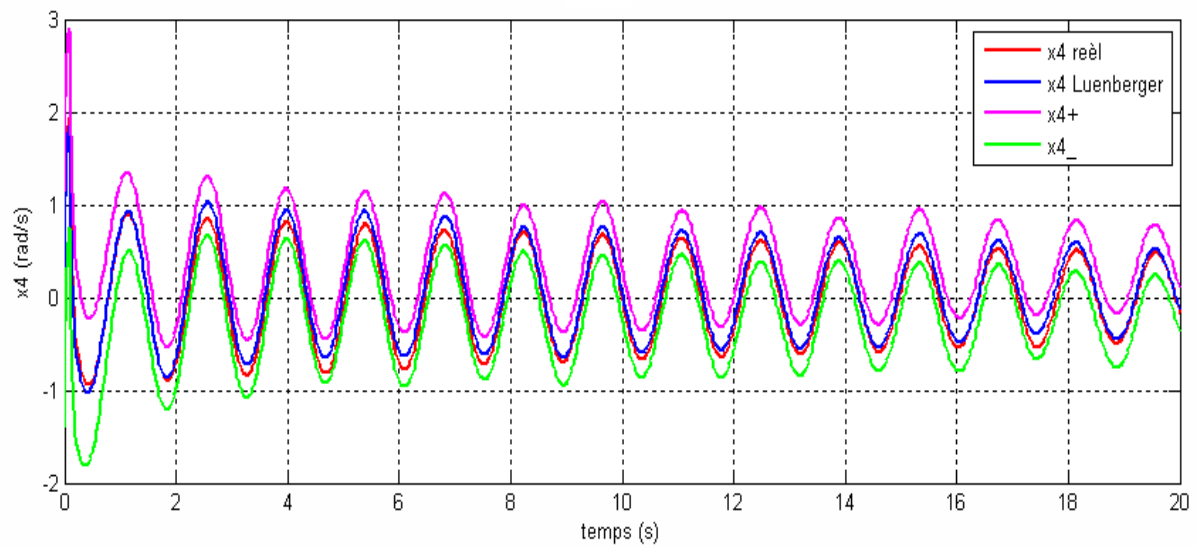


Figure 4.18 : Estimation d'état de la vitesse angulaire du pendule (x_4), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet, pour $w^+ = w(t) + 0.015$ et $w^- = w(t) - 0.015$.

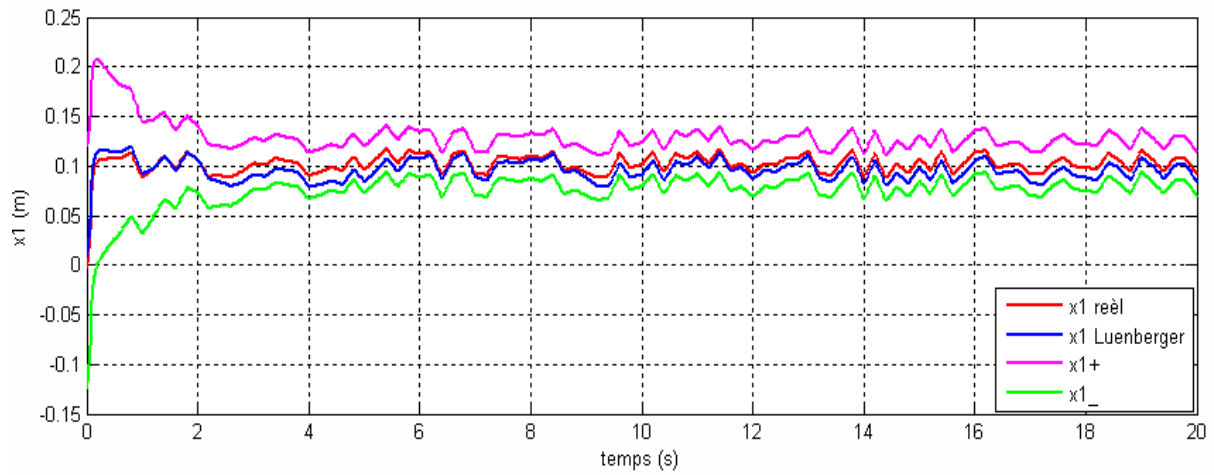


Figure 4.19: Estimation d'état de la position du chariot (x_1), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet, pour $w^+ = w(t) + 0.01$, $w^- = w(t) - 0.01$

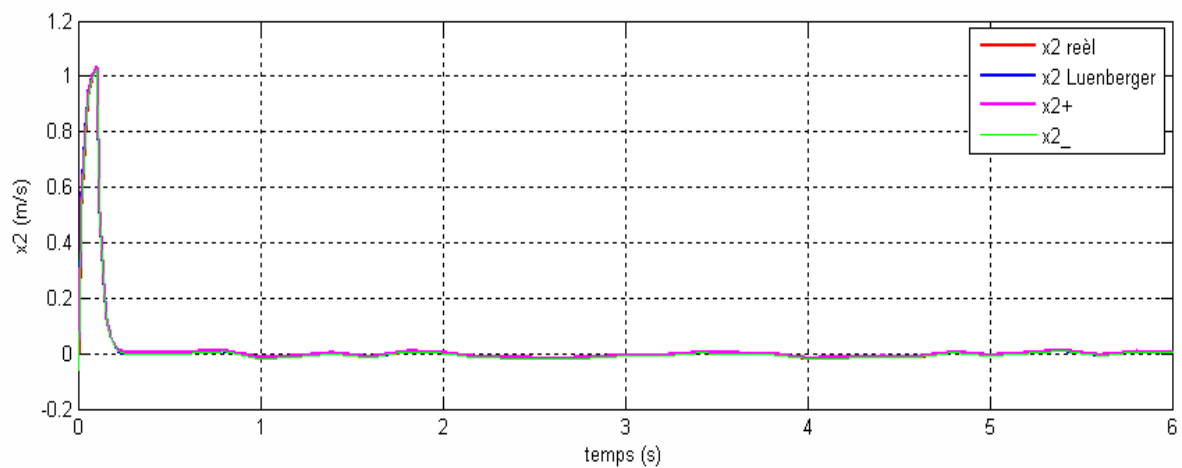


Figure 4.20 : Estimation d'état de la vitesse du chariot (x_2), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par l'observateur intervalle en vert et en violet, pour $w^+ = w(t) + 0.01$ et $w^- = w(t) - 0.01$.

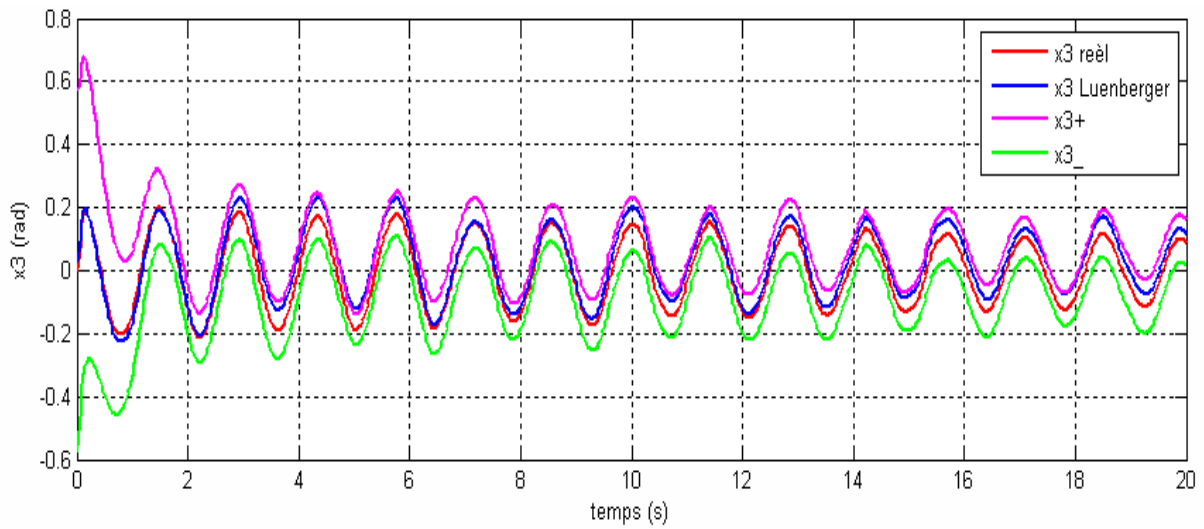


Figure 4.21 : Estimation d'état de l'angle du pendule (x_3), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet et pour $w^+ = w(t) + 0.01$ et $w^- = w(t) - 0.01$

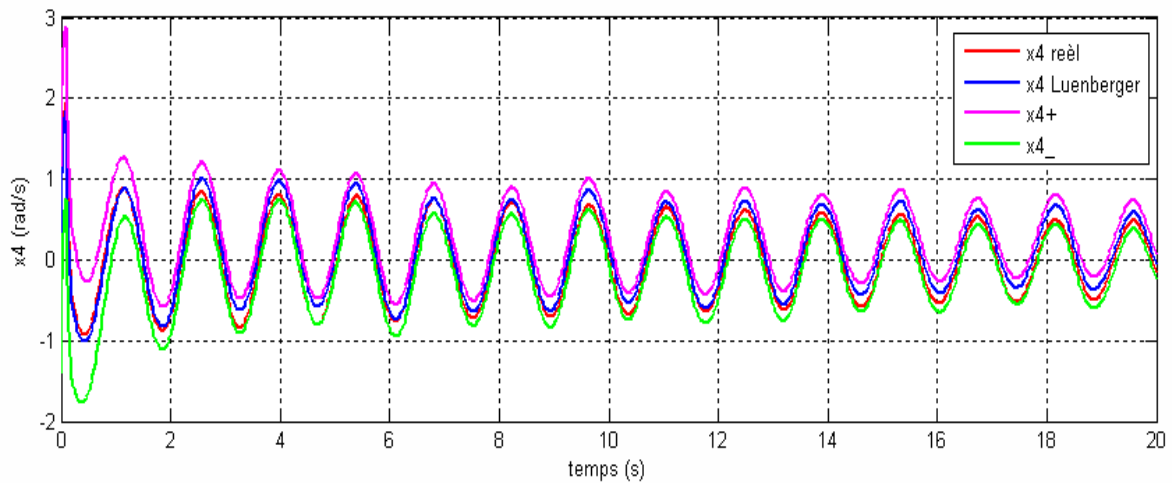


Figure 4.22: Estimation d'état de la vitesse angulaire du pendule (x_4), l'état réel est représenté en rouge, son estimation avec Luenberger classique en bleu et par un observateur intervalle en vert et en violet et pour $w^+ = w(t) + 0.01$ et $w^- = w(t) - 0.01$.

Nous remarquons d'après les tracés des figures ci-dessus, qu'en variant la valeur du paramètre Γ des bornes w^+ et w^- de la perturbation $w(t)$ en le réduisant, nous obtenons la convergence de l'estimation d'état supérieure $x^+(t)$ (en violet) et l'estimation inférieure (en vert) $x^-(t)$ obtenues par l'observateur intervalle vers l'état réel du système.

L'observateur intervalle conçu pour le modèle pendule-chariot en boucle ouverte a pu prendre en charge ces incertitudes influant sur le système (conditions initiales inconnues et perturbations incertaines sur les mesures de sortie). Et ceci en enfermant l'évaluation de l'état réel entre deux limites qui sont les estimations supérieure et inférieure retrouvées par l'observateur supérieur et inférieur de l'observateur intervalle, offrant ainsi la garantie de contenir l'état réel du système incertain.

4.5 Commande par retour d'état du pendule inversé

4.5.1 Description de la méthode

La méthode que nous développons pour la commande du pendule, se base sur une approche classique d'une commande par retour d'état à la quelle nous allons introduire cet aspect d'incertitudes influant sur le système par l'approche intervalle.

La méthode est décrite de la manière suivante :

Nous disposons du modèle d'état linéaire (autour du point d'équilibre $\pi = 0$) en boucle ouverte du pendule inversé avec les perturbations additives en sorties (4.39) que nous réécrivons :

$$\begin{cases} \dot{x} = Ax + Bu \\ y = Cx + \varphi(t) \end{cases}$$

Un régulateur d'état pour le système est un régulateur linéaire fournissant un signal de commande $u(t)$ proportionnel à chacune des variables d'état le décrivant, elle est donnée par :

$$u = -K x(t)$$

Le vecteur ligne K contient les gains associés à chaque variable d'état. Dans ce cas, le régulateur d'état est utilisé pour régler l'inclinaison du pendule et la position du chariot.

Il est à noter que l'entrée du modèle d'état est une force $F(t)$ (4.26), qui est reliée à l'entrée réelle du système $u(t)$ (tension) par le coefficient de force ks qu'il faut alors ajouter dans le retour d'état.

Alors la loi de commande est donnée par :

$$u = -ks K x(t) \quad (4.72)$$

Si les états du système sont tous disponibles à la mesure, nous aurons le modèle d'état en boucle fermé suivant :

$$\begin{cases} \dot{x} = (A - ks B K)x \\ y = Cx + \varphi(t) \end{cases} \quad (4.73)$$

Le choix du gain K se fait par placement de pôles en choisissant une dynamique désirée décrite par la matrice :

$$A_c = (A - ks B K)$$

Comme les états du système ne sont pas tous disponibles à la mesure, nous avons développé dans la section (4.4.1) un observateur de Luenberger pour reconstruire les états du système perturbé sans considérer les incertitudes.

Nous réécrivons l'observateur de Luenberger:

$$\dot{\hat{x}} = A\hat{x} + Bu + L((Cx + \varphi(t)) - C\hat{x})$$

$\hat{x}$ est l'état estimé.

Nous avons la commandabilité et l'observabilité du système (4.39) en boucle ouverte qui sont vérifiées, alors le signal de commande $u(t)$ pour le retour d'état reconstruit par l'observateur de Luenberger est donné comme suit :

$$u = -ks K \hat{x}(t) \quad (4.74)$$

Alors, le système bouclé par la loi de commande ci-dessus est donnée comme suit :

$$\dot{x} = A x(t) - ks B K \hat{x}(t) \quad (4.75)$$

Et l'observateur en présence du retour d'état s'écrit :

$$\dot{\hat{x}} = (A - ks B K - L C) \hat{x} + L Cx + L\varphi(t) \quad (4.76)$$

L'erreur dynamique $e = (x - \hat{x})$ en boucle fermée correspondante satisfait l'équation différentielle suivante :

$$\dot{e} = (A - L C) e - L\varphi(t) \quad (4.77)$$

Posons

$$A_{ce} = (A - L C) \text{ et } w(t) = -L\varphi(t)$$

alors l'erreur dynamique est réécrite comme suit :

$$\dot{e} = A_{ce} e + w(t) \quad (4.77)'$$

Dans le cas d'une connaissance parfaite des conditions initiales et des perturbations $\varphi(t)$ additives en sortie du système, l'objectif de la commande par retour d'état reconstruit par l'observateur de Luenberger est de maintenir les sorties du système à zéro.

Sauf que dans le cas étudié, les perturbations en sortie du système sont incertaines et les conditions initiales du vecteur d'état sont inconnues. Alors, afin de pouvoir réaliser une commande par retour d'état qui prend en compte ces incertitudes, nous allons faire introduire un observateur intervalle pour l'erreur dynamique.

Cette procédure nous garantit d'avoir une évolution de l'état corrigé enfermé par deux limites supérieure et inférieure en présence de ces incertitudes.

Nous avons l'erreur dynamique (4.77)'qui satisfait toutes les conditions du théorème 3.5 [Mazenc et Bernard, 2011], qui sont principalement:

- La stabilité de la matrice A_{ce} qui est garantie par le gain de l'observateur L calculé précédemment (4.36).

La matrice A_{ce} est égale à la matrice $A_e = (A - LC)$ donnée par (4.54) dont les valeurs propres sont à parties réelles strictement négatives :

$$(\lambda_1 = -1, \lambda_2 = -34.9734 \text{ et } \lambda_{3,4} = -1.5 \pm 0.5 i)$$

- $w(t)$ est une fonction Lipschitz continue inconnue bornée par deux fonctions Lipschitz continues connues $w^-(t)$ et $w^+(t)$.
- Les conditions initiales de l'erreur e_0 sont bornées par e_0^- et e_0^+ qui sont les bornes inférieure et supérieure respectivement.

Alors cette erreur dynamique admet l'observateur intervalle définit par (4.45) et (4.46).

Par la suite, nous obtenons à partir de cet enveloppement de l'erreur un état corrigé borné par deux trajectoires supérieure et inférieure. Ainsi, nous assurons d'enfermer la trajectoire réelle de l'état corrigé du système qui est donnée comme suit :

$$e^-(t) + \hat{x}(t) \leq x(t) \leq e^+(t) + \hat{x}(t) \quad (4.78)$$

avec

$\hat{x}(t)$ est l'état corrigé obtenu par le retour d'état reconstruit par l'observateur de Luenberger en connaissance parfaite des perturbations et des conditions initiales du système.

et

$e^+(t) + \hat{x}(t)$: est la limite supérieure.

$e^-(t) + \hat{x}(t)$: est la limite inférieure.

4.5.2 Commande par retour d'état reconstruit par l'observateur de Luenberger et l'observateur intervalle

a) Commande par retour d'état reconstruit par l'observateur de Luenberger en absence d'incertitudes

Le gain du retour d'état K est choisi de telle manière à avoir les états du système bouclé maintenus à zéro en régime permanent et réduire l'effet des perturbations en sortie. En choisissant une dynamique du système bouclé qui est décrite par la matrice $A_c = (A - ks B K)$ donnée par le choix des pôles suivants:

$$\begin{aligned} p1 &= -35.7199 \\ p2 &= -16 \\ p3 &= -1 + 0.1i \\ p4 &= -1 - 0.1i \end{aligned} \quad (4.79)$$

nous aurons le gain du correcteur donné par le vecteur ligne suivant :

$$K = (0.7846 \quad 0.6897 \quad 8.3141 \quad -0.0958) \quad (4.80)$$

La matrice A_c calculée est donnée par :

$$A_c = \begin{pmatrix} 0 & 1.0000 & 0 & 0 \\ -29.4207 & -60.8632 & -312.2858 & 3.5910 \\ 0 & 0 & 0 & 1.0000 \\ -58.8413 & -121.7263 & -644.1915 & 7.1433 \end{pmatrix} \quad (4.81)$$

Nous avons obtenu les réponses représentées par la figure 4.23 pour les conditions initiales du système réel bouclé et celles de l'observateur de Luenberger par le retour d'état (4.74) respectivement :

$$\begin{aligned} x_0 &= (0, 0, 0.17, 0)^T \\ \hat{x}_0 &= (0.01, 0, 0.1, 0)^T \end{aligned} \quad (4.82)$$

en prenant $\varphi(t)$ un bruit de mesure uniformément distribué sur l'intervalle $[-0.02, 0.01]$.

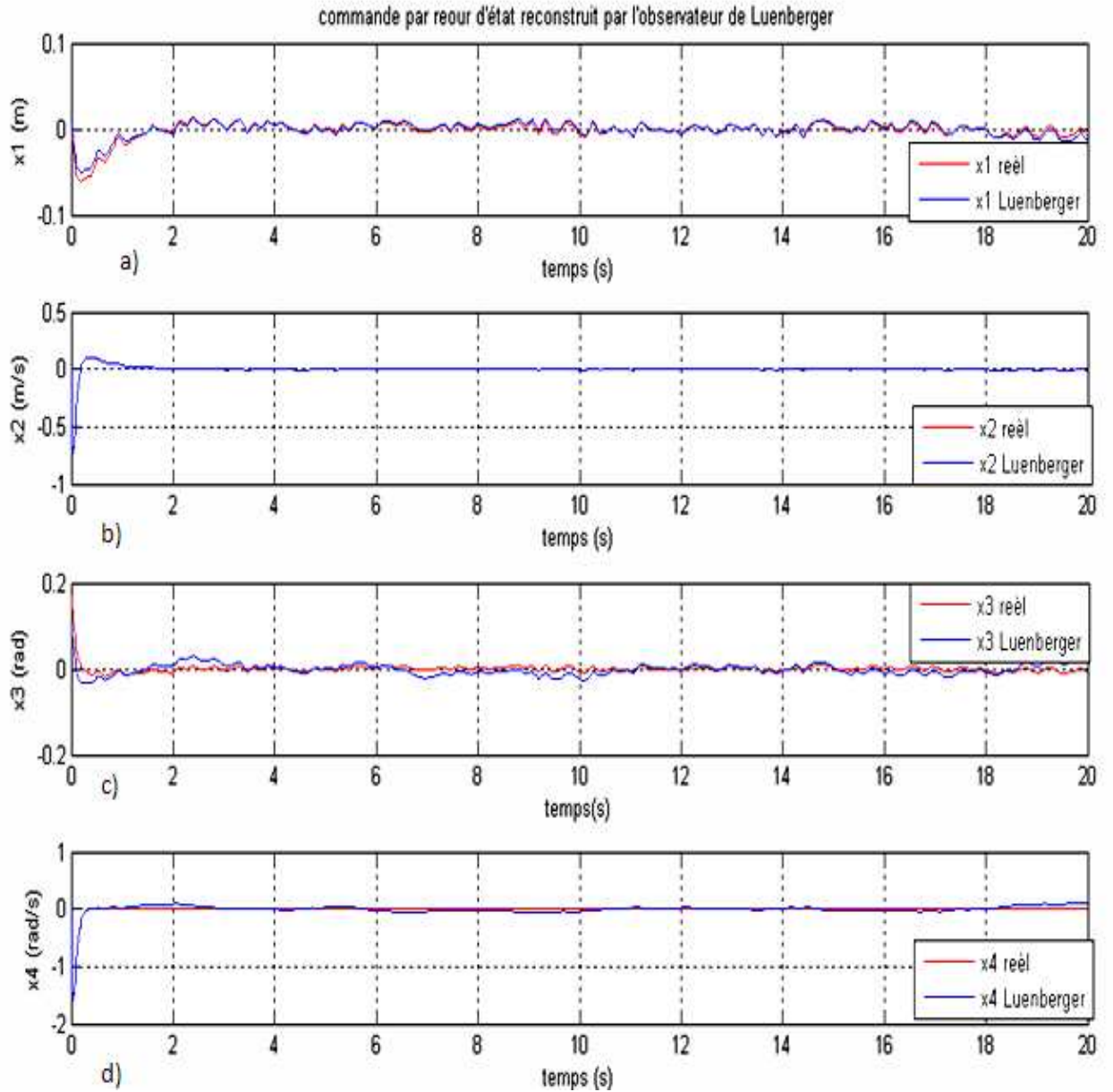


Figure 4.23 : Réponses du système réel bouclé en rouge, l'état corrigé reconstruit par l'observateur de Luenberger en bleu. a) et b) Représentent la position et la vitesse du chariot respectivement, c) et d) représentent la position angulaire et la vitesse angulaire du pendule respectivement.

L'évolution de l'erreur dynamique en boucle fermée des quatre états du système avec comme condition initiale pour l'erreur :

$$e_0 = x_0 - \hat{x}_0 = (-0.0100, 0, 0.0700, 0)^T \quad (4.83)$$

est donnée par la figure ci-dessous :

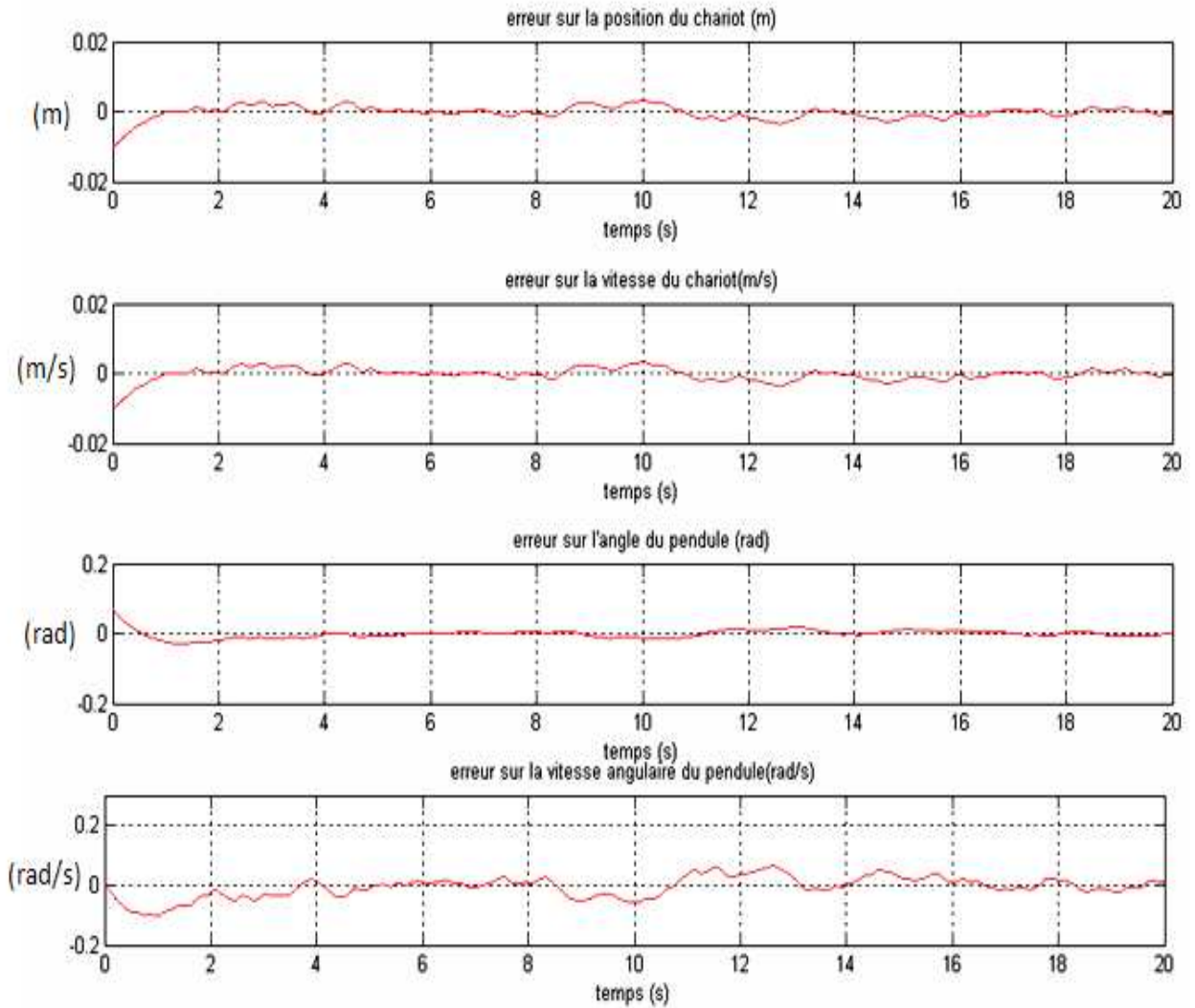


Figure 4.24 : Evolution de l'erreur dynamique en boucle fermée

Nous remarquons d'après les tracés des réponses du système bouclé figure 4.23 que celles-ci sont maintenues à zéro avec une réduction de l'effet des perturbations en sortie. Les réponses obtenues par le retour d'état reconstruit par l'observateur de Luenberger suivent les réponses du système réel bouclé.

Et d'après l'évolution de l'erreur dynamique $e = x - \hat{x}$ (figure 2.24), celle-ci tend vers zéro. Donc, les états estimés par l'observateur classique de Luenberger convergent vers les états réels du système bouclé.

b) Commande par retour d'état en présence d'incertitudes

La commande par retour d'état reconstruit par l'observateur de Luenberger ne prend pas en considération les incertitudes sur les perturbations en sortie et des conditions initiales qui sont inconnues. Alors, la méthode que nous allons utiliser se base sur l'observateur intervalle construit pour l'erreur dynamique en boucle fermée (4.77)'. A partir des limites supérieure et inférieure retrouvées pour cette erreur, nous obtenons un enveloppement (4.78) qui nous garantit d'enfermer l'évolution de l'état réel corrigé du système.

b.1) L'observateur intervalle de l'erreur dynamique en boucle fermée

L'observateur intervalle pour l'erreur dynamique en boucle fermée (4.77)' est le même que celui construit en boucle ouverte (section 4.4.2) avec des conditions initiales différentes. Nous réécrivons la structure de l'observateur intervalle de l'erreur dynamique en boucle fermée qui est donnée par le système différentiel suivant :

$$\begin{cases} \dot{z}^+ = N z^+ + P^+(t) w^+(t) - P^-(t) w^-(t) \\ \dot{z}^- = N z^- + P^+(t) w^-(t) - P^-(t) w^+(t) \end{cases} \quad (4.84)$$

avec N , $P^+(t)$, $P^-(t)$ sont les matrices calculées précédemment données respectivement par (4.61), (4.63) et $P^-(t)$ qui est égale à $P^+(t) - P(t)$ ($P(t)$ est donnée par (4.60)).

Les conditions initiales pour la résolution du système différentiel (4.84) sont calculées par les deux équations suivantes :

$$\begin{aligned} z^+(t_0) &= P^+(t_0)e_0^+ - P^-(t_0)e_0^- \\ z^-(t_0) &= P^+(t_0)e_0^- - P^-(t_0)e_0^+ \end{aligned}$$

Avec

$P^+(t_0)$, $P^-(t_0)$ à l'instant ($t_0 = 0$) sont les matrices calculées précédemment données par (4.65) et (4.66) respectivement

et

e_0^+ , e_0^- sont les bornes supérieure et inférieure des conditions initiales e_0 de l'erreur dynamique en boucle fermée (4.83), qui sont données par les vecteurs colonnes suivants :

$$\begin{aligned} e_0^+ &= (0.0100, 0.0200, 0.0900, 0.0500)^T \\ e_0^- &= (-0.0300, -0.0200, 0.0500, -0.0500)^T \end{aligned}$$

qui doivent contenir e_0 , vérifiant ainsi :

$$e_0^- \leq e_0 \leq e_0^+$$

Nous obtenons alors les conditions initiales du système différentiel (4.84) données par les vecteurs colonnes ci-dessous :

$$\begin{aligned} z^+(t_0) &= (0.1094, 0.0454, 0.4764, 0.5703)^T \\ z^-(t_0) &= (-0.1439, -0.0468, -0.1234, 0.0107)^T \end{aligned}$$

Finalement, les limites supérieure et inférieure enferment l'erreur en boucle fermée sont données par le système d'équations suivant :

$$\begin{aligned} e^+ &= M^+(t)z^+ - M^-(t)z^- \\ e^- &= M^+(t)z^- - M^-(t)z^+ \end{aligned} \quad (4.85)$$

Avec $M^+(t)$ est donnée par (4.67) et $M^-(t)$ est égale à $(M^+(t) - M(t))$ ($M(t)$ est donnée par (4.62) qui est l'inverse de $P(t)$)

b.2) Limites supérieure et inférieure pour l'état corrigé

A partir de e^+ et e^- qui sont les bornes supérieure et inférieure de l'erreur dynamique e en boucle fermée, nous obtenons les trajectoires supérieure $e^+ + \hat{x}$ et inférieure $e^- + \hat{x}$ (4.78) enfermant la trajectoire de l'état réel corrigé x .

Les incertitudes sont introduites dans $w(t) = -L\varphi(t)$, vérifiant :

$$w^-(t) \leq w(t) \leq w^+(t)$$

avec $w^+(t) = w(t) + \Gamma$ et $w^-(t) = w(t) - \Gamma$

Nous avons obtenu les tracés des réponses en boucle fermée et de l'erreur associée pour $\Gamma = 0.005$:

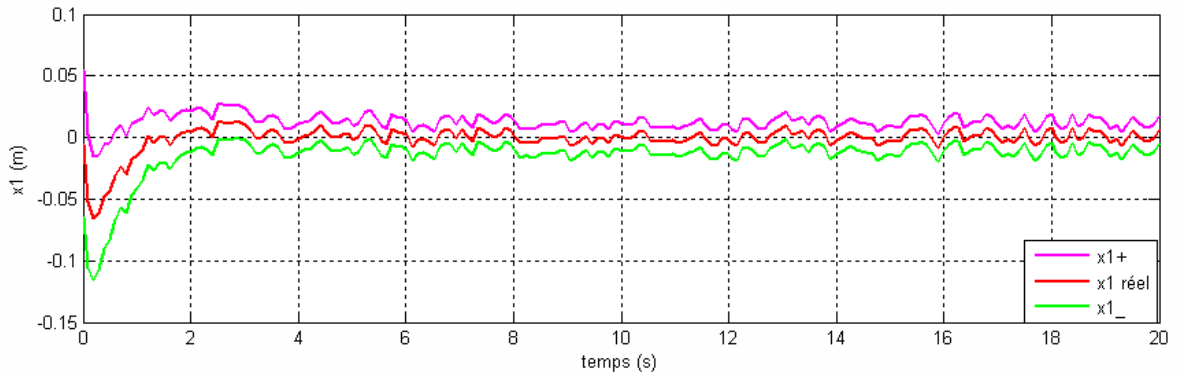


Figure 4.25 : Réponse de la position du chariot par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.005$

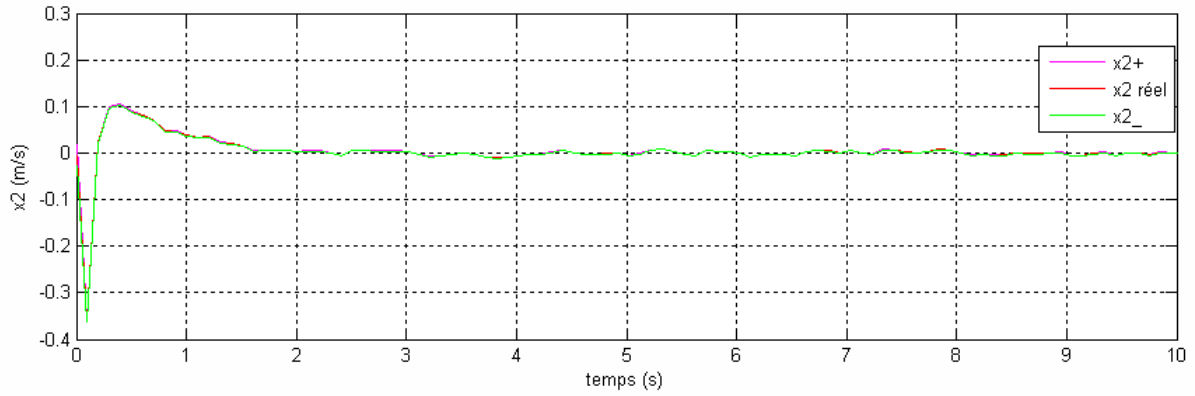


Figure 4.26 : Réponse de la vitesse du chariot par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.005$

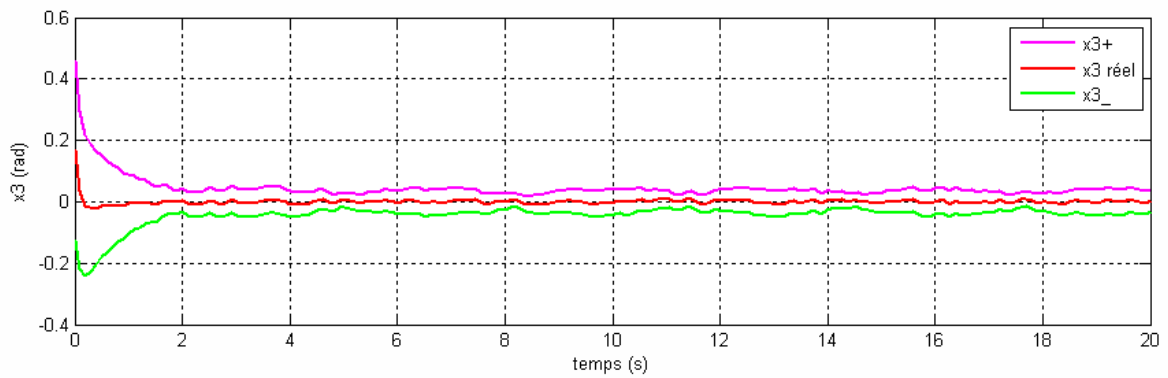


Figure 4.27 : Réponse de la position angulaire du pendule par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.005$.

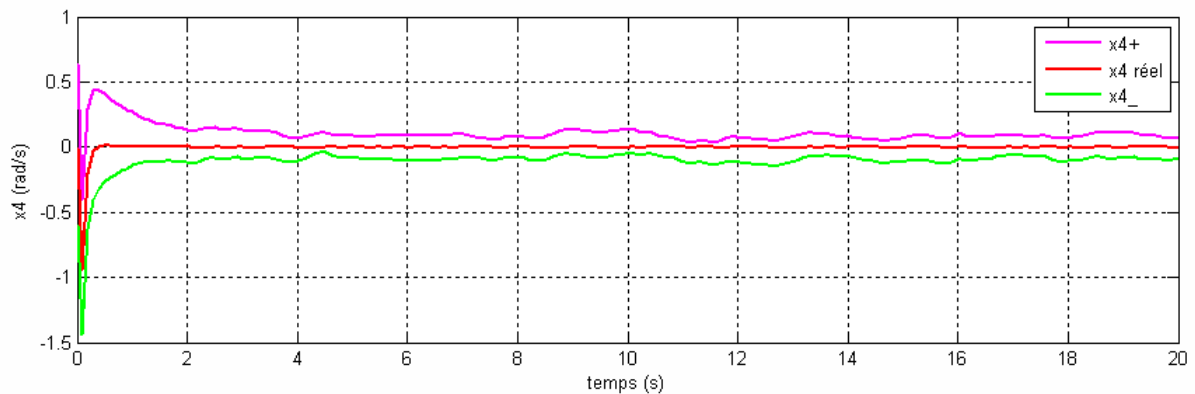


Figure 4.28 : Réponse de la vitesse angulaire du pendule par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.005$.

Les limites supérieure et inférieure de l'évolution de l'erreur en boucle fermée, obtenues par l'observateur intervalle sont données par la figure suivante:

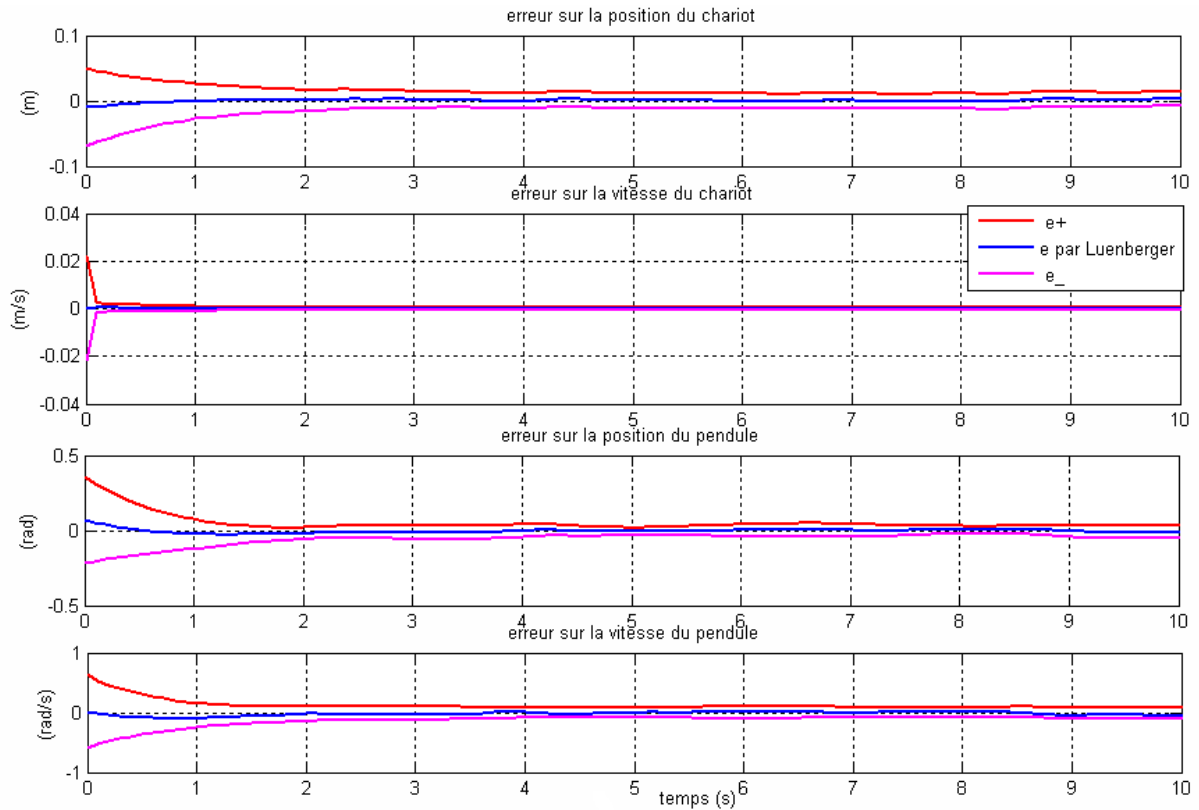


Figure 4. 29: Limites supérieure et inférieure de l'erreur en boucle fermée obtenues par l'observateur intervalle en rouge et en violet enfermant l'erreur associée à l'observateur de Luenberger classique en bleu pour une valeur de Γ égale à 0.005

Et pour la valeur de $\Gamma = 0.002$, nous avons obtenu les tracés des réponses en boucle fermée et de l'erreur associée représentés par les figures ci-dessous :

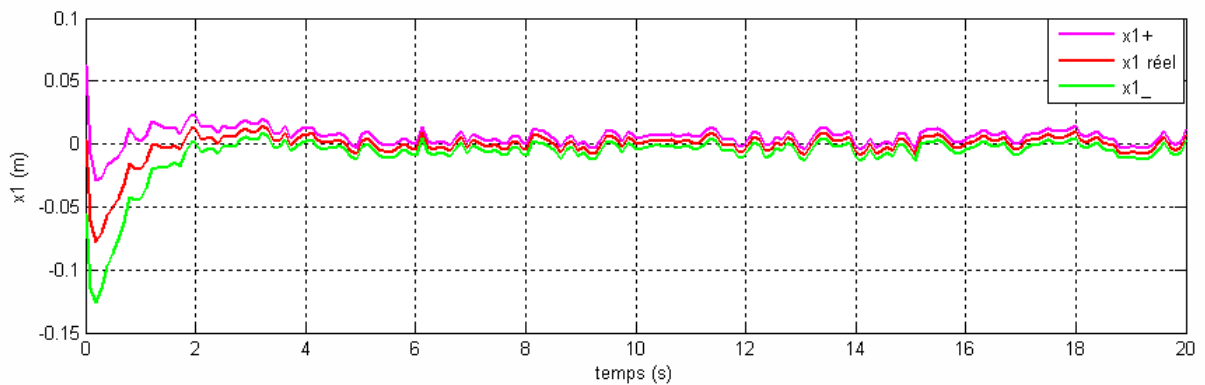


Figure 4.30: Réponse de la position du chariot par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.002$

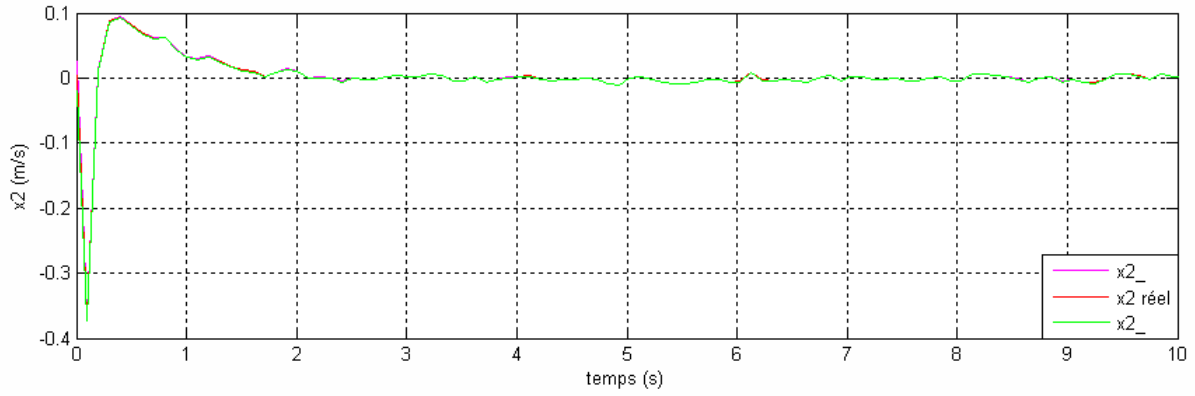


Figure 4.31: Réponse de la vitesse du chariot par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.002$.

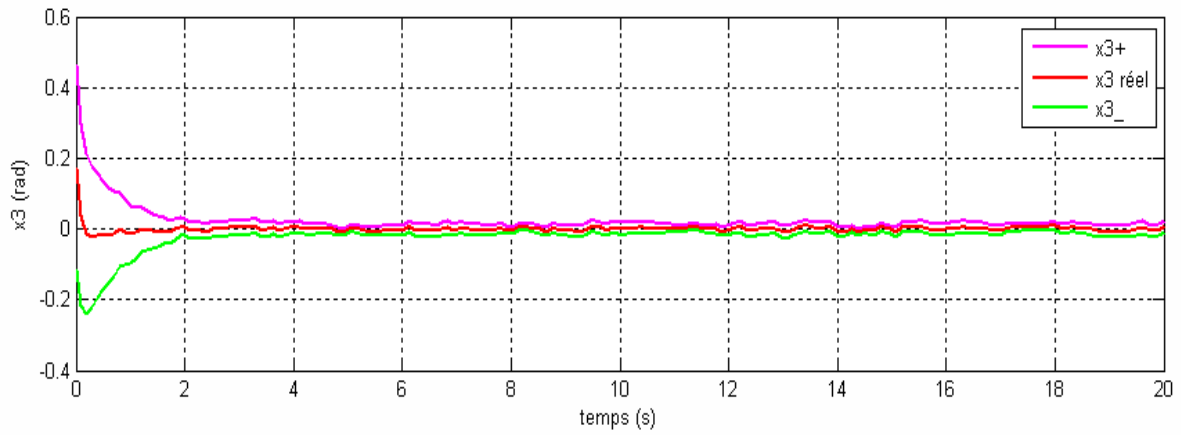


Figure 4.32 : Réponse de la position angulaire du pendule par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.002$.

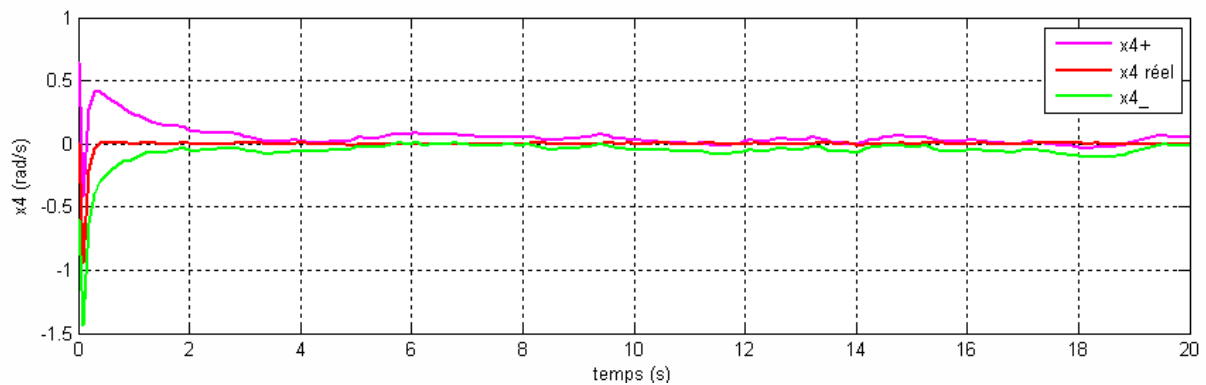


Figure 4.33 : Réponse de la vitesse angulaire du pendule par le retour d'état, en violet et en vert les trajectoires supérieure et inférieure obtenues par l'observateur intervalle qui enferment l'état réel corrigé en rouge pour $\Gamma = 0.002$.

Les limites supérieure et inférieure de l'évolution de l'erreur en boucle fermée obtenues par l'observateur intervalle sont données par la figure suivante:

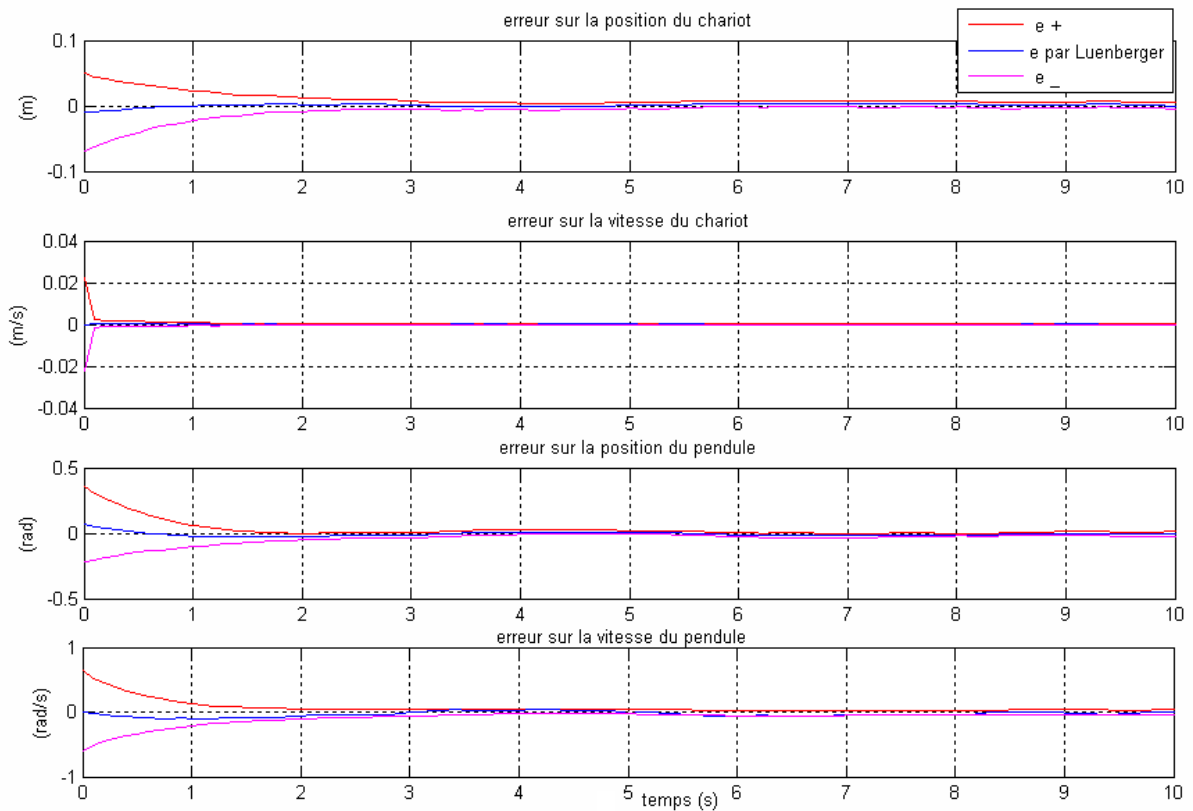


Figure 4. 34: Limites supérieure et inférieure de l'erreur en boucle fermée obtenues par l'observateur intervalle en rouge et en violet enfermant l'erreur associée à l'observateur de Luenberger classique en bleu pour une valeur de Γ égale à 0.002 .

D'après les tracés des réponses obtenues pour le système bouclé en présence des incertitudes sur les perturbations additives en sortie et les conditions initiales, nous avons la garantie que l'état réel corrigé évolue à l'intérieur de deux limites supérieure et inférieure obtenues à l'aide de l'observateur intervalle.

Le problème que nous avons rencontré lors des simulations de cette application est que la variation de Γ des incertitudes sur les perturbations est faible (elle ne doit pas dépasser 0.01) et ceci afin d'avoir la convergence de la limite supérieure et inférieure vers l'état réel corrigé.

Cependant, le fait de prendre Γ faible dans cet exemple n'est pas une limite à la faisabilité de la méthode, car en introduisant des incertitudes sur le système, l'analyse par la combinaison « observateur classique-observateur intervalle » semble le meilleur moyen de reconstruire et corriger l'état d'un système incertain. Et cette technique pourrait être plus intéressante dans d'autres applications.

4.6 Conclusion

Le quatrième chapitre a été consacré à mettre en application l'observateur intervalle temps variant développé dans le chapitre précédent, qui était conçu pour la classe fondamentale des systèmes LTI non coopératifs exponentiellement stable, présentant des perturbations additives.

Les différentes techniques d'estimation d'état utilisant l'analyse par intervalles, ont été développées pour une classe restreinte de systèmes, exigeant cette condition de coopérativité. L'observateur intervalle dont est basé notre travail fait face à cette difficulté, en offrant un moyen de transformer ces systèmes par des changements de coordonnées temps-variant en systèmes coopératifs.

Cette approche a été appliquée au modèle d'état linéaire d'un pendule inversé en boucle ouverte, qui présentait des perturbations additives mal connues en sortie du système et des incertitudes sur les conditions initiales des variables d'état.

La technique d'estimation d'état que nous avons développé dans notre application, consistait à l'utilisation simultanée « observateur classique-observateur intervalle ». Nous avons constaté qu'en combinant ces deux observateurs, nous obtenons une estimation d'état offrant la certitude et la garantie d'envelopper l'état réel du système entre deux trajectoires, qui représentent l'estimation supérieure et inférieure obtenues par l'observateur intervalle.

Par la suite, nous avons bouclé le système incertain en utilisant une commande par retour d'état reconstruit par l'utilisation simultanée de l'observateur classique de Luenberger et de l'observateur intervalle. Et ceci afin d'avoir la certitude de contenir l'état réel corrigé entre deux trajectoires supérieure et inférieure.

Conclusion générale et perspectives

Conclusion générale et perspectives

Dans notre travail présenté dans ce mémoire, l'objectif était de développer une approche capable de tenir compte des incertitudes (numériques ou physiques) inévitables influant sur les données d'un système, qui pouvaient devenir un sérieux handicap dans le traitement des données et dans l'interprétation des résultats. Le fait qu'elle présentait des caractéristiques à mieux modéliser ces incertitudes, et qui offrait une manipulation aisée des calculs, ne réclamant aucune information statistique sur les variables incertaines comparées aux autres méthodes statistiques développées dans ce sens, l'approche ensembliste a été le meilleur choix pour caractériser plus explicitement ces données entachées d'erreurs dans notre thème d'étude.

Après avoir présenté les fondements de base de cette approche telles que les notions principales reliées à la manipulation des ensembles et illustré quelques représentations ensemblistes, nous avons constaté que la majorité de ces représentations se basent sur un outil de calcul qui est l'arithmétique des intervalles.

Par la suite, nous avons donné un aperçu sur l'historique de cette arithmétique, sa nécessité et ses débuts d'application en automatique. L'avantage principal des intervalles résidait dans le fait qu'ils fournissaient des limites bornées finies connues a priori, garantissant un enveloppement aux données incertaines. Puis nous avons abordé les principes de base de cet outil, nous avons souligné l'une des difficultés engendrée par la manipulation des vecteurs intervalles qui était le pessimisme présent dans le résultat des calculs, nous avons présenté les solutions apportées à ce problème. Des outils mettant en application cette arithmétique tels que l'inversion ensembliste par l'arithmétique des intervalles et la caractérisation d'ensembles par contraction ont été illustrés, qui sont généralement utilisés pour caractériser les systèmes d'équations linéaires et non linéaires.

Dans le cadre de notre travail, nous avons introduit l'outil de l'arithmétique d'intervalles pour redéfinir les propriétés structurelles telles que la commandabilité et l'observabilité des systèmes LTI à paramètres incertains et ceci en se basant sur la propriété de l'indépendance (dépendance) linéaire des vecteurs intervalles. L'utilisation de cette arithmétique dans le domaine d'estimation d'état a permis de prendre en charge les incertitudes influant sur un système donné.

Nous avons développé dans notre travail, deux observateurs intervalles appliqués à la classe des systèmes LTI incertain :

Le premier observateur intervalle a été conçu pour la classe des systèmes LTI à incertitudes paramétriques, l'avantage qu'offrait cet observateur était dans le fait qu'il fournissait une estimation d'état qui n'évolue plus dans un vecteur intervalle. Et ceci au moyen d'une estimation pondérée, calculée en utilisant la sortie disponible à la mesure et les bornes supérieure et inférieure des estimations retrouvées par l'observateur intervalle. Nous avons

illustré cette méthode sur un exemple d'application d'un réacteur chimique continu, qui présentait des incertitudes paramétriques dans son modèle d'état.

La limitation rencontrée dans cette méthode d'estimation, était dans le fait qu'elle exigeait de poser une hypothèse restreinte qui était la condition de coopérativité du système. Le deuxième observateur intervalle développé, permettait de palier à cet obstacle par le moyen d'un changement de coordonnées « temps-variant » qui transformait le système en un système coopératif. Cependant, cette méthode ne concernait qu'une classe particulière des systèmes LTI stable en boucle ouverte présentant des perturbations additives. Nous avons montré toutes les étapes de construction de cet observateur intervalle sur un exemple d'un système LTI stable d'ordre trois.

Enfin, nous avons mis en application la deuxième méthode d'estimation pour la conception d'un observateur intervalle exponentiellement stable pour le modèle d'état d'un pendule inversé. Le modèle que nous avons pris présentait des perturbations additives mal connues en sortie et les conditions initiales du vecteur d'état sont incertaines. L'approche adoptée a été de combiner à un observateur classique (Luenberger), un observateur intervalle construit pour l'erreur dynamique. Cette utilisation simultanée des observateurs classiques et intervalles a permis d'obtenir des informations complémentaires sur le système perturbé incertain. Cette technique a permis d'offrir une certitude concernant l'évolution de l'état réel en l'enfermant par des limites supérieure et inférieure fournies par l'observateur intervalle.

Par la suite, nous avons bouclé le système incertain en utilisant une commande par retour d'état classique toujours en utilisant « observateur classique-observateur intervalle » pour reconstruire et corriger l'état réel du système incertain.

Les résultats obtenus dans ce mémoire offrent des perspectives intéressantes de développements ultérieurs. Comme travaux futurs, nous pouvons étendre la deuxième méthode d'estimation basée sur la combinaison observateur intervalle et observateur classique de type Luenberger comme suit :

- Synthétiser un observateur classique ayant une structure plus complexe, intégrant par exemple un gain dynamique, offrant ainsi plus de garantie d'enveloppement de l'état réel par les deux trajectoires supérieure et inférieure et offrant la convergence de celles-ci vers l'état réel du système incertain.
- Par les changements de coordonnées « temps-variant » pour transformer les systèmes non coopératifs en des systèmes coopératifs, il serait intéressant d'utiliser cette technique pour la classe des systèmes non-linéaires temps-variant et pour les systèmes à retard.
- Une autre direction à développer serait l'application au diagnostic des fautes et à la maintenance des procédés industriels.

Bibliographie

Bibliographie

- [Ahn et al., 2005] H-S.Ahn, K.L. Moore and Y.Q Chen, Linear Independency of Interval Vectors and Its Applications to Robust Controllability Tests. *Proceedings of the 44th IEEE Conference of Decision and Control, and the European Control Conference*, Seville, Spain, pp.8070-8075, December 2005.
- [Alcaraz-Gonzalez et al., 2002] V. Alcaraz-Gonzalez, J. Harmand, A. Rapaport, J.P. Steyer, V. Gonzalez-Alvarez and C. Pelayo-Ortiz, Software sensors for highly uncertain WWTPs: a new approach based on interval observers, *Water Research*, vol. 36 no.10, pp.2515-2524, 2002.
- [Athans et al., 1977] M. Athans, D. Castanon, K. Dunn, C. Greene, W. Lee, N. Sandell and A. Willsky, The Stochastic Control of the F-8C Aircraft Using a Multiple Model Adaptive Control (MMAC) Method - Part I: Equilibrium Flight, *IEEE Transactions on Automatic Control*, vol. 22 no.5, pp. 68-780, 1977.
- [Banerjee et al., 1977] A. Banerjee, Y. Arkun, B. Ogunnaike and R. Pearson, Estimation of Nonlinear Systems Using Linear Multiple Models, *AIChE Journal*, vol. 43 no.5, pp. 1204-1226, 1977.
- [Berger, 1979] M. Berger, *Espaces euclidiens, triangles, cercles et sphères*, vol. 2, Cedic Fernand Nathan, 1979.
- [Bernard et Gouzé, 2004] O. Bernard, J-L Gouzé, Closed loop observers bundle for uncertain biotechnological models, *Journal of Process Control*, vol. 14 no.7, pp. 765-774, 2004.
- [Gouzé et al., 2000] J-L.Gouzé, A.Rapaport, Z.Hadj-Sadok, Interval observers for uncertain biological systems, *Ecological Modelling*, vol. 133 no.1, pp. 45-56, 2000.
- [Hansen, 1992] E.Hansen, *Global Optimization using Interval Arithmetic*. Marcel Dekker, 1992.
- [Hargreaves, 2002] G. I. Hargreaves, *Interval Analysis in MATLAB, Numerical Analysis Report no.416*, Manchester Centre for Computational Mathematics, Université de Manchester, 2002.
- [Jaulin et Walter, 1993] L. Jaulin and E. Walter, Set inversion via interval analysis for nonlinear bounded-error estimation, *Automatica*, vol. 29 no.4, pp.1053-1064, 1993.
- [Jaulin, 1994] L.Jaulin, *Solution globale et garantie de problèmes ensemblistes : application à l'estimation de problèmes ensemblistes et à la commande*, Phd Thesis, Université de Paris-Sud, centre d'Orsay, France, 1994.
- [Jaulin et al., 2001] L. Jaulin, M. L. Kieffer, O. Didrit et E.Walter, *Applied interval analysis with Examples in Parameter and State Estimation, Robust Control and Robotics*. Springer verlag, London, 2001.

- [Jaulin, 2004] L.Jaulin, *Interval contractors and their applications*, Université d'Angers, lesson 9-12 February 2004 à l'Université Politècnica de Catalunya (Barcelone), 2004.
- [Jaulin et Chabert, 2008] L.Jaulin et G. Chabert, « Quimper : un langage de programmation pour le calcul ensembliste, Application à l'automatique », *CIFA 5ème Conférence Internationale Francophone d'Automatique*, Bucarest, Roumanie, 2008.
- [Lalami, 2008] A.Lalami, *Diagnostic et approches ensemblistes à base des zonotopes*, Phd Thesis, Ecole Nationale Supérieure de l'Electronique et de ses Applications (ENSEA) de l'université de Cergy-Pontoise, Paris, France, 2008.
- [Labarrer et al., 1993] M. Labarrer, J.P.Krief et B.Gimonet, *Le filtrage et ses applications*, Cépaduès-Edition, 1993.
- [Li et Bar-Shalom, 1966] X. Li and Y. Bar-Shalom, Multiple-Model Estimation with Variable Structure, *IEEE Transactions on Automatic Control*, vol. 41 no.4, pp.478-493, 1966.
- [Luenberger, 1964] D. Luenberger, Observing the state of a linear system, *IEEE Transaction Military Electronics*, vol.MIL-8, pp.74-80, 1964.
- [Luenberger, 1966] D.G Luenberger, Observers for multivariable systems, *IEEE Transactions on Automatic Control*, vol.11 no.2, pp.190-197, 1966.
- [Magill, 1965] D.Magill, Optimal adaptive estimation of sampled stochastic processes, *IEEE Transactions on Automatic Control*, vol.10 no.4, pp.434-439, 1965.
- [Mazenc et Bernard, 2010] F.Mazenc, O.Bernard, Asymptotically Stable Interval Observers for Planar systems with Complex Poles, *IEEE Trans on Automatic Control*, vol.55 no.2, pp.523-527,2010.
- [Mazenc et Bernard, 2011] F. Mazenc, O.Bernard, Interval observers for linear time-invariant systems with disturbances, *Automatica*, vol. 47 no.1, pp.140–147, 2011.
- [Moisan, 2007] M. Moisan, *Synthèse d'observateurs par intervalles pour des systèmes biologiques mal connus*, Phd Thesis, Université de Nice Sophia_Antipolis, France, 2007.
- [Moisan et al., 2009] M. Moisan, O. Bernard, J-L Gouzé, Near optimal interval observers bundle for uncertain bioreactors, *Automatica*, vol.45 no.1, pp.291-295, 2009.
- [Moore, 1962] R. Moore, *Interval arithmetic and automatic error analysis in digital computing*, Phd Thesis, Department of Mathematics, Stanford University, California, Published as Applied Mathematics and Statistics Laboratories Technical Report no.25, 1962.
- [Moore, 1966] R. Moore, *Interval analysis*, Prentice Hall, 1966.
- [Moussouami et Aupoix, 2007] G.Moussouami, P.Aupoix, *Etude du pendule inversé*, PhD Thesis. Université de Caen Basse-Normandie UFR de Sciences, France, 2007.

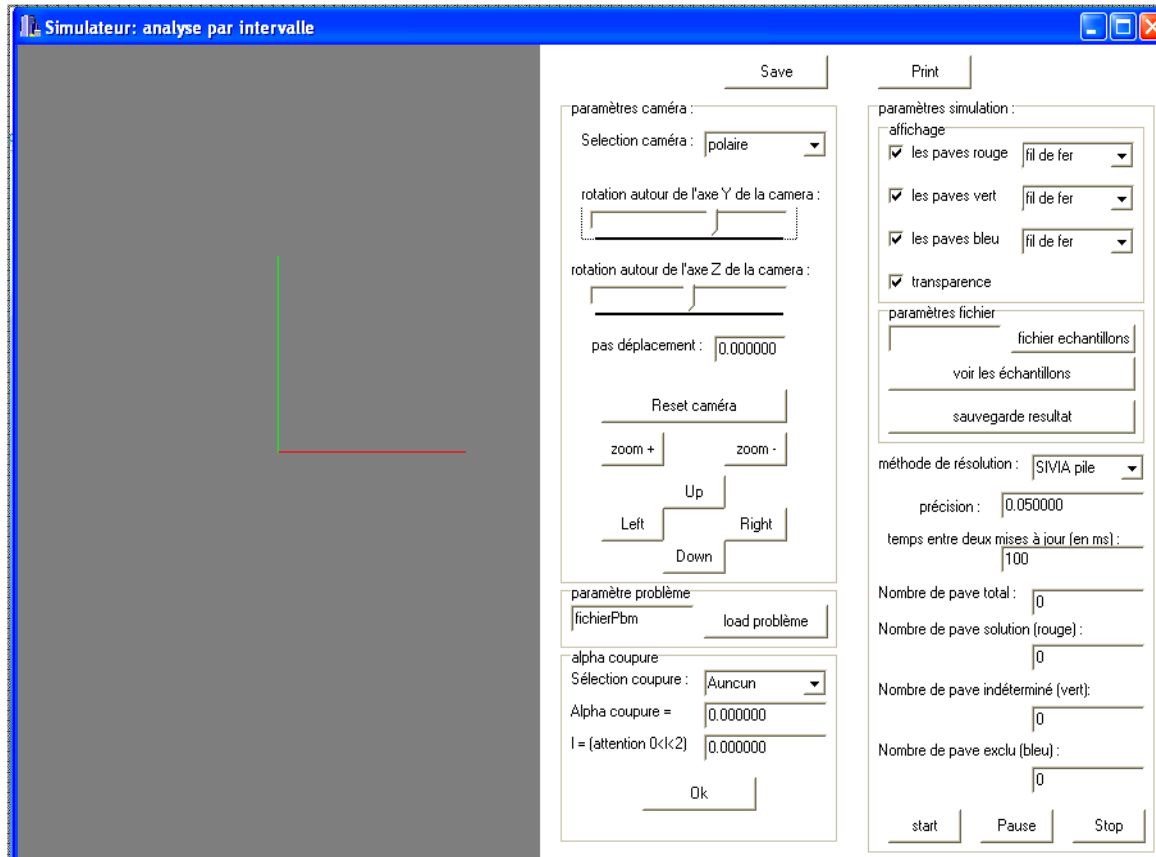
- [Neumaier,1991] A. Neumaier, *Interval Methods For Systems Of Equations*, Cambridge Univ Pr, 1991.
- [Neumaier, 1999] A. Neumaier, A simple derivation of the Hansen-Bliek-Rohn-Ning-Kearfott enclosure for linear interval equations, *Reliable Computing*, vol.52 no.5, pp.131-136, 1999.
- [Perko, 1991] L. Perko, *Differential equations and dynamical systems*, 3rd edition, text in Applied Mathematics 7, Springer, 1991.
- [Raissi et al., 2003] T.Raissi, N.Ramdani and Y.Candau, Guaranteed state estimation for nonlinear continuous time with Taylor models, 13th IFAC Symposium on system identification, Rotterdam, SYSID2003, pp.1725-1730, 2003.
- [Raissi et al., 2005] T. Raissi, N. Ramdani, Y. Candau, Bounded error moving horizon state estimation for non-linear continuous time systems: application to a bioprocess system, *Journal of Process Control*, vol.15 no.5, pp.537-545, 2005.
- [Rapaport et Gouzé, 2003] A. Rapaport, J-L Gouzé, Parallelotopic and practical observers for nonlinear uncertain systems, *International Journal Control*, vol.76 no.3, pp.237-251, 2003.
- [Rapaport et Dochain, 2005] A. Rapaport and D. Dochain, Interval Observers for Biochemical Processes with Uncertain Kinetics and Inputs, *Mathematical Biosciences* , vol.193 no.2, pp.235-253, 2005.
- [Sauvage et al., 2007] F.Sauvage, D.Dochain, M.Perré, State estimation within guaranteed interval, *Proceedings of the European Control Conference, Kos, Greece*, pp.5108-5114, 2007.
- [Schichl et Neumaier, 2005] H.Schichl and A.Neumaier, Interval Analysis on Directed Graphs for Global Optimization, *Journal of Global Optimization*, vol.33, pp.541-562, 2005.
- [Videau, 2005] G.Videau, *Méthodes garanties pour l'estimation d'état et le contrôle de cohérence des systèmes non linéaires à temps continu*, Phd Thesis, École doctorale des sciences physiques et de l'ingénieur de l'Université Bordeaux I, France, 2005.
- [Walter et Jaulin, 1994] E.Walter and L. Jaulin, Guaranteed characterization of stability domains via set inversion, *IEEE Transaction on Automatic Control*, vol.39 no.4, pp.886-889, 1994.

Annexe

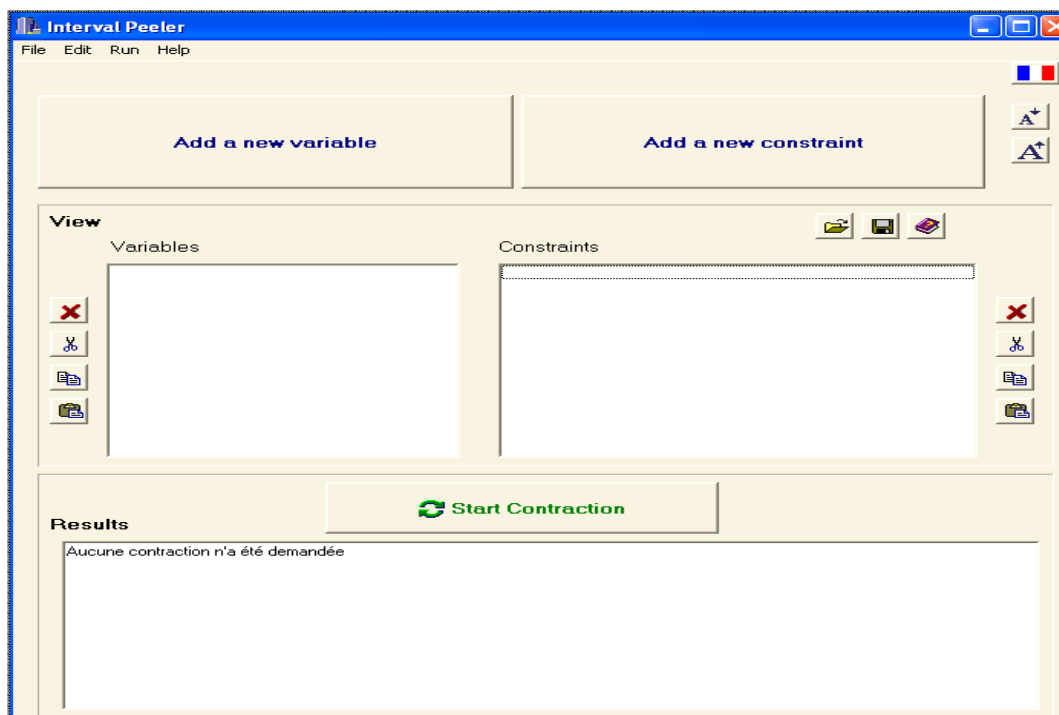
Annexe A

Logiciels utilisés :

« Simulateur : analyse par intervalle » réalisé par Luc Jaulin pour caractériser des ensembles en 3D , qui satisfait des inégalités non linéaires en utilisant l'algorithme SIVIA.



INTERVAL PEELER, pour l'utilisation de l'outil « contracteurs », réalisé par L.Jaulin, X. Baguenard et M.Dao.



INTLAB- INTerval LABoratory : toolbox de MATLAB pour la manipulation des intervalles, développée par Siegfried M. Rump, Insitute for Reliable Computing Hamburg University of Technology, Germany.

Annexe B

Fonctions MATLAB/INTLAB utilisées :

La fonction `kraw.m` pour la résolution des systèmes d'équations linéaires à intervalles par la méthode de Krawczyk

```
function x= kraw(A,b)
n = length(A);
midA = mid(A);
% Preconditions system.
C = inv(midA);
b = C*b;
CA = C*intval(A);
A = eye(n)-CA;
beta = norm(A,inf);
if beta >= 1;
disp('error Algorithm not suitable for this A')
end;
alpha = norm(b,inf)/(1-beta);
x(1:n,1) = intval( [ -alpha ; alpha ] );
s_old = inf;
s = sum(rad(x));
mult = (1+beta)/2;
while s < mult*s_old
x = intersect(b+A*x,x); % Krawczyk's iteration.
s_old = s;
s = sum(rad(x));
end
```

La fonction `hsolve.m` pour la résolution des systèmes d'équations linéaires à intervalles par la méthode de Hansen, Bliek, Rohn, Ning, Kearfott and Neumaier.

```

hsolve.m
function x = hsolve(A,b)
n = dim(A);
C = inv(mid(A));
A = C*A;
b = C*b;
dA = diag(A); % Diagonal entries of A.
A = compmat(A); % Comparison matrix.
B = inv(A);
v = abss(B*ones(n,1));
setround(-1)
u = A*v;
if ~all(min(u)>0) % Check positivity of u.
error('A is not an H-matrix')
else
dAc = diag(A);
A = A*B-eye(n); % A contains the residual matrix.
setround(1)
w = zeros(1,n);
for i=1:n,
w = max(w,(-A(i,:))/u(i));
end;
dlow = v.*w'-diag(B);
dlow = -dlow;
B = B+v*w; % Rigorous upper bound for exact B.
u = B*abss(b);
d = diag(B);
alpha = dAc+(-1)./d;
beta = u./dlow-abss(b);
x = (b+midrad(0,beta))./(dA+midrad(0,alpha));
end

```