



Titre: Identification et caractérisation des éléments manuscrits d'une
Title: signature

Auteur: Justin Picard
Author:

Date: 1997

Type: Mémoire ou thèse / Dissertation or Thesis

Référence: Picard, J. (1997). Identification et caractérisation des éléments manuscrits d'une
Citation: signature [Mémoire de maîtrise, École Polytechnique de Montréal]. PolyPublie.
<https://publications.polymtl.ca/6733/>

 Document en libre accès dans PolyPublie
Open Access document in PolyPublie

URL de PolyPublie: <https://publications.polymtl.ca/6733/>
PolyPublie URL:

**Directeurs de
recherche:** Jean-Jules Brault
Advisors:

Programme: Non spécifié
Program:

UNIVERSITÉ DE MONTRÉAL

IDENTIFICATION ET CARACTÉRISATION
DES ÉLÉMENTS MANUSCRITS D'UNE SIGNATURE

JUSTIN PICARD

DÉPARTEMENT DE GÉNIE ÉLECTRIQUE ET DE GÉNIE INFORMATIQUE
ÉCOLE POLYTECHNIQUE DE MONTRÉAL

MÉMOIRE PRÉSENTÉ EN VUE DE L'OBTENTION DU DIPLÔME DE
MAITRISE ÈS SCIENCES APPLIQUÉES (M.SC.A.)
(GÉNIE ÉLECTRIQUE)

AOÛT 1997

© Justin Picard, 1997.



National Library
of Canada

Acquisitions and
Bibliographic Services

395 Wellington Street
Ottawa ON K1A 0N4
Canada

Bibliothèque nationale
du Canada

Acquisitions et
services bibliographiques

395, rue Wellington
Ottawa ON K1A 0N4
Canada

Your file Votre référence

Our file Notre référence

The author has granted a non-exclusive licence allowing the National Library of Canada to reproduce, loan, distribute or sell copies of this thesis in microform, paper or electronic formats.

The author retains ownership of the copyright in this thesis. Neither the thesis nor substantial extracts from it may be printed or otherwise reproduced without the author's permission.

L'auteur a accordé une licence non exclusive permettant à la Bibliothèque nationale du Canada de reproduire, prêter, distribuer ou vendre des copies de cette thèse sous la forme de microfiche/film, de reproduction sur papier ou sur format électronique.

L'auteur conserve la propriété du droit d'auteur qui protège cette thèse. Ni la thèse ni des extraits substantiels de celle-ci ne doivent être imprimés ou autrement reproduits sans son autorisation.

0-612-33176-8

UNIVERSITÉ DE MONTREAL

ÉCOLE POLYTECHNIQUE DE MONTRÉAL

Ce mémoire intitulé :

IDENTIFICATION ET CARACTÉRISATION
DES ÉLÉMENTS MANUSCRITS D'UNE SIGNATURE

présenté par : JUSTIN PICARD

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

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

M. HERVÉ Jean-Yves, Ph.D., président

M. BRAULT Jean-Jules, Ph.D., membre et directeur de recherche

M. ELIE Daniel, Ph.D., membre

REMERCIEMENTS

Avant tout, je dois remercier mon directeur de recherche Jean-Jules Brault pour les excellents rapports que l'on a conservés tout au long de ma thèse, pour les après-midi complets passés à discuter dans son bureau, pour ses commentaires et critiques bien placés qui m'ont été très utiles.

Je remercie vivement les membres du jury, M. Jean-Yves Hervé et M. Daniel Elie, d'avoir pris le temps de lire ce mémoire.

Merci aux membres du laboratoire Scribens, toujours prêts à m'aider en cas de problème.

Je tiens finalement à remercier toutes les personnes qui ont participé à l'expérience, et en particulier aux faussaires, Lydie, Christine et Olivier, en espérant qu'ils n'auront pas l'idée de mettre à profit leur talent certain.

RÉSUMÉ

Les systèmes de vérification de signature (VdS) ont pour but de déterminer l'origine (authentique ou non-authentique) d'une signature à partir du signal codé de la signature originale. En VdS, on distingue les systèmes de type dynamique et statique, selon la méthode d'acquisition de la signature originale. Ce mémoire porte sur les systèmes dynamiques, dans lesquels on code la séquence de mouvements génératrice du tracé de signature en conservant l'information de vitesse. Les recherches menées pour améliorer les performances des systèmes de VdS, c'est-à-dire pour minimiser le taux d'erreur de reconnaissance, sont actuellement d'un grand intérêt pratique et théorique.

En reconnaissance de formes (RF), on classe une forme d'origine inconnue parmi un ensemble de classes, en s'appuyant sur une certaine connaissance de chacune des classes. La VdS n'est pas un problème classique de RF dans la mesure où la classe des signatures non-authentiques n'existe pas vraiment, donc ne peut être « connue » par le système. Ce problème est fondamental car, en théorie, on ne peut classer adéquatement une forme parmi deux classes sans une connaissance minimale de ces deux classes.

Un système de VdS classe des objets très similaires, les signatures authentiques et les imitations. Pour distinguer des objets très similaires, on doit rechercher des différences locales au niveau des sous-objets qui constituent la structure de ces objets, ce qui nécessite une étape préliminaire d'identification de ces sous-objets. En VdS, cette étape, appelée **identification des éléments manuscrits**, est très complexe. L'étape

suivante est de trouver une représentation de l'information contenue dans les éléments manuscrits d'une signature, qui permette la meilleure discrimination possible entre signatures authentiques et non-authentiques. Cette étape est appelée **caractérisation des éléments manuscrits**.

Ce mémoire propose une méthode très robuste d'identification des éléments manuscrits. Cette méthode est basée sur l'alignement élastique, déjà utilisé en VdS pour le calcul d'une distance entre deux signatures. L'algorithme d'alignement élastique recherche la séquence de comparaison des éléments manuscrits minimisant un coût total de comparaison. On démontre que cette séquence de comparaison peut être utilisée pour identifier les éléments manuscrits de façon optimale selon un certain critère. La méthode développée peut également servir à segmenter des signatures d'une manière concertée, évitant les problèmes habituels d'instabilité de la segmentation.

Une modélisation statistique de la VdS, qui s'appuie essentiellement sur l'hypothèse que les distributions de probabilité des deux classes respectent une loi normale, permet d'approcher une connaissance de la classe des signatures non-authentiques. On en déduit de nombreux résultats théoriques intéressants, notamment sur la variation du taux d'erreur d'un système de VdS en fonction de certains paramètres d'un scripteur. Le modèle peut aussi être utilisé comme méthode de génération de signatures authentiques et d'imitations.

Une étude expérimentale a été faite sur 480 signatures, soit huit scripteurs qui ont produit 30 exemplaires de leur signature et 3 imitateurs qui ont fourni 10 imitations par

scripteur. Ces imitations devaient être acceptées par un système de VdS breveté. Le taux d'erreur obtenu par le système est de 4.17%, contre 10.21% pour la méthode d'alignement élastique classique. Si on choisit un seuil de décision personnalisé pour le système de VdS, ce taux passe à 1.04%.

De cette expérience, des conclusions importantes sont tirées sur la caractérisation adéquate d'une signature. Notamment, on conclut que la vitesse horizontale locale est l'information la plus discriminante, et que la segmentation doit autant que possible être évitée car elle nuit à la reconnaissance.

ABSTRACT

Automatic signature verification (ASV) systems is concerned with determining the authenticity of a signature from its encoded signal. Dynamic and static ASV systems are distinguished by the way original signatures are acquired. This research is concerned with dynamic systems, where the sequence of moves of the hand generating the signature is encoded. Research aimed at improving SV error rate is of a great practical and theoretical interest.

In pattern recognition, a feature of unknown origin must be classified in a set of classes, based on a certain knowledge of these classes. ASV is not a classical problem of pattern recognition since the class of signatures that are not authentic (forgeries and others) does not really exist, or at least there is no suitable training set for that class. This is a fundamental problem since it is theoretically not possible to classify a feature between two classes without a minimal knowledge of each class.

ASV systems classify authentic signatures and forgeries, which are very similar objects. To classify very similar objects, differences must be found in the local features that constitute the structure of the objects. For this, a preliminary step of identifying local features of signatures must be done. The next step is to find a proper representation of the local information contained in a signature. This representation aims at an optimal classification of authentic signatures and forgeries.

This research proposes a robust method to identify local handwritten features. It is based on Dynamic Time Warping, a technique normally used in ASV to calculate a distance between two signatures. The algorithm searches the sequence of comparisons that minimizes a global cost of comparison. Based on a certain criteria, we show that this method can be used to identify optimally local features. The method can also be used as a segmentation method that avoids the classical problem of variability of segmentation.

A statistical model of SV, mostly based on the hypothesis that the probability distribution of the classes have a normal density, is proposed. It improves knowledge of the forgery class. Theoretical results are deducted on the error rate of SV system, which is shown to depend on the duration and stability of the signer. Based on the probability distribution, a method of generating authentic signatures and forgeries is developed.

An experience was done on 480 signatures, that is eight different signers with 30 signatures each, and three imitators with 10 forgeries for each signer. Error rates achieved with our system is 4.17%, compared to 10.21% with Dynamic Time Warping. With a personalized decision rule, the error rates is 1.04%.

From this experience, important deductions are done on the adequate way to represent local features of signatures. Notably, we conclude that horizontal speed is the most discriminant information, and that segmentation must be avoided if possible, since it has negative effects on signature identification.

TABLE DES MATIÈRES

REMERCIEMENT.....	iv
RÉSUMÉ.....	v
ABSTRACT.....	viii
TABLE DES MATIÈRES.....	x
LISTE DES TABLEAUX.....	xvi
LISTE DES FIGURES.....	xvii
LISTE DES SYMBOLES.....	xx
1. INTRODUCTION.....	1
1.1 Introduction à la VdS.....	1
1.1.1 Présentation générale des systèmes de VdS	1
1.1.2 Composants et sous-systèmes d'un système de VdS.....	5
<i>1.1.2.1 L'acquisition des signatures.....</i>	<i>7</i>
<i>1.1.2.2 Prétraitements des signatures.....</i>	<i>8</i>
<i>1.1.2.3 L'analyse des signatures.....</i>	<i>11</i>
<i>1.1.2.4 L'apprentissage</i>	<i>11</i>
<i>1.1.2.5 Comparaison et décision</i>	<i>12</i>
1.1.3 Quelques systèmes de VdS dynamique rencontrés dans la littérature.	12
1.1.4 Problèmes inhérents à la VdS	14
<i>1.1.4.1 Représentation des signatures</i>	<i>14</i>

1.1.4.2	<i>Quelles sont les classes ?</i>	15
1.2	Objectifs du mémoire	16
1.3	sommaire de la méthode utilisée	19
1.4	Plan du mémoire	22
2.	L'IDENTIFICATION DES ÉLÉMENTS MANUSCRITS .	24
2.1	Introduction	24
2.1.1	But et outil de base	24
2.1.2	Plan du chapitre	25
2.2	Synchronisation de deux signatures	27
2.2.1	But et problématique de la synchronisation de deux signatures.....	27
2.2.2	Identification des éléments manuscrits des signatures: une revue de la littérature	34
2.2.3	Présentation informelle de l'alignement élastique	35
2.2.4	Présentation formelle de l'alignement élastique pour la VdS	41
2.2.5	Calcul d'une distance normalisée entre deux signatures	45
2.2.6	Justification de l'utilisation de l'alignement élastique comme méthode d'identification des éléments manuscrits	45
2.3	Synchronisation de N signatures	47
2.3.1	But et problématique de la synchronisation de N signatures.....	47
2.3.2	Description de la méthode	48
2.3.3	Choix d'une signature comme espace temporel de projection.....	49

2.3.4	Détermination de la signature moyenne et resynchronisation	50
2.4	Segmentation	54
2.4.1	Présentation de la segmentation de séquences.....	55
2.4.1.1	<i>Pourquoi segmenter une signature?</i>	55
2.4.1.2	<i>Quelques méthodes connues de segmentation</i>	55
2.4.2	Outils de base de la méthode de segmentation	57
2.4.2.1	<i>Les minima de vitesse</i>	58
2.4.2.2	<i>Le filtre gaussien</i>	60
2.4.3	Les dangers de la pré-segmentation.....	63
2.4.3.1	<i>Problème d'échelle</i>	65
2.4.4	Approche retenue: la post-segmentation	69
2.5	Conclusion	71
3.	MODÉLISATION STATISTIQUE DE LA VDS	72
3.1	Introduction	72
3.2	Développement du modèle	73
3.2.1	Approches possibles	73
3.2.2	L'approche statistique à la VdS.....	74
3.2.3	Estimation de $P(\omega_1)$	76
3.2.4	Estimation de $P(X \omega_1)$	77
3.2.5	Estimation de $P(X \omega_2)$	87
3.3	Résultats théoriques tirés du modèle	97

3.3.1	Estimation du taux d'erreur de reconnaissance d'un scripteur par le système.....	97
3.3.2	Difficulté d'imitation d'un scripteur.....	102
3.3.3	Qualité d'un système de VdS	104
3.4	Conclusion	107
4.	CONCEPTION GÉNÉRALE DU SYSTÈME DE VDS ET MÉTHODOLOGIE D'ÉVALUATION	109
4.1	Introduction	109
4.2	Conception du système de VdS.....	111
4.2.1	Plan de la section	111
4.2.2	Acquisition des signatures	112
4.2.3	Prétraitements des données.....	113
4.2.3.1	<i>Filtrage des signatures</i>	<i>113</i>
4.2.3.2	<i>Normalisation des signatures</i>	<i>114</i>
4.2.3.3	<i>Synchronisation</i>	<i>116</i>
4.2.3.4	<i>Segmentation des signatures.....</i>	<i>117</i>
4.2.4	Analyse	117
4.2.5	Apprentissage	120
4.2.6	Comparaison et décision.....	122
4.2.6.1	<i>Comparaison par mesure d'un distance.....</i>	<i>122</i>

4.2.6.2	<i>Prise de décision à partir de la mesure de distance</i>	
	à la classe ω_1	127
4.3	Méthodologie d'évaluation	130
4.3.1	Introduction	130
4.3.2	Acquisition des signatures	131
4.3.2.1	<i>Principes respectés lors de l'acquisition des données</i>	131
4.3.2.2	<i>Établissement de classes de signatures</i>	132
4.3.2.3	<i>Acquisition des signatures authentiques</i>	136
4.3.2.4	<i>Acquisition des imitations</i>	137
4.3.2.5	<i>Analyses humaines</i>	139
4.3.3	Évaluation des performances d'un système	140
4.3.4	Validité des résultats obtenus et comparaison de systèmes	144
4.4	Conclusion	147
5.	ANALYSE DES RÉSULTATS	148
5.1	Introduction	148
5.2	Résultats de référence	149
5.2.1	Alignement élastique classique	149
5.2.2	Résultats des analystes humains	152
5.3	Comparaison des résultats pour différents classificateurs	153
5.3.1	Influence des prétraitements	154
5.3.1.1	<i>Effet d'une resynchronisation</i>	154

5.3.1.2	<i>Segments et points</i>	155
5.3.2	Choix des caractéristiques	159
5.3.2.1	<i>Comparaison des différentes caractéristiques locales</i>	159
5.3.2.2	<i>Combinaison de caractéristiques</i>	165
5.3.3	Mesure de distance	167
5.3.4	Conclusion	168
5.4	Vérification expérimentale des hypothèses du modèle de système	
	de VdS	169
5.4.1	Traitement des portions non-identifiées	169
5.4.2	Détermination du seuil.....	172
5.4.3	Détermination du taux d'erreur en utilisation normale :	175
5.5	Conclusion	175
6.	CONCLUSION	177
6.1	Contributions	177
6.2	Problèmes non-résolus et recherche future	179
	ANNEXE	191

LISTE DES TABLEAUX

Tableau 4-1 : Typologie des scripteurs	135
Tableau 5-1 : Taux d'erreur de la méthode d'alignement élastique.....	150
Tableau 5-2 : Comparaison des taux d'erreurs avec et sans resynchronisation.....	155
Tableau 5-3 : Taux d'erreur obtenue par la représentation par segments	157
Tableau 5-4 : Taux d'erreur (%) par coordonnée et type d'information	160
Tableau 5-5 : Taux d'erreur pour des signatures caractérisées par leurs éléments non identifiés.....	170

LISTE DES FIGURES

Figure 1-1: Prototype d'un système de VdS.....	2
Figure 1-1: Schéma général d'un système de reconnaissance de forme	6
Figure 1-1 : Signaux $x(t)$ et $y(t)$ d'une signature	8
Figure 1-1: le filtrage	9
Figure 1-2: Segmentation d'une signature.....	10
Figure 1-1 : En vous inspirant de la signature authentique située en haut, trouvez laquelle des signatures (a) et (b) est une imitation.	17
Figure 1-1 : Prototype du système de RF élaboré dans cette recherche.....	20
Figure 2-1 : Points de même nature sur deux signatures	28
Figure 2-2 : La distorsion.....	30
Figure 2-3 - La variabilité	31
Figure 2-4 : L'inconstance	32
Figure 2-5 : Inconstances d'un signataire	32
Figure 2-6 : L'élément manuscrit, un point ou un segment ?	33
Figure 2-1 : Visualisation de certains points synchronisés par alignement élastique. Les points de même numéro sont de même nature	40
Figure 2-1 : Deux signatures d'un scripteur instable et la politique de comparaison.	44
Figure 2-1 : 15 des 30 Signatures d'un scripteur instable et inconstant, et la	

signature moyenne qui en résulte.....	53
Figure 2-2 : Signature moyenne avec écart-type représenté	54
Figure 2-1 : Signal de vitesse non-filtré (a), même signal filtré (b) et (c)	61
Figure 2-1 : Instabilité de la segmentation sur deux signatures très semblables.	
La flèche indique qu'un point de segmentation est absent sur (a)	
et présent sur (b).....	64
Figure 2-1 : Où segmenter ? Une question de point de vue.	67
Figure 2-2 : Points de segmentation à différentes échelles	68
Figure 2-1 : Segmentation à $\sigma=3$ par la méthode de post-segmentation.....	70
Figure 2-1 : Prototype de l'apprentissage pour l'identification des éléments	
manuscrits	71
Figure 3-1 : Variation dans le début de la signature d'un scripteur	78
Figure 3-2 : Distribution de caractéristiques locale v_x et v_y et loi normale.	81
Figure 3-3 : Corrélation selon la distance entre points, pour le scripteur no 3.	83
Figure 3-4 : Corrélation de position (p), vitesse (v) et accélération (a)	85
Figure 3-5 : 4 signatures générées aléatoirement et 4 signatures authentiques.....	86
Figure 3-1 : Distribution des caractéristiques locales v_x et loi normale.....	89
Figure 3-2 : A gauche, classe ω_2 polarisée. A droite, classe ω_2 non-polarisée.	91
Figure 3-3 : Position des signatures hybrides par rapport au scripteur	92
Figure 3-4 : Signatures hybrides de deux scripteurs	94
Figure 3-5 : Imitations « originales » et imitations générées	96

Figure 3-1 : Évolution de pmc en fonction de n pour différentes valeurs de α	102
Figure 4-1 : Étapes du système de VdS et plan de la section.....	111
Figure 4-1 : Exemple de la signature de chacun des 8 scripteurs	136
Figure 4-1: Exemple d'un scripteur difficile à imiter.....	139
Figure 4-1 : Évolution des taux d'erreur en fonction du seuil	142
Figure 5-1 : Deux scripteurs problématiques.....	151
Figure 5-1 : Caractéristiques extraites des segments	156
Figure 5-1 : Direction de l'accélération sur le tracé manuscrit.....	162
Figure 5-2 : Le résidu inconscient.....	164
Figure 5-1 :Taux d'erreur pour la représentation $[V_x, V_y]$ des signatures	166
Figure 5-2 : Taux d'erreur pour la représentation $[V_x, V_\alpha]$ des signatures	166
Figure 5-1 :Variation du taux d'erreur en fonction du seuil	173
Figure 5-2:Variation du taux d'erreur en fonction du seuil pour l'alignement élastique	174
Figure 6-1 : Signature moyenne et rapport des écarts-types locaux	183

LISTE DES SYMBOLES

X	Vecteur de représentation d'une signature
ω_1	Classe authentique
ω_2	Classe non-authentique
x, y, r, α	Représentations horizontale, verticale, radiale et angulaire de l'information locale
V, A, P	Information locale de vitesse, accélération et position
V_x	Signature représentée par les caractéristiques locales de vitesse horizontale ; de même pour $V_y, V_r, V_\alpha, A_x, A_y, A_r, A_\alpha, P_x, P_y, P_r, P_\alpha,$
n	Nombre d'éléments manuscrits représentant la signature d'un scripteur
N	Nombre de signatures de référence

Note : il arrive que N ou n prenne momentanément une autre signification.

1. INTRODUCTION

Ce chapitre commence par une introduction à la vérification de signatures (VdS) dans le but d'établir un contexte pour la recherche décrite ici (section 1.1). On présente ensuite les objectifs de cette recherche (1.2), et le sommaire de la méthode utilisée pour les atteindre (1.3). On termine en présentant un plan général du mémoire (1.4).

1.1 Introduction à la VdS

1.1.1 Présentation générale des systèmes de VdS

L'expertise d'une signature permet d'en vérifier l'authenticité dans des cas litigieux. Elle est faite par un expert en analyse de signature qui en fait une étude approfondie. Si l'expertise est très coûteuse et ne se justifie que pour des cas où l'enjeu est important (par ex. validité d'un testament, d'un contrat ou d'un chèque à montant élevé), la décision de l'expert est très fiable et possède une valeur juridique.

Les systèmes de VdS sont une tentative d'automatisation du processus d'expertise des signatures, et ne demandent habituellement pas d'intervention humaine lors de leur utilisation. Lors de la vérification d'une signature, le système commence par en faire l'acquisition, sous forme dynamique ou statique (ces termes seront éclaircis en prochaine sous-section). L'information contenue dans la signature subit ensuite de multiples

traitements à travers les composants du système, dont on parlera en détail à la prochaine sous-section. La sortie du système est une décision binaire: la signature est acceptée (authentique) ou rejetée (imitation ou origine inconnue).

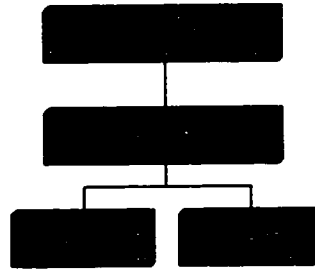


Figure 1-1: Prototype d'un système de VdS

L'émergence et le développement actuel des systèmes de VdS est essentiellement dû aux trois facteurs suivants:

- Le développement de technologies permettant l'acquisition (ou le codage) de signatures, soit par lecture optique, soit par enregistrement sur une tablette graphiques, à l'aide d'un stylo électromagnétique.
- La puissance de calcul des microprocesseurs actuels permet d'effectuer les différents traitements de l'information en un laps de temps suffisamment court, et la taille des mémoires informatiques permet le stockage d'information en grande quantité (pour les paramètres décrivant les différents signataires utilisant le système, par exemple).
- Les techniques développées pour le traitement du signal (filtrage, transformée de Fourier), les méthodes de l'intelligence artificielle et de la reconnaissance de formes (RF) profitent énormément à la VdS. Comme la signature présente des similitudes avec différents types de signaux 1-D et 2-D, ainsi qu'avec les séquences, on peut

adapter certaines méthodes élaborées initialement pour d'autres problèmes. La VdS profite donc de la recherche effectuée dans des domaines connexes.

Si les systèmes de VdS n'ont pas encore suffisamment fait leur preuve pour qu'on accorde à leur décision une valeur juridique, il est probable que l'on fasse de plus en plus appel à eux dans les prochaines années pour la vérification d'identité. En effet, la signature est selon plusieurs auteurs un des meilleurs moyens d'identification d'une personne [Raphael 74, Worthington 85].

Il y a relativement peu d'espoir que les signatures codées par lecture optique permettent des taux de reconnaissance équivalents à ceux des experts humains, en tout cas dans les prochaines années. Par contre, les signatures enregistrées sur des tablettes numérisées, appelées signatures dynamiques, ont un grand potentiel : elles donnent l'accès à des données séquentielles plus faciles à analyser que des images, et permettent d'extraire une information non accessible à l'expert humain et au faussaire qui n'ont en théorie accès qu'à l'image. La vitesse en chaque point du tracé, la pression du stylo sur la tablette et son orientation dans l'espace sont autant d'informations supplémentaires caractéristiques d'un signataire, accessibles par les tablettes graphiques et des stylos spécialisés. Par une utilisation adéquate de ces informations, les systèmes de VdS dynamique pourraient atteindre ou même dépasser les performances des experts dans les prochaines années. **Ce mémoire est consacré à l'étude des systèmes dynamiques.**

La technologie actuelle permet l'acquisition et la détermination de l'origine d'une signature en des temps assez court. Cela rend possible des applications à grande échelle

de la VdS : vérification systématique de chèques, vérification immédiate de la signature lors de l'utilisation d'une carte de crédit (il faudrait pour cela la présence de systèmes d'acquisition dans les commerces). Ces applications n'ont pas à détecter toutes les fausses signatures pour être rentables, mais simplement une bonne part d'entre elles (ordre de grandeur de 80% de détection des faux pour les cartes de crédit, selon [Nelson 94]) pour éviter autant de transactions illégales. Le coût d'achat et d'installation des dispositifs de VdS serait compensé par les économies faites sur les imitations détectées.

L'avènement d'une société où les rapports humains directs seront de plus en plus remplacés par des rapports humains à distance et par des rapports homme-machine nécessite l'implantation de techniques de vérification d'identité des individus. Le code d'accès ou le mot de passe sont les moyens les plus utilisés actuellement, mais ils ne peuvent être utilisés pour toutes les vérifications d'identité. La signature, dont l'utilisation est entrée dans les mœurs et qui possède pour chaque personne des caractéristiques uniques tant au niveau de l'image que de la génération du tracé, pourrait devenir un des moyens d'identification de l'avenir. L'accès à certains locaux, la signature de contrats à distance sont deux des multiples exemples d'utilisation possible de la VdS.

L'utilisation des systèmes de VdS se justifie complètement en théorie. Leur utilisation massive ou leur oubli au profit d'autres technologies de l'identification personnelle (empreintes digitales, examen de l'iris de l'œil) dépend en bonne partie de leurs performances, évaluables en termes de coût d'implémentation, de vitesse

d'exécution et surtout de taux d'erreur de reconnaissance. **Ce mémoire a pour but ultime d'améliorer ce dernier aspect, essentiel aux systèmes de VdS.**

1.1.2 Composants et sous-systèmes d'un système de VdS

La VdS est, au même titre que le traitement et analyse d'images et la reconnaissance de la parole, un sujet de la reconnaissance de formes (RF). La RF est une activité fondamentale à l'être humain (et l'animal) qui lui permet de se repérer en identifiant des signaux du monde extérieur à des situations antérieures. La tentative de reproduction des activités perceptuelles humaines par des machines (reconnaissance de la parole, de l'écriture, d'objets dans des images, etc.) avec l'avènement de l'informatique a donné lieu à la naissance du domaine de la reconnaissance de forme. Dans ce domaine, la démarche classique consiste à séparer un système de reconnaissance en les étapes suivantes :

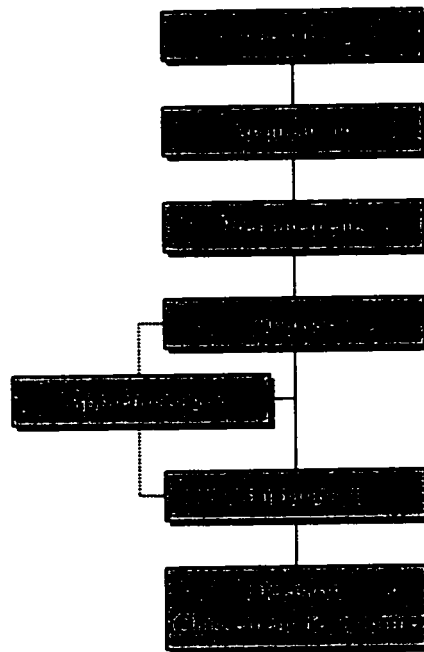


Figure 1-2: Schéma général d'un système de reconnaissance de forme

Dans le cas de la VdS, le système doit déterminer si une signature est authentique ou non en transformant l'information initiale (soit la signature codée lors de l'étape d'acquisition) au cours des différentes étapes jusqu'à l'étape de décision. Nous allons maintenant expliquer le rôle de chaque étape d'un système de RF, en nous focalisant sur les systèmes de VdS. Évidemment, chaque étape dépend de la précédente : par exemple, la signature est codée selon son image, les prétraitements seront adaptés à la nature bidimensionnelle des données.

Ce schéma général est une sorte d'archétype des systèmes de RF, qui ne suivent pas tous exactement la même procédure. Dans notre cas, l'étape de prétraitement inclut un sous-système de RF servant à l'identification des éléments manuscrits de la signature.

1.1.2.1 L'acquisition des signatures

Dans le monde physique, les signatures sont des images statiques mais aussi une séquence de mouvements de la main génératrice du tracé. Selon la méthode d'acquisition des signatures, on parle de systèmes statiques ou dynamiques.

Dans le cas des systèmes statiques, les données sont les images numérisées des signatures, enregistrées par exemple sur des chèques. Les systèmes de VdS statiques n'ont pas accès à une information directe sur la séquence de mouvements qui a généré le tracé, ni sur la vitesse à laquelle chacune des portions du tracé a été effectuée.

Pour les systèmes dynamiques, l'acquisition des données se fait au moment de la signature, qui est exécutée en général sur une tablette numérique à l'aide d'un crayon électromagnétique. Les systèmes dynamiques traitent des signatures dites « dynamiques » car ils bénéficient des informations temporelles. Les systèmes de vérification dynamique ont accès à davantage d'information que les systèmes statique (vitesse, accélération, durée) il leur est plus facile d'extraire l'information pertinente (ordre dans lequel a été exécuté le tracé).

En VdS, on distingue totalement les deux types de systèmes, qui sont utilisés dans des cas différents et font appel à des techniques très différentes. Les signatures statiques se rapprochent du traitement et analyse d'images, alors que les signatures dynamiques sont plus proches des séquences et des signaux 1-D comme le signal de parole.

Une signature est un signal échantillonné à une fréquence en général autour de 100 Hz. Elle se présente sous forme d'une séquence de positions $\{x(t), y(t)\}$, auquel on ajoute d'autres types d'information tel que la pression.

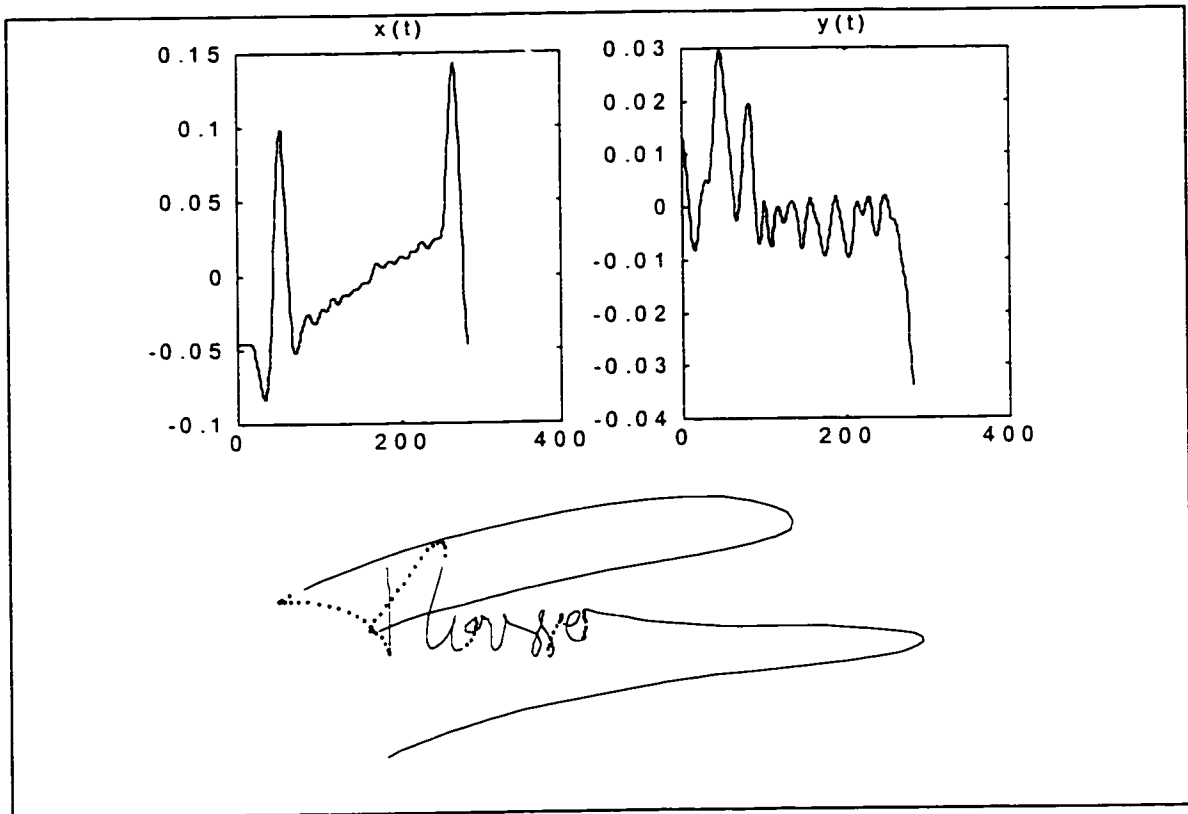


Figure 1-3 : Signaux $x(t)$ et $y(t)$ d'une signature

Ce mémoire porte sur les signatures dynamiques uniquement. Ce qui suit ne concernera plus que les signatures dynamiques.

1.1.2.2 Prétraitements des signatures

Après l'étape d'acquisition des signatures dynamiques, celles-ci se présentent sous forme d'un signal échantillonné dans le temps. Ce signal contient des informations jugées

non-nécessaires ou même nuisibles à la reconnaissance. Les prétraitements ont pour but d'éliminer ces informations. Parmi les opérations de prétraitements, on retrouve :

- **Le filtrage** : les bruits dus aux tremblements de la main, aux frottements du stylo ou à d'autres causes extérieures sont généralement localisés dans les hautes fréquences. On les réduit par l'application d'un filtre passe-bas.

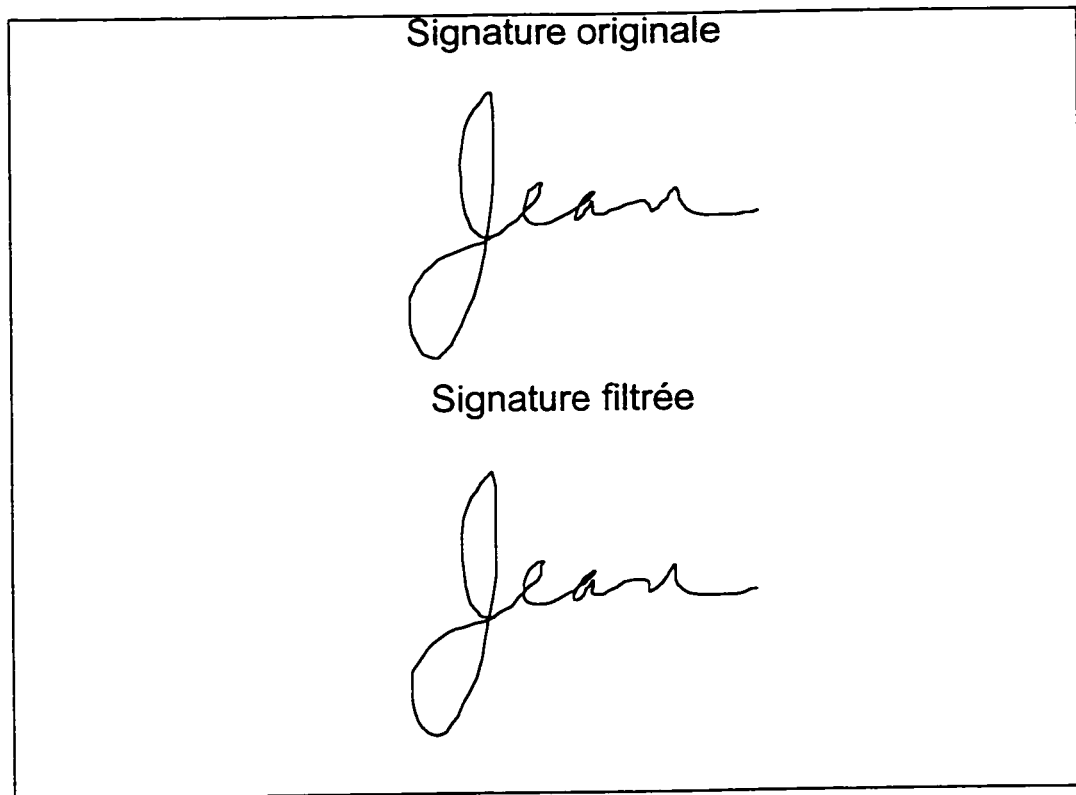


Figure 1-4: le filtrage

- **La normalisation** : les signatures d'un même scripteur pouvant présenter des différences de taille notables et peu significative peuvent être normalisés, par exemple selon la longueur du tracé ou l'énergie totale contenue dans le signal.

- **La segmentation** : une signature échantillonnée à 100 Hz et d'une durée typique de 3 secondes contient 300 points. Afin d'éliminer la redondance ou de réduire la dimension de l'espace de représentation des données, on peut segmenter la signature en portions de taille égale ou en segments représentatifs du tracé (parties du tracé sans changement de direction majeur). La segmentation est un traitement assez risqué car elle produit un résultat généralement instable, i.e. deux signatures similaires ne produiront pas les mêmes séquences de segments. Cette instabilité aura des répercussions négatives sur les étapes suivantes de la RF. Ce mémoire propose une méthode de segmentation comportant des risques d'erreur très faibles.

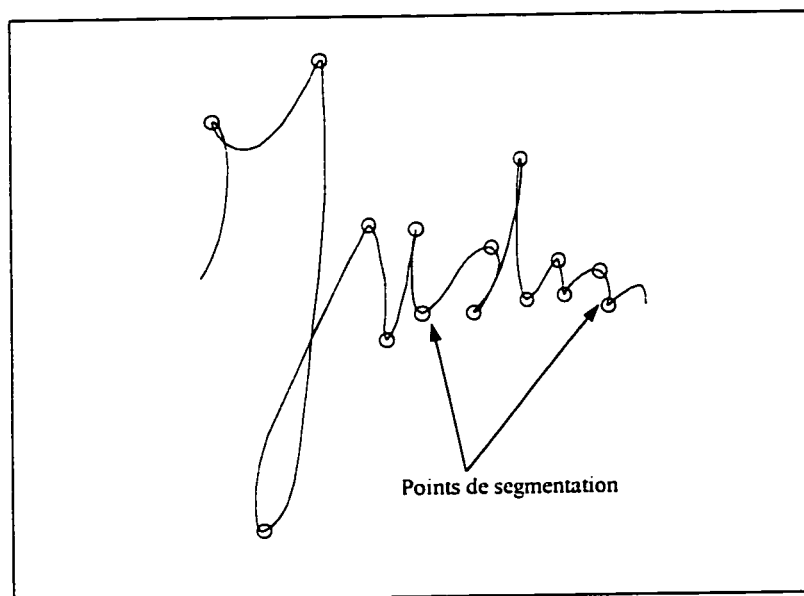


Figure 1-5: Segmentation d'une signature

1.1.2.3 L'analyse des signatures

L'étape d'analyse sert à calculer un certain nombre de caractéristiques ou paramètres décrivant la signature. Ces caractéristiques peuvent être globales (vitesse moyenne, durée de la signature) ou locales. Dans ce dernier cas, les caractéristiques peuvent porter sur les points échantillonnés ou sur des segments du tracé. Cette étape est évidemment déterminante car sans une description adéquate des signatures, on ne peut avoir un bon taux de reconnaissance. Un des apports les plus importants de ce mémoire est justement de proposer une description des signatures par ses caractéristiques locales, plus adéquates que les caractéristiques globales pour la reconnaissance.

1.1.2.4 L'apprentissage

Dans cette étape, le système de RF ou de VdS construit sa base de connaissance modélisant les différentes classes de formes qu'il aura à reconnaître. Dans le cas des signatures, il n'y a qu'une seule classe connue, la classe des signatures authentiques. En fait, on doit plutôt parler de méta-classe, car certains scripteurs signent de différentes façons, et une classe générale ne pourrait rendre compte fidèlement des différents types de signature de ces scripteurs. L'apprentissage est supervisé dans la mesure où le système connaît la classe de la forme servant à l'apprentissage. Il se fait lors d'une phase préliminaire où le signataire donne un échantillon aussi représentatif que possible de ses signatures. Il peut se poursuivre lors de l'utilisation ultérieure du système, ce qui permet de tenir compte des évolutions possibles de la signature dans le temps.

Notons que la phase d'apprentissage est facultative : le système peut conserver les signatures d'apprentissage tel quel. On parle alors d'approche « par prototype ».

1.1.2.5 Comparaison et décision

Le rôle de cette étape est de déterminer si une signature-test d'origine inconnue est authentique ou non, à partir des caractéristiques extraites dans l'étape d'analyse. D'abord une comparaison est faite entre la signature-test et les différentes signatures de référence (approche par prototype) ou, après apprentissage, entre la signature-test et une forme représentative de la classe. Le résultat de la comparaison est généralement une mesure de distance, de similarité ou de probabilité, à partir de laquelle une décision sera prise. La décision est souvent basée sur un seuil, qui est ajusté afin de minimiser le risque d'erreur, ou plus généralement l'espérance du coût, à chaque erreur étant rattaché un coût. Dans le cas de la VdS, il y a deux types d'erreurs : rejeter une signature authentique (type I) ou accepter une imitation (type II).

1.1.3 Quelques systèmes de VdS dynamique rencontrés dans la littérature

Les quelques systèmes de VdS dynamique suivants sont assez représentatifs de la recherche qui se fait dans le domaine. Notre approche s'inspire en partie de certains d'entre eux :

- **Système statistique (ou décisionnel) de Nelson [Nelson 94]** : ce système, orienté vers une utilisation commerciale pour cartes de crédit, extrait un ensemble de caractéristiques globales de la signature, très rapides à calculer (vitesse moyenne,

centre de gravité..). La décision, et donc la fixation du seuil de décision, se fait selon les règles de la statistique bayésienne. Pour que le système soit bien accepté par la population des scripteurs, on choisit un seuil peu sévère qui réduit les risques de rejet d'une signature authentique : les gens n'apprécient pas de voir leur signature refusée. Les performances de ce système sont fondamentalement limitées car les paramètres globaux sont facilement imitables et relativement variables pour un même scripteur.

- **Alignement élastique de Brault [Brault 87b, Brault 88]** : Dans ce système, la signature-test est comparée avec chacune des signatures de référence par un algorithme de mesure de distance entre séquences, élaboré sur les principes de la programmation dynamique. Cet algorithme a également été utilisé pour la correction d'erreur typographiques [Wagner 74] et pour la reconnaissance de la parole. L'idée de la méthode est de trouver la séquence de comparaisons locales qui minimise le coût total de comparaison entre deux séquences. Le taux de reconnaissance est un des meilleurs obtenus en VdS [Parizeau 90]. Ce mémoire s'inspire fortement de cet algorithme et cherche à améliorer ses performances. L'algorithme qui était utilisé pour l'étape de comparaison, servira ici comme étape préliminaire d'identification des portions élémentaires de la signature. Le chapitre 2 traitera de cet algorithme en détail.
- **Réseaux Neuronaux** : l'essor des méthodes connexionnistes a conduit à des tentatives d'applications des réseaux de neurones dans de très nombreuses disciplines scientifiques. La VdS n'est pas en reste, notamment avec Lalonde [Lalonde 93] qui a élaboré un système de VdS entièrement basé sur les réseaux de neurones (réseau

multicouches pour la segmentation et carte topographique de Kohonen pour la reconnaissance), ainsi que Brault qui implémente actuellement le réseau de Hopfield-Tank conçu initialement pour le signal de parole [Brault 96]. La VdS n'est peut-être pas un domaine d'application idéal des réseaux neuronaux, en raison de la spécificité de l'apprentissage aux données (le réseau ne peut généraliser à un nouveau scripteur), de l'absence de données représentatives de la classe non-authentique, et de la structure interne de la signature très complexe à modéliser.

Dans notre approche, l'alignement élastique sert à synchroniser les signatures afin d'identifier leurs éléments manuscrits. On transforme ainsi les signatures en vecteurs de taille constante contenant l'ensemble des caractéristiques locales de la signature. On peut ensuite traiter ces vecteurs selon la RF statistique, d'une manière similaire à Nelson, mais pour des paramètres locaux.

1.1.4 Problèmes inhérents à la VdS

1.1.4.1 Représentation des signatures

Pour distinguer une souris d'un éléphant, nous n'avons pas besoin de faire des comparaisons internes du nez, des yeux, etc. La différence de taille suffit largement à les distinguer ! Par contre, il faudra être beaucoup plus attentif si on désire distinguer deux frères jumeaux : l'information globale ou superficielle ne suffisant pas, il faudra aller chercher l'information locale au niveau de la « structure » des frères jumeaux. Sans en

être conscient, on comparera le nez, les yeux, les différentes parties de la peau. Sur la base d'un grain de beauté, on pourra faire la distinction entre les deux.

Il en va de même en VdS : une signature authentique et un faux réussi sont deux objets très similaires, et les comparer globalement (durée, vitesse moyenne du tracé) ne suffit pas. Or il est très facile pour un système de VdS d'extraire des caractéristiques globales, et c'est pourquoi plusieurs systèmes fonctionnent ainsi. Ces systèmes interceptent en général les imitations de mauvaise qualité mais ne peuvent espérer rivaliser avec les experts humains en VdS.

Pourquoi ne pas caractériser localement les signatures (c'est-à-dire évaluer, pour chacune des portions du tracé, la vitesse, l'orientation, etc.) ? C'est que, avant de caractériser les éléments manuscrits d'une signature, il faut les identifier. L'identification des éléments manuscrits est un problème complexe (voir chapitre 2), et les méthodes d'identification existantes comportent une part non négligeable d'incertitude. Si l'identification est fautive, toute analyse ultérieure est non-valable (« garbage in, garbage out »).

Ce mémoire propose une méthode robuste d'identification des éléments manuscrits, permettant une caractérisation locale des signatures.

1.1.4.2 Quelles sont les classes ?

La RF est une discipline qui a pour but de classer un objet parmi un ensemble de classes connues. En VdS, on cherche ainsi à classer une signature parmi deux classes : la classe « authentique » ou la classe « non-authentique ». A partir d'un ensemble de

signatures authentiques d'apprentissage, le système extrait certaines caractéristiques générales qui lui permettent de décrire la classe authentique.

Mais la classe des signatures non-authentiques, dite classe « fantôme », est difficilement descriptible. En effet, une signature non-authentique peut être n'importe quoi : une signature d'un autre scripteur entrée par erreur ou une imitation. Il y a à la limite autant de classes d'imitateurs que de personnes à la surface de la terre. On préfère plutôt parler d'une méta-classe non-authentique regroupant toutes ces classes.

De plus, si pour tester un système de VdS on dispose d'un ensemble de faux, on ne dispose en général pas d'imitations lors de l'utilisation du système: il serait complètement irréaliste de collecter, pour chacun des utilisateurs un ensemble de faux servant à l'apprentissage.

Comment décrire une classe qui n'existe pour ainsi dire pas (puisque'elle est à la limite constituée de toute signature qui ne provient pas du signataire authentique) et dont l'acquisition d'un échantillon est trop coûteuse en pratique?

Ce mémoire propose une modélisation statistique de la VdS dans lequel est proposée une représentation de la classe des imitations, s'appuyant sur certaines hypothèses justifiées. Le modèle peut entre autres servir de générateur d'imitations.

1.2 Objectifs du mémoire

La discussion de la section précédente a souligné le fait que la qualité essentielle d'un système de VdS est son taux de reconnaissance et que, dans ce but, la modélisation

adéquate de l'information contenue dans une signature est indispensable. Certains auteurs [Brault 88, Xiao 95] ainsi que l'auteur de ce mémoire croient que l'extraction de caractéristiques locales et leur utilisation efficace sont une des avenues les plus prometteuses en VdS. Sur la figure suivante, le lecteur peut vérifier que, pour déterminer une signature-test, il faut effectivement faire des comparaisons locales entre une signature authentique et cette signature. Ces comparaisons entre éléments manuscrits locaux sont faites une fois que ces éléments sont identifiés.

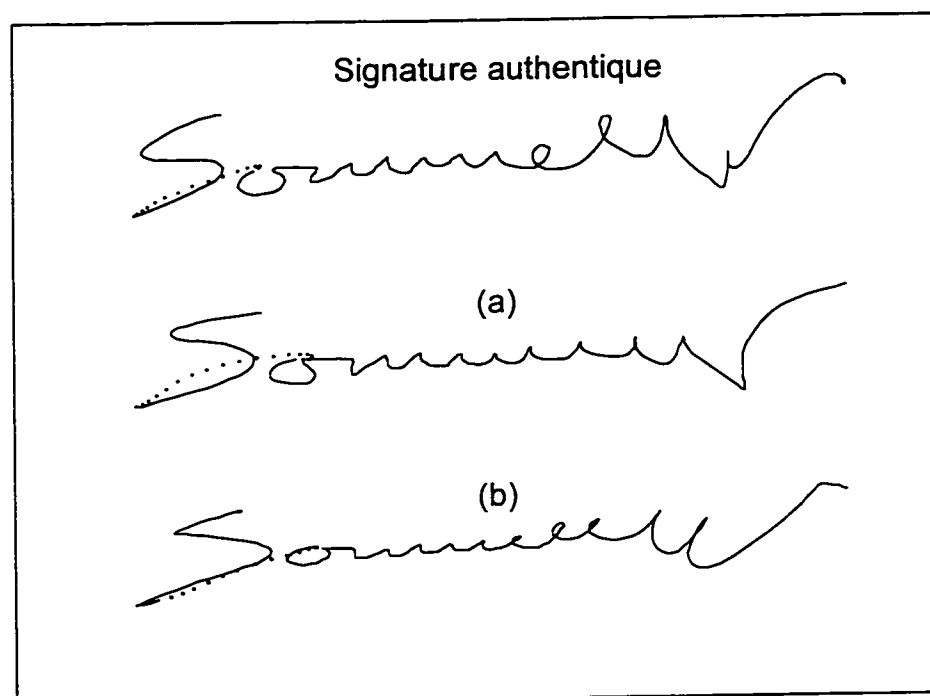


Figure 1-6 : En vous inspirant de la signature authentique située en haut, trouvez laquelle des signatures (a) et (b) est une imitation.

Réponse : (b) est une imitation, (a) est authentique.

L'objectif général de cette recherche est donc la conception d'un système de VdS basé sur l'extraction de caractéristiques locales des éléments manuscrits.

Cet objectif général se divise en 4 objectifs plus précis :

1. **Améliorer les performances de l'alignement élastique** : L'idée à la base de cette recherche est d'améliorer les performances de la méthode d'alignement élastique, une des meilleures méthodes de VdS, par l'utilisation de statistiques locales. La stabilité de la signature d'un scripteur varie selon les régions, il est donc normal de pondérer les coûts locaux de comparaison d'une signature-test proportionnellement à cette stabilité locale évaluée statistiquement. Bien que cette idée de base ait évolué au cours des recherches, améliorer les performances de l'alignement élastique reste un de nos objectifs.
2. **Élaborer une méthode de détection et identification des éléments manuscrits** d'une signature, en les identifiant aux éléments manuscrits de même nature sur différentes signatures de référence. Cette méthode de détection devrait aussi permettre d'avoir une segmentation stable.
3. **Modéliser statistiquement la VdS**. Cette modélisation serait utile pour avoir une meilleure connaissance de la classe des imitations, pour obtenir des résultats théoriques sur les taux d'erreur espérés avec un certain système de VdS et un certain scripteur. Éventuellement, ce modèle pourrait servir comme générateur de signatures vraies et fausses.
4. **Déterminer la meilleure représentation locale des signatures**. Le choix des paramètres locaux décrivant une signature est très vaste . Il s'agit de déterminer quelle

information locale permet la meilleure discrimination entre signatures vraies et fausses.

Pour tester la validité des méthodes développées, une étude a été réalisée avec 8 scripteurs donnant 30 exemplaires de leur signature, et 3 imitateurs donnant 10 imitations de chacune des signatures. Cela donne 60 signatures par scripteur et 480 au total. Le meilleur taux d'erreur obtenu sur l'ensemble des signatures est de 1.04%, contre 2.08% pour la méthode d'alignement élastique. A titre comparatif, une expérience d'expertise humaine sur les mêmes signatures a été soumise à des volontaires (qui ne sont pas des experts), avec un montant d'argent en jeu : les taux d'erreur vont 10% à 33%.

Cette expérimentation permet de déduire des résultats importants : entre autres, on y apprend que la vitesse locale selon l'axe x (appelée v_x) est l'information la plus discriminante et que la segmentation est une opération nuisible.

1.3 sommaire de la méthode utilisée

Les systèmes de reconnaissance des formes ne respectent pas toujours à la lettre le schéma vu à la section 1.1. Le système de RF que nous proposons est en fait constitué de deux systèmes de RF imbriqués l'un dans l'autre :

- Le sous-système d'identification des éléments manuscrits détermine l'ensemble des classes existantes par apprentissage, ainsi que la classe de chacun des éléments manuscrits d'une signature test. Pour chacune des portions de cette signature-test il détermine la portion des signatures de référence qui coïncide avec elle.

- Le système de VdS utilise les identifications faites par le sous-système d'identification pour extraire les caractéristiques de la signature-test ou d'apprentissage. Dans le cas d'une signature-test, il évalue une distance entre cette signature et la classe authentique. La signature est acceptée si elle est inférieure à un certain seuil.

Nous allons maintenant présenter succinctement les deux systèmes de RF. La figure ci-dessous résume les opérations effectuées.

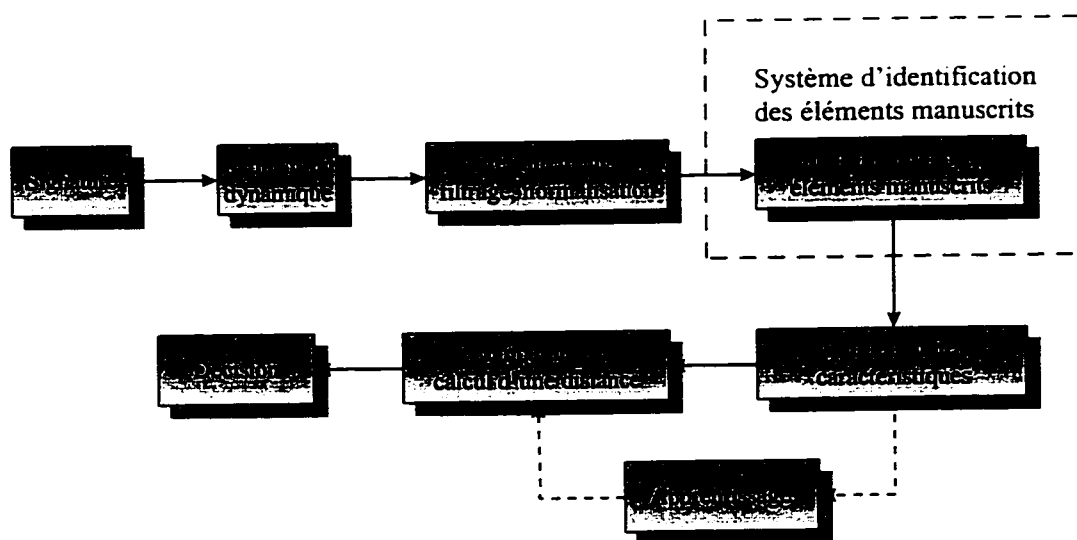


Figure 1-7 : Prototype du système de RF élaboré dans cette recherche

- **Acquisition** : l'acquisition des signatures se fait sur une tablette graphique avec un crayon électromagnétique. Les informations de départ sont la séquence temporelle de positions du crayon, échantillonnée à 100 Hz.
- **Prétraitements** : certains prétraitements ne posent pas de problème majeur : l'élimination du bruit est faite avec le filtre gaussien choisi pour ses propriétés remarquable (voir chapitre 2). La segmentation, traitement censé favoriser la reconnaissance en dégagant l'information significative, peut au contraire nuire en

donnant des résultats « incohérents » si elle est appliquée à l'aveugle, sans connaissance préalable des formes à segmenter. La segmentation sera appliquée lors de l'opération d'identification des éléments manuscrits, en tenant compte de l'ensemble des signatures d'apprentissage. Les signatures sont normalisées, bien que cette opération soit controversée en VdS (voir chapitre 4).

- **Identification des éléments manuscrits** : Il s'agit de déterminer quelles sont les classes d'éléments manuscrits existant dans une certaine signature. L'idée de la méthode consiste à sélectionner une signature authentique de référence, puis à « synchroniser » toutes les autres signatures de référence sur cette signature. Synchroniser signifie éliminer les distorsions temporelles, donc trouver quel sont les points de même nature ou de même classe. Il y a autant de classes que de points échantillonnés dans la signature de référence. On applique pour cela l'algorithme d'alignement élastique, qui donne pour résultat une séquence de comparaison des points échantillonnés minimisant une distance globale.. D'autres manipulations sont faites pour optimiser cette synchronisation, mais il est trop tôt pour en parler. La segmentation, qui n'est pas indispensable, se fait après cette synchronisation qui sert de référence. On utilise la technique des minima de vitesse, mais en l'appliquant sur la vitesse moyenne de toutes les signatures en un point synchronisé. On obtient ainsi un nombre de segments égal pour chaque signature, les segments de même position étant identiques.

- **Analyse** : certaines caractéristiques locales de la signature sont extraites, variant selon le système choisi. Notamment, on peut extraire la vitesse de chacun des points ou segments composant la signature.
- **Apprentissage** : l'apprentissage est effectué sur les signatures fournies en référence par le scripteur. Chacun des éléments manuscrits étant identifié sur l'ensemble des signatures de référence, on peut évaluer localement la moyenne et l'écart-type de la vitesse, de la position ou de l'accélération.
- **Comparaison et décision** : on calcule une distance entre la signature-test et la forme décrivant la classe authentique, forme évaluée lors de l'étape d'apprentissage. Différentes mesures de distances seront évaluées. On peut ensuite évaluer le seuil de décision en se basant sur des tests antérieurs ou sur la modélisation statistique de la VdS.

1.4 Plan du mémoire

- **Chapitre 2 : Identification des éléments manuscrits**. Les problèmes présents lors de l'identification des éléments manuscrits sont présentés, puis la méthode de synchronisation de deux signatures par alignement élastique est proposée comme solution (section 2.2). Une extension de la méthode à la synchronisation de N signatures est faite, permettant de pallier aux nouveaux problèmes qui surgissent lorsque les éléments manuscrits de plusieurs signatures à la fois doivent être identifiés (2.3). La problématique de la segmentation exposée, on voit comment utiliser cette méthode pour segmenter de façon stable les signatures (2.4).

- **Chapitre 3 : Modélisation statistique de la VdS.** Un modèle statistique des classes authentiques et non authentiques s'appuyant sur certaines hypothèses est proposé. Ce modèle permet une meilleure compréhension de la classe non-authentique (section 3.2). Des conclusions en sont tirées sur les performances des différents types de systèmes, et de scripteurs (3.3). Ce modèle peut servir également de générateur de signatures vraies et fausses.
- **Chapitre 4 : Conception du système de VdS et méthodologie d'évaluation.** L'ensemble des opérations effectuées sur les données sont présentées en détail. Certains paramètres du systèmes ne sont pas fixés, parce que l'on ne connaît pas la valeur optimisant les performances du système. L'expérimentation permettra de choisir les paramètres optimaux. La méthodologie d'évaluation des systèmes de VdS est présentée.
- **Chapitre 5 : Résultats de l'expérimentation:** L'analyse des résultats est faite en détail dans ce chapitre. On commence par exposer les résultats de référence : alignement élastique et test passés par des humains (section 5.2). Puis les performances obtenues selon les différentes paramétrisation du système (information de vitesse ou de position, choix d'un type d'élément manuscrit, etc.) sont présentées sous forme comparative (5.3). Une dernière section est réservée à des études plus poussées (5.4).
- **Chapitre 6 : Conclusion générale.** Les contributions de cette recherche sont présentées, et une discussion sur les axes de recherche future clôt cette dissertation.

2. L'IDENTIFICATION DES ÉLÉMENTS MANUSCRITS

2.1 Introduction

2.1.1 But et outil de base

Dans ce chapitre, on explique la problématique générale de la phase complexe de détection et identification des éléments manuscrits d'une signature, et on expose la méthode par laquelle on parvient à nos fins avec des risques d'erreurs très faibles. La synchronisation de 2 signatures, la synchronisation de N signatures ($N > 2$) et le problème spécifique de la segmentation de signatures sont 3 problèmes différents et font l'objet d'une section chacun.

L'outil de base utilisé est l'alignement élastique, bien connu dans le domaine de la VdS [Brault 88, Plamondon 94, Wirtz 95]. Cette méthode d'alignement, de deux séquences est issue des techniques de la programmation dynamique. Cependant les raisons pour lesquelles nous l'utilisons diffèrent de sa vocation première, qui est de mesurer une distance entre deux séquences. Au lieu d'en tirer une décision sur la classe à laquelle appartient une signature, l'alignement élastique n'est que la première étape vers la reconnaissance de la forme, permettant par la suite de faire des analyses locales. L'alignement élastique a donc pour but une représentation plus « fine » des signatures.

La méthode de synchronisation de plusieurs signatures et la méthode de segmentation reposent également sur l'alignement élastique. Cette méthode de segmentation est une fusion de trois techniques, soit l'alignement élastique, la segmentation aux minima de vitesse et le filtrage gaussien. La segmentation aux minima de vitesse est très utilisée en VdS car elle correspond à des nécessités morphologiques : il y a un lien très fort entre les ralentissement de vitesse et les changements de direction du tracé. Comme le signal contient du bruit et que ce bruit engendre la présence de nombreux minima de vitesse indésirable, on élimine ceux-ci en appliquant le filtre gaussien, qui est un filtre passe-bas aux propriétés remarquables (celles-ci seront exposées dans ce chapitre). Afin d'obtenir des segments semblables sur différentes signatures, on utilise au préalable l'alignement élastique pour synchroniser la segmentation.

2.1.2 Plan du chapitre

- La **section 2** est consacrée à la synchronisation de deux signatures. On débute par une présentation de la problématique de l'identification des éléments manuscrits sur deux séquences. Puis on donne une présentation de la méthode d'alignement élastique, succincte mais largement suffisante pour la compréhension de son utilisation. On verra ensuite de quelle façon on parvient avec cette technique à synchroniser deux signatures.
- En **section 3**, on commence par exposer le but et la problématique de la synchronisation de N signatures. En effet, de nouveaux problèmes apparaissent lorsque

plus de deux signatures doivent être synchronisées : explosion combinatoire des solutions, possibilités de contradiction entre la synchronisation de plusieurs signatures. On verra comment l'utilisation d'une signature de référence nous permet de pallier à ces problèmes.

- La **section 4** débute par une présentation générale de la segmentation. La segmentation aux minima de vitesse et le filtrage gaussien sont introduits, suivis de la problématique de la pré-segmentation. La pré-segmentation correspond à une segmentation avant toute connaissance de la forme. On termine en présentant la méthode de segmentation de N signatures, qui a la particularité d'être postérieure à l'étape de synchronisation (la majorité des méthodes de segmentation sont préliminaires à toute reconnaissance de forme), évitant ainsi les difficultés auxquelles la pré-segmentation doit faire face : l'identification des segments implique une mesure de similarité entre segments rendant la reconnaissance de la forme hypothétique, explosion combinatoire de possibilités lorsque plusieurs signatures sont segmentées. Comme dans la plupart des cas, deux signatures différentes ne posséderont pas le même nombre de segments, on doit considérer les éventualités suivantes (entre autres): segments ne correspondant pas à une forme connue (décision de rejeter le segment), deux segments consécutifs correspondant à une seule forme connue (nécessité de prendre la décision de fusionner ces segments). Notons que si deux signatures possèdent le même nombre de segments, cela n'entraîne pas pour autant que les segments doivent être associés un à un.

2.2 Synchronisation de deux signatures

2.2.1 But et problématique de la synchronisation de deux signatures

Le but de la synchronisation de deux signatures peut s'exprimer ainsi: chacune des signatures étant composée d'éléments manuscrits, on cherche à associer les éléments manuscrits de même nature sur les deux signatures. De façon plus formelle, la synchronisation doit répondre à cette question: pour tout point p_i d'une signature de référence S_i quel est, s'il existe, le point de même nature p_j sur une signature S_j . Sur la figure suivante, on a un exemple de points de même nature.

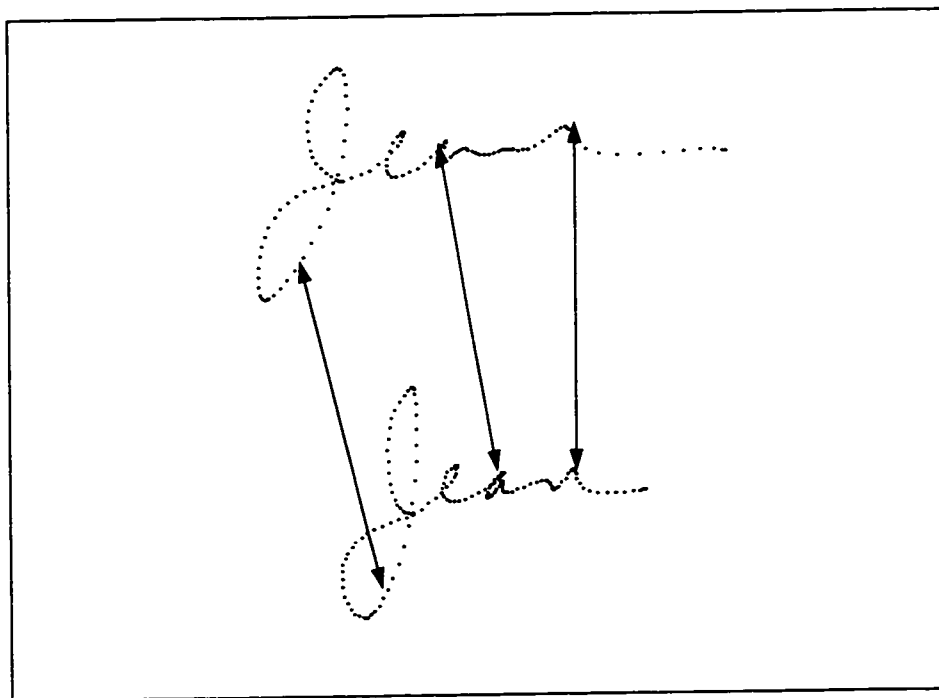


Figure 2-1 : Points de même nature sur deux signatures

Une méthode d'identification des éléments manuscrits doit tenir compte de l'ensemble du contexte dans lequel se situent les deux éléments manuscrits en question. Pour que deux éléments manuscrits soient associés, ils doivent avoir des caractéristiques semblables, et les portions de signatures situées de part et d'autre de ces deux points doivent également avoir des caractéristiques semblables.

On peut distinguer trois grandes classes de problèmes qui rendent le problème de la synchronisation de deux signatures si complexe, soit ceux de distorsion, de variabilité et d'inconstance :

Distorsion: Le premier problème qui complexifie la synchronisation de deux signatures est la distorsion temporelle des signatures d'un même signataire. Soit un certain tracé effectué par un signataire en un temps t . S'il effectue à nouveau le même tracé, il y a de fortes chances que ce soit en un temps t' différent de t . Comme ce tracé est échantillonné, le même tracé aura donné des séquences de points différentes.

Ainsi, un signataire produit des signatures similaires en des temps plus ou moins différents. Bien sûr, si la distorsion est globale le problème se résout en rééchantillonnant le signal. Mais pour les signatures, cette distorsion est présente localement et est non-uniforme : alors qu'une région r_j d'une signature S_j sera de durée plus longue que la région associée r_i d'une signature S_i , la région annexe r_{j-1} de S_j pourra très bien être de durée plus courte que la région équivalente r_{i-1} de S_i .

On peut observer le problème de distorsion sur la figure suivante. Sur la signature (1), les portions situées entre 2.26 s et 2.80 s, et entre 2.80 s et 3.34 s durent chacune 0.54 s et 0.54 s. Sur la signature (2), les portions équivalentes, entre 2.55 et 3.07 et entre 3.07 et 3.65, durent respectivement 0.52 et 0.58 secondes. Ainsi la 1^{ère} portion est compressée alors que la deuxième est dilatée, par rapport à la signature 1.

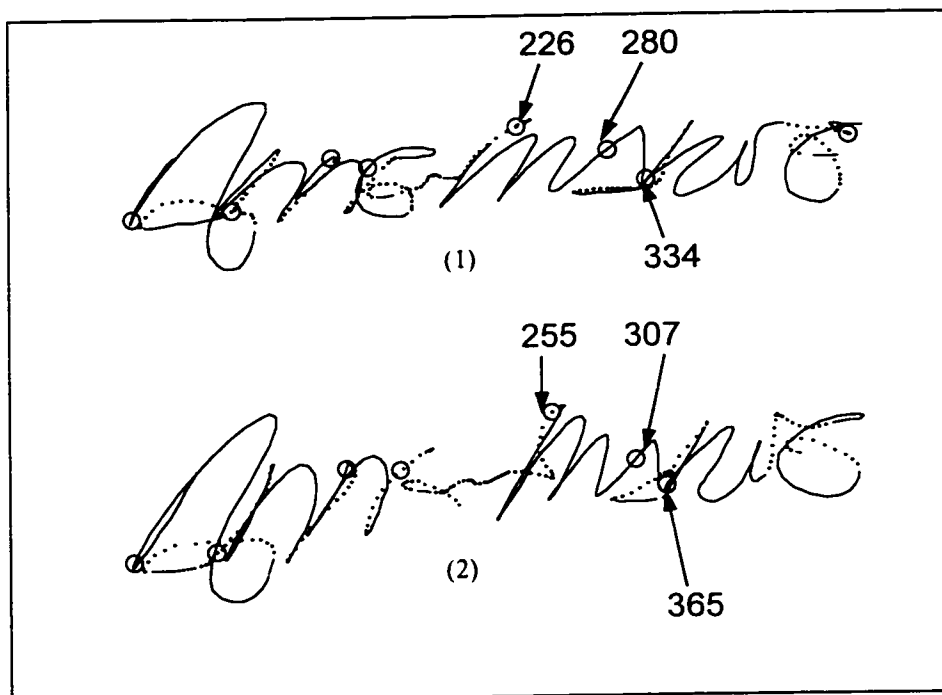


Figure 2-2 : La distorsion

Variabilité: Le deuxième problème est la variabilité d'un même signataire. La variabilité correspond aux variations du tracé d'un signataire sur différentes signatures. Ces variations sont dues aux instabilités du mouvement de la main. Sur la figure suivante, la séquence de droite a subi de légères dilatations et rotations locales, illustrant le problème de variabilité.

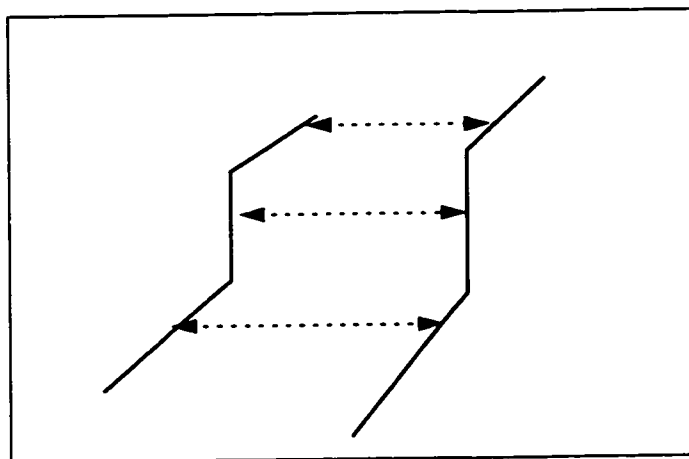


Figure 2-3 - La variabilité

- **Inconstance :** Le troisième problème est particulièrement présent chez certains signataires instables, qui n'ont pas encore de signature bien déterminée : ceux-ci rajoutent ou retranchent des parties importantes de leur signature d'une fois à l'autre. Cette partie peut par exemple être une lettre ou un aller-retour dans le sens vertical (un " m " est parfois tronqué à un " n "). On peut voir ce problème illustré sur les 2 figures suivantes. Sur la Figure 2-4, le problème de l'inconstance est illustré par l'ajout d'un segment à la séquence de droite.

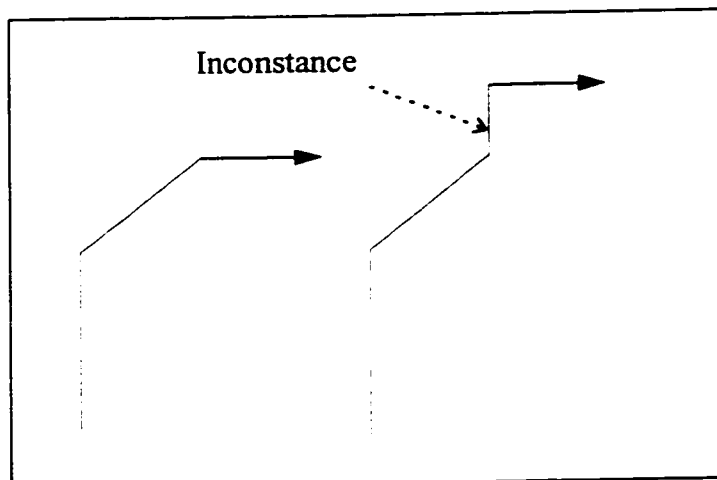


Figure 2-4 : L'inconstance

Sur la Figure 2-5, les inconstances d'un signataire (« Fred ») sont illustrées : les parties (1) et (2) ne sont pas présentes sur la signature de gauche.

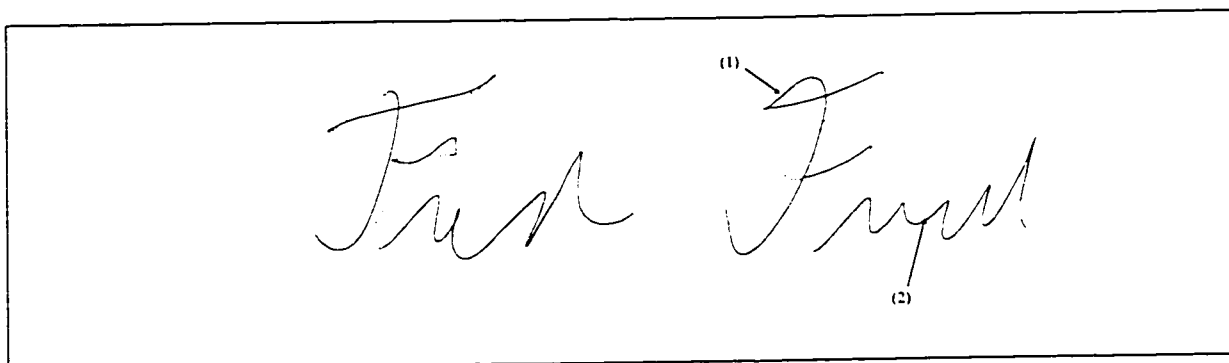


Figure 2-5 : Inconstances d'un signataire

Nous n'avons toujours pas défini ce qu'est un élément manuscrit, ou du moins ce qu'on entend par cette expression. On sait qu'un élément manuscrit est une portion de signature, mais cela n'est pas une définition suffisante. Est-ce-que chacun des points composant la signature peut être considéré comme élément manuscrit, ou est-ce plutôt

une portion du tracé de même direction, appelée segment? Nous estimons qu'il est trop tôt encore pour répondre complètement à cette question (si tant est qu'elle ait une réponse).

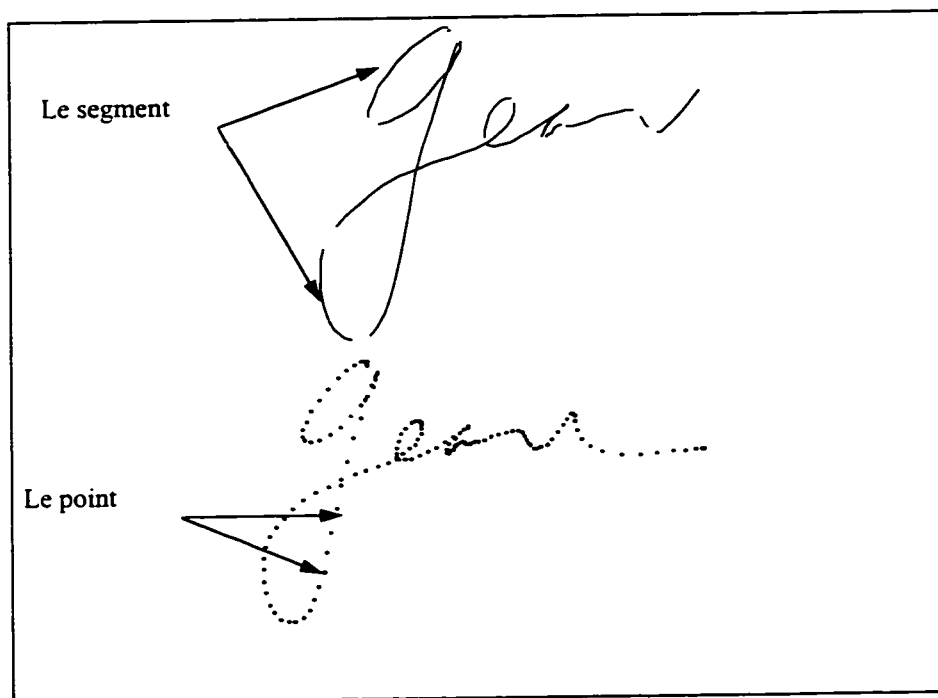


Figure 2-6 : L'élément manuscrit, un point ou un segment ?

Nous donnerons donc pour l'instant la définition non-contraignante suivante: *un élément manuscrit est toute portion d'une signature qui est identifiée sur une ou plusieurs autres signatures d'un même scripteur*. Ainsi cela peut être autant un segment qu'un point. Lors de la phase de synchronisation, comme il vaut mieux ne pas segmenter pour éviter de déformer l'information, les éléments manuscrits seront définis comme étant les points. Rien ne nous empêchera après la synchronisation de segmenter les signatures, si cela permet une meilleure caractérisation des éléments manuscrits.

2.2.2 Identification des éléments manuscrits des signatures: une revue de la littérature

L'identification des éléments manuscrits d'une signature n'est pas un domaine qui a été intensivement traité. Néanmoins, diverses approches ont été tentées, que nous allons décrire brièvement dans cette section:

- Nous avons consulté un livre constitué d'articles résumant ce qui se fait actuellement en VdS [Plamondon 94b]. Ce livre permet de voir le peu de cas que l'on fait du problème de l'identification des éléments manuscrits. Dimauro [Dimauro 94], qui est le seul à en parler, propose une méthode pour extraire ce qu'il appelle "les composantes" d'une signature. Ces "composantes" sont en réalité non pas des segments ou encore des lettres, mais carrément des groupes de lettre séparés par des levers de crayon. Dans une signature sans lever de crayon, la composante serait la signature elle-même ! Il est bien évident que nous désirons une description beaucoup plus fine des signature .
- Différents auteurs [Plamondon 94, Lalonde 93] proposent de pré-segmenter les signatures. Les segments sont décrits chacun par un vecteur de caractéristiques qui doit permettre de les identifier. On verra en section 4 quels problèmes implique toute méthode basée sur une pré-segmentation.
- Le réseau de neurones à délais de Hoppfield-Tank [Unnikrishnan 91] a été conçu initialement pour reconnaître les signaux issus de la parole. Il ne nécessite aucune segmentation du signal, et reconnaît des mots dilatés ou compressés jusqu'à 50%

par rapport aux mots ayant servi à son apprentissage, par une identification des constituantes de ce mot. Son usage pour la VdS a d'abord été suggéré par Lalonde [Lalonde 94]. Il est présentement implanté par Brault [Brault 96] pour la VdS. C'est selon nous une approche neuronale prometteuse pour l'identification des éléments manuscrits. Néanmoins il converge, sans nécessairement l'atteindre, la solution optimale trouvée par un algorithme de programmation dynamique (que nous allons utiliser!).

- Finalement, l'approche de Menier pour la reconnaissance d'écriture [Menier 95] mérite d'être citée. En effet, pour la phase de reconnaissance des lettres il utilise l'alignement élastique entre le mot inconnu et les mots de référence (technique que nous allons largement décrire plus loin), afin de pouvoir comparer les vecteurs de caractéristiques. Cette approche ressemble à la nôtre mais est adaptée au problème de reconnaissance d'écriture, qui diffère de la VdS.

2.2.3 Présentation informelle de l'alignement élastique

Un problème de la RF est de trouver une mesure de distance valable entre deux séquences. L'algorithme d'alignement élastique (appelé « Dynamic time warping » en anglais) évalue la différence entre deux signatures S_i et S_j comme étant le coût total des modifications minimales que l'on doit apporter à la signature S_j pour obtenir la signature S_i .

La méthode d'alignement élastique des signatures développée par Brault [Brault 87, Brault 88] s'inspire de l'algorithme de Wagner-Fischer [Wagner 74] pour la

comparaison de chaînes de caractères. Dans cet algorithme, on cherche la séquence de modifications minimales à appliquer à une chaîne de caractère C_1 pour obtenir une chaîne de caractères C_2 . Les modifications sont appliqués sur les **primitives**, qui sont ici les caractères. Trois types de modifications, appelées opérations d'édition, sont permises: la **suppression** d'un caractère de C_1 , l'**insertion** d'un caractère de C_2 et la **substitution** d'un caractère de C_1 à un caractère de C_2 . A chaque opération est rattachée un coût, et il s'agit de trouver la séquence de modifications qui minimise le coût total.

Par exemple, avec les chaînes de caractères $C_1 = \text{GBCD}$ et $C_2 = \text{A'B'C'}$, une séquence de modifications possibles sur C_1 pour obtenir C_2 serait dans l'ordre:

- supprimer G sur C_1 ;
- insérer A' de C_2 ;
- substituer B de C_1 à B' de C_2 ;
- substituer C de C_1 à C' de C_2 ;
- supprimer D de C_1 .

Il y aurait plusieurs autres séquences de modifications possibles pour transformer C_1 en C_2 . On comprend que le choix d'une séquence de modifications n'est pas indifférent: si on définit un coût pour chaque type de modification, l'algorithme trouve la séquence de modifications qui minimise le coût total de modifications. Ce coût est la somme de chacun des coûts individuels rattachés à chacune des modifications. Il faut donc décider du coût rattaché à une modification.

Il faut ensuite trouver la séquence de modifications qui minimise le coût total. Une possibilité est de toutes les essayer. Cela est possible dans ce cas-ci. Pour des séquences de la taille des signatures (200 à 600 points), il y a un nombre gigantesque de comparaisons possibles. Brault [Brault 88] a évalué que pour deux signatures de 500 points, il y aurait 2^{380} séquences de modifications possibles!

Un algorithme inspiré des techniques de la programmation dynamique permet de trouver la séquence de modification optimale en un temps proportionnel à $(n_i * n_j)$ où n_i et n_j sont les tailles des séquences S_i et S_j respectivement. Pour expliquer le principe général de la programmation dynamique [Bellman 57], nous citons [Belaïd 92, page 136] :

« La programmation dynamique est une méthode de recherche d'optimum fondée sur le principe d'optimalité locale de Bellman. [...] Ce principe indique que tout chemin optimal est constitué de portions elles-mêmes optimales. En effet si on considère un chemin optimal reliant A à B et un point quelconque M de ce chemin, les chemins de A à M et de M à B sont tous deux optimaux. »

Dans la méthode d'alignement élastique des signatures, les primitives qui étaient les lettres dans les chaînes de caractères deviennent les segments de droite reliant deux points successifs d'une signature. Les règles de modification (suppression, substitution et insertion) restent les mêmes et les coûts rattachés aux modifications sont des distances métriques euclidiennes. Le choix d'une distance métrique euclidienne par rapport à d'autres types de distances a été largement justifié dans [Brault 88] et nous ne jugeons pas nécessaire de revenir sur ce point. Le calcul d'une distance entre deux primitives

vectérielles est une soustraction de deux vecteurs quelque soit le type de modification: dans le cas d'une substitution, on soustrait les segments de droite substitués, et dans les cas d'insertion et suppression, la soustraction du segment de droite inséré ou supprimé se fait avec le vecteur nul.

Les bases de la méthode d'alignement élastique ont été fermement établies et nous nous y tenons, à l'exception du choix de la primitive qui est la vitesse évaluée au point lui-même, et non le segment de droite entre deux points. Un segment de droite entre deux points i et $(i+1)$ correspond en fait à la vitesse du point $(i + \frac{1}{2})$ évaluée par une différence finie d'ordre 1. Puisque nous cherchons à identifier des points, il est plus précis de prendre comme primitive les points eux-mêmes. De plus, la formule par laquelle nous calculons la vitesse (voir section 4 du même chapitre) étant d'ordre 2, nous aurons des mesures plus précises des coûts rattachés aux modifications.

Jusqu'à maintenant, l'intérêt principal du calcul d'une distance par alignement élastique résidait dans son efficacité à distinguer les imitations des signatures authentiques. Même si, d'après Parizeau [Parizeau 90], cette méthode n'a pas démontré clairement qu'elle était supérieure à d'autres méthodes connues (corrélation régionale et alignement par arbres squelettiques), nous doutons que l'on puisse généraliser à toute primitive car l'alignement élastique a été fait en séparant l'information selon l'axe X de celle selon l'axe Y. Nous sommes convaincus que l'alignement élastique est une des meilleures méthode de VdS existantes, et l'utilisation qui en est faite ici se veut une amélioration de la méthode originale.

Dans l'utilisation classique de la méthode, la séquence de comparaisons n'était intéressante que dans la mesure où on voulait s'assurer qu'elle associait des parties équivalentes des deux signatures, afin que la mesure de distance soit bel et bien représentative d'un écart entre elles. Cependant le simple calcul d'une distance est loin d'exploiter toutes les possibilités de l'alignement élastique : comme il synchronise les signatures de façon optimale, on peut en faire une méthode préliminaire d'identification des éléments manuscrits. On peut alors évaluer les caractéristiques de chacun des éléments manuscrits sur un ensemble d'apprentissage et obtenir une description très précise de la forme générale de la signature d'un scripteur et de ses variations. Lors de l'identification d'une signature-test, on identifie les éléments manuscrits qu'elle contient puis on peut appliquer une méthode de reconnaissance indépendante de la phase d'identification (méthode statistique bayésienne, réseaux neuronaux, k plus proches voisins, ...).

Il est difficile d'évaluer quantitativement la qualité de la synchronisation par alignement élastique car pour cela, il faut faire appel au jugement humain. Mais de fait, les cas sont très rares où le résultat de l'alignement élastique entre en contradiction avec le jugement humain. Sur la figure suivante, on peut voir de quelle façon deux signatures sont alignées.

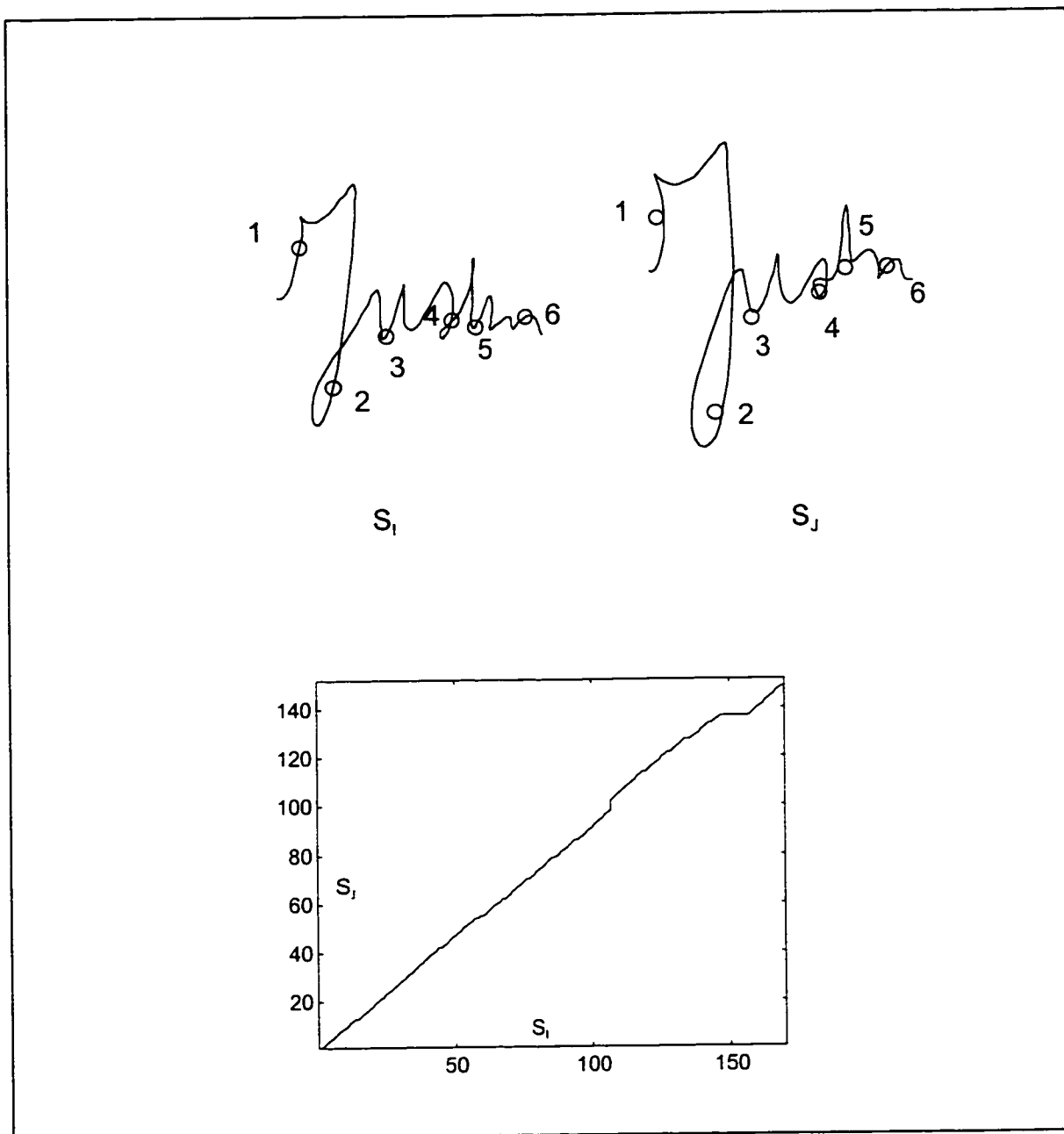


Figure 2-7 : Visualisation de certains points synchronisés par alignement élastique. Les points de même numéro sont de même nature

L'alignement élastique répond aux exigences que nous demandons à une méthode d'alignement de deux signatures. Premièrement, il ne pose pas de contrainte sur le choix

de l'élément manuscrit: nous pouvons comme nous le désirions porter notre choix sur les points. Ensuite, toutes les décisions prises minimisent un coût total de modification: deux éléments manuscrits sont associés s'ils sont suffisamment semblables, et si les conséquences globales de cette association sont les moins coûteuses possibles. Du même coup, la robustesse de la méthode est assurée, puisqu'il est garanti que la solution optimale est trouvée.

La politique optimale de comparaison nous assure que l'association de deux éléments est celle qui est le plus en rapport avec le contexte : toute autre association d'éléments aurait entraîné une séquence de modification plus coûteuse. Nous en donnerons une preuve plus loin. Passons maintenant à la description méthodique de l'alignement élastique.

2.2.4 Présentation formelle de l'alignement élastique pour la VdS

Soient deux signatures $S_I = (p'_1, \dots, p'_{N_I})$ et $S_J = (p'_1, \dots, p'_{N_J})$ composées respectivement de N_I et N_J points. On cherche la séquence de comparaisons qui minimise une distance totale entre S_I et S_J , cette distance étant la sommation des distances rattachées à chacune des comparaisons. Les comparaisons sont de trois types: substitution de p'_i en p'_j , suppression de p'_i et insertion de p'_j . Les comparaisons peuvent être vues comme des associations:

- pour une substitution, association de p'_i avec p'_j

- pour une insertion, p_j' n'a pas de points associé, association de p_j' avec $\emptyset$
- pour une suppression, p_i' n'a pas d'éléments associé, association de p_i' avec $\emptyset$

Les comparaisons permettent de progresser dans le chemin de comparaison des deux séquences :

- une substitution permet de passer de (p_{i-1}', p_{j-1}') à (p_i', p_j') ;
- une insertion permet de passer de (p_i', p_{j-1}') à (p_i', p_j')
- une suppression permet de passer de (p_{i-1}', p_j') à (p_i', p_j')

On fait appel aux vitesses des points pour calculer les coûts des modifications.

Soient $S_I = (\bar{v}_1', \dots, \bar{v}_{N_I}')$ et $S_J = (\bar{v}_1', \dots, \bar{v}_{N_J}')$ dans l'espace de représentation des vitesses.

Les $\bar{v}$ sont alors des vecteurs à deux dimensions.

Les coûts rattachés aux modifications sont:

- pour une substitution (sub): $|\bar{v}_i' - \bar{v}_j'|$,
- pour une insertion (ins): $|\bar{v}_j'|$
- pour une suppression (sup): $|\bar{v}_i'|$

L'algorithme de programmation dynamique fonctionne de façon récursive. Soit $c(i,j)$ le coût total de comparaison des séquences de points 1 à i de S_I et 1 à j de S_J . Alors:

$$c(i, j) = \begin{cases} c(i-1, j-1) + |\bar{v}_i' - \bar{v}_j'| \\ c(i-1, j) + |\bar{v}_i'| \\ c(i, j-1) + |\bar{v}_j'| \end{cases}$$

La distance entre S_i et N_j est le coût total de comparaisons $c(N_i, N_j)$ aux derniers points N_i et N_j des séquences. La séquence de comparaisons qui s'y rattache est $t(N_i, N_j)$.

On peut voir sur la figure suivante, la politique de comparaisons optimale entre deux signatures d'un même scripteur: remarquons les "sauts" (parties du chemin verticales ou horizontales) sur le chemin, qui correspondent à de petites inconstances du scripteur, soit de petites régions de la signature présente sur une seule des deux signatures.

La méthode d'identification des éléments manuscrits considère qu'une substitution correspond à l'identification d'un point.

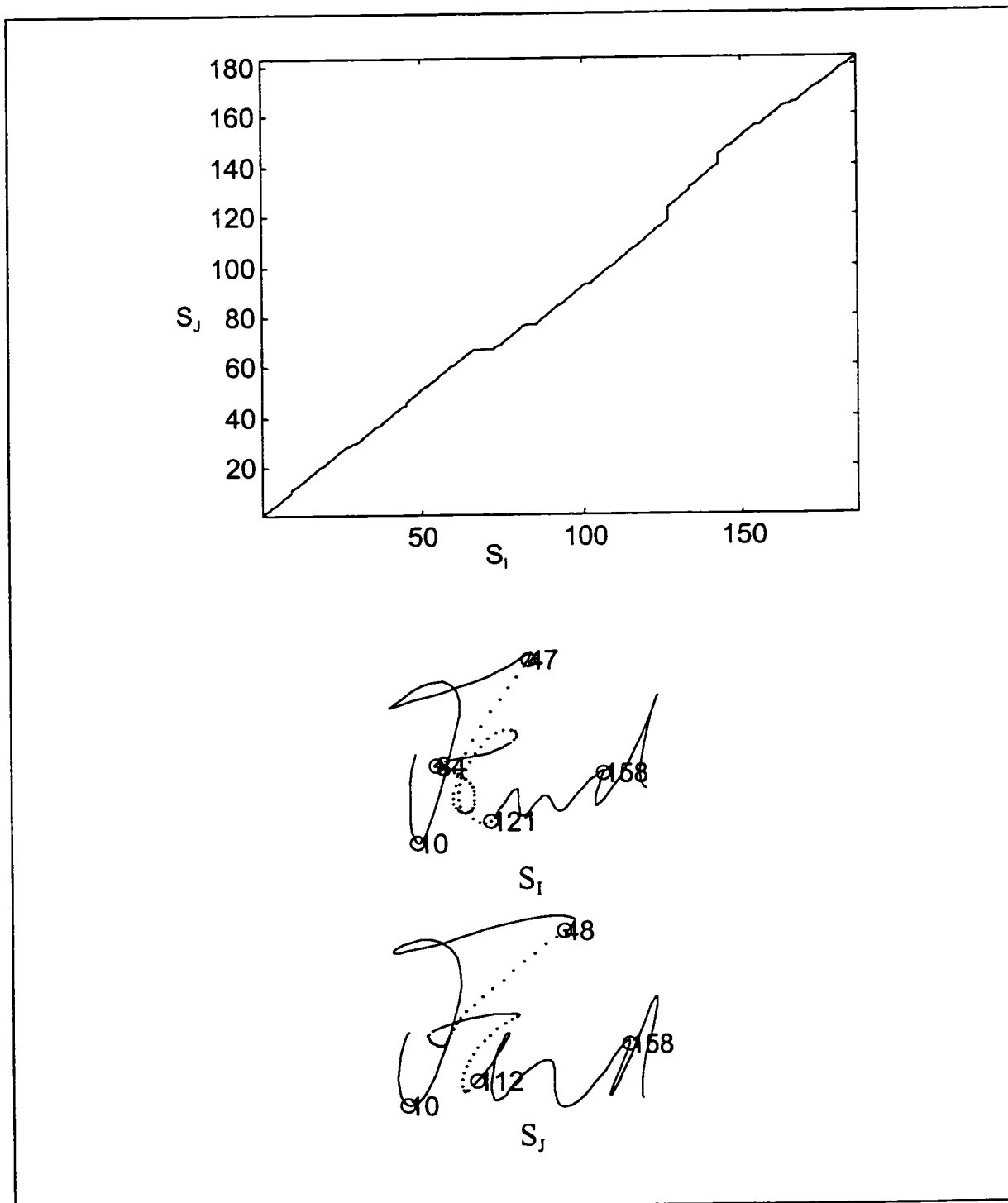


Figure 2-1 : Deux signatures d'un scripteur instable et la politique de comparaison.

2.2.5 Calcul d'une distance normalisée entre deux signatures

Brault [Brault 88] a utilisé la programmation dynamique pour le calcul d'une distance normalisée entre deux signatures. Le coût total est divisé par la racine des longueurs respectives des signatures S_i et S_j qui correspondent aux coût de comparaison de ces signatures avec des séquences vides. Cette normalisation permet de tenir compte des tailles respectives des signatures. Ainsi une imitation de petite taille n'aura pas une faible distance. On a donc :

$$D(S_i, S_j) = \frac{c(N_i, N_j)}{\sqrt{c(N_i, \emptyset) \times c(N_j, \emptyset)}}$$

Cette mesure de distance nous servira lors de l'expérimentation afin de comparer les résultats obtenus par notre méthode avec l'alignement élastique classique.

2.2.6 Justification de l'utilisation de l'alignement élastique comme méthode d'identification des éléments manuscrits

Jusqu'à maintenant, les justifications apportées à l'utilisation de l'algorithme d'alignement élastique n'étaient qu'intuitives. En effet, pourquoi tient-il compte du contexte d'une façon adéquate dans sa décision d'associer deux points? Nous apportons la justification dans cette section.

La question d'identification des points d'une signature S_j , S_i étant considérée comme signature de référence, peut se poser ainsi. Étant donné un point i d'une signature

S_j , on cherche, s'il existe, le point j de S_i auquel il devrait être identifié. Le point de S_i auquel il devrait être identifié est celui qui minimise une distance qui est composée :

1. d'une mesure du coût de comparaison $c(i,j)$ des deux points i et j
2. d'une mesure du coût de comparaison $c(S_i', S_j')$ des séquences de points $S_i' = \{1, 2, \dots, i-1\}$ avec $S_j' = \{1, 2, \dots, j-1\}$
3. d'une mesure du coût de comparaison $c(S_i'', S_j'')$ des séquences de points $S_i'' = \{i+1, \dots, N_i\}$ avec $S_j'' = \{j+1, \dots, N_j\}$

On a ainsi une mesure de la distance entre les deux points qui tient compte du contexte:

$$d(i,j) = c(S_i', S_j') + c(i,j) + c(S_i'', S_j'')$$

Le point j de S_j associé à i de S_i devrait être celui minimisant cette distance contextuelle:

$$j | d(i,j) = \min_j [c(S_i', S_j') + c(i,j) + c(S_i'', S_j'')] \\ S_j' = (1, 2, \dots, j-1) \\ S_j'' = (j+1, N_j)$$

Si on évalue $c(.,.)$ par la méthode d'alignement élastique, le point j de distance minimale au point i est celui se trouvant sur le chemin de comparaison optimal de S_i et S_j . Dans le cas où le point j a été inséré dans S_i , aucun point dans S_j n'a été trouvé pour être identifié au point i .

2.3 Synchronisation de N signatures

2.3.1 But et problématique de la synchronisation de N signatures

Le but de la synchronisation de N signatures $S_1, S_2, \dots, S_N$ est d'associer les éléments manuscrits de même nature sur ces N signatures. Cette synchronisation est nécessaire lors de l'apprentissage par le système des classes d'éléments manuscrits formant la signature d'un scripteur, et des caractéristiques de ces classes. Autrement dit, on désire avoir une estimation des caractéristiques locales de la signature d'un scripteur par un apprentissage sur un échantillon de N signatures.

Ce problème est sensiblement plus complexe qu'avec deux signatures. Il n'y a apparemment pas de moyen connu de comparer plus de deux séquences à la fois. Et nous aimerions conserver l'outil présenté à la section précédente, conçu que pour deux séquences à la fois mais qui nous garantit que les identifications faites sont optimales au sens d'un certain critère.

Il y a au moins 3 nouveaux problèmes qui surgissent lors de l'identification des éléments manuscrits sur N signatures :

- Pour N signatures, il existe $\frac{N \cdot (N - 1)}{2}$ comparaisons de signatures possibles. Si la méthode élaborée exige de faire toutes les comparaisons, on a un temps de calcul en $O(N^2)$, ce qui limite la taille de la banque de signatures de référence.
- Les identifications d'éléments manuscrits peuvent mener à des contradictions. Prenons seulement trois signatures S_1, S_2 et S_3 , il est possible que les comparaisons entre S_1 et

S_2 , S_1 et S_3 , S_2 et S_3 mènent à des identifications d'éléments manuscrits qui soient contradictoires entre elles. Un exemple simple: si par alignement élastique, le point i de S_1 est associé au point j de S_2 , j de S_2 à k de S_3 , et i de S_1 à $(k+1)$ de S_3 , devrait-on associer i à k ou à $(k+1)$? Si toutes les signatures sont synchronisées entre elles, on risque d'aboutir à un résultat non-cohérent.

- Pour deux signatures, on choisit une des deux signatures dont les points constitueront les classes d'éléments manuscrits. Pour N signatures, comment décider quels sont les classes d'éléments manuscrits? Comme certains éléments manuscrits ne se retrouvent que sur une partie des signatures, peut-on simplement choisir une signature dont les éléments manuscrits formeront les classes de référence ?

2.3.2 Description de la méthode

Faire toutes les synchronisations possibles entre N signatures apparaît comme étant inutilement coûteux et complexe. En choisissant une signature dont les éléments manuscrits serviront de référence, il n'y a que $(N-1)$ synchronisations à faire. Si le scripteur signe essentiellement d'une seule façon, la signature choisie contiendra l'essentiel des éléments manuscrits du scripteur.

La méthode choisie se divise en 4 étapes: (1) choix d'une signature S dont les éléments manuscrits serviront de référence, (2) synchronisation des autres signatures sur S , (3) génération d'une signature "moyenne" davantage représentative de la classe, (4) resynchronisation sur cette signature moyenne.

Les prochaines sections expliquent en détail la méthode.

2.3.3 Choix d'une signature comme espace temporel de projection

Le choix d'une signature de référence impose que seuls les éléments manuscrits de cette signature composeront l'ensemble des éléments manuscrits décrivant une signature authentique. Si cette signature ne comporte pas certains éléments manuscrits importants, ceux-ci seront irrémédiablement perdus. Selon quels critères choisir la signature de référence, soit celle qui permettra la « meilleure » synchronisation?

Il est clair que la meilleure synchronisation est celle qui donne la meilleure classification de signatures authentiques et imitées. Seulement en pratique on ne dispose pas d'imitations pour évaluer la qualité de la classification. Un autre problème se pose: si on doit faire tous les alignements élastiques possibles, pour connaître la meilleure synchronisation, les calculs deviennent très coûteux.

Cependant, le problème devient caduque si le choix de la signature de référence n'a pas une influence déterminante sur les performances du système.

Nous avons fait l'expérience suivante. Pour un scripteur assez stable, nous avons fait tous les alignements possibles, soit $\frac{30 \cdot 29}{2} = 435$ pour 30 signatures. Ensuite on a évalué l'écart-type (normalisé par la longueur) moyen de v_y (la vitesse selon l'axe y) de chacun des éléments synchronisés, pour chaque espace temporel de projection. L'écart-type, qui exprime la stabilité d'une caractéristique, est en effet la mesure fondamentale qui permet d'évaluer la difficulté d'imitation d'une signature: l'imitateur doit respecter la

précision imposée par l'écart-type. Plus celui-ci est petit, plus les éléments ont été synchronisés correctement et plus il sera dur d'imiter les signatures authentiques.

Nous n'avons mesuré aucune différence notable (variation de moins de 5%) entre les différents écarts-type moyens selon l'espace temporel de projection. Si le scripteur est stable, la synchronisation est suffisamment stable pour ne pas varier notablement en fonction du choix de la signature d'alignement.

Nous pensons qu'il vaut mieux tout de même prendre une signature de durée près de la durée moyenne de l'ensemble des signatures. Une signature beaucoup plus courte ou plus longue que la durée moyenne entraînerait un plus grand nombre d'insertions ou suppressions. De plus, les signatures de durée éloignée de la moyenne ont plus de chance d'être pathologiques.

La signature choisie est donc celle de durée la plus proche de la durée moyenne des signatures.

2.3.4 Détermination de la signature moyenne et resynchronisation

La méthode pourrait s'arrêter ici, et serait tout à fait fonctionnelle. En effet pour chaque point de la signature de référence, on connaît les points de même nature sur les autres signatures. On pourrait directement passer à l'apprentissage des caractéristiques de chacune des classes d'éléments manuscrits. C'est d'ailleurs ce que nous avons fait dans un premier temps, avant de tenter d'améliorer la synchronisation par la méthode suivante. Comme on le verra dans l'analyse des résultats, il semble bien que les résultats finaux (en

terme d'erreur) soient améliorés par cette extension de la méthode que nous présentons maintenant.

Soit $S_1, S_2, \dots, S_N$ N signatures de référence. S_1 est la signature choisie comme espace temporel de synchronisation. Soit, pour chacun des points i de SI , $\{i_1, i_2, \dots, i_N\}$ l'ensemble des points de même classe des autres signatures (dans certains cas, aucun point ne sera associé sur une autre signature, étant donné le fonctionnement de la méthode d'alignement élastique).

Pour un ensemble de points identifiés comme étant de même nature, on peut évaluer la moyenne et l'écart-type d'une certaine caractéristique de ce point. On peut notamment évaluer la vitesse moyenne pour chacun des groupes de points de même classe :

$$\bar{v}_{\text{moy}} = \frac{\sum_{i=1}^N \bar{v}_i}{N'}$$

avec N' étant le nombre de signatures qui ont un point identifié de cette classe.

Pour chacun des éléments manuscrits de référence, on a une estimation plus précise qu'auparavant de leur caractéristique de vitesse. On appelle « signature moyenne » la signature constituée de ces vitesses moyennes... Cette « signature moyenne » est plus représentative de la classe authentique que chacune des signatures de l'échantillon.

Pour chacune des N signatures, on peut évaluer une signature moyenne rattachée à cette signature. Afin de vérifier la stabilité de la signature moyenne, nous avons calculé,

pour un scripteur, les N signatures moyennes. On peut observer sur la figure suivante cinq de ces différentes signatures. Elles sont, à peu de chose près, identiques. Cela signifie que la synchronisation est très stable et varie très peu par rapport au choix de l'espace temporel de projection.

La distance $c(N_i, N_j)$ telle que définie à la section précédente entre les différentes signatures moyennes est en moyenne de 0.17, ce qui est une mesure très faible, les plus petites distances entre signatures étant généralement de 0.3. La signature moyenne, quel que soit la signature de référence, offre un moyen de synchronisation extrêmement stable. L'ensemble des prises de décision qui mène à la séquence de comparaisons s'effectue en tenant compte de toutes les signatures de référence et non d'une seule. Les distorsions locales ne risquent pas de dérégler la synchronisation comme cela pourrait être (rarement) le cas si l'alignement élastique était fait sur une signature quelconque.

Le lecteur peut se demander quelle apparence a une signature moyenne, en particulier si les signatures sont très instables et inconstantes. Nous avons choisi le scripteur le plus instable (« Fred »), dont un exemple de signature moyenne est donné à la figure suivante.

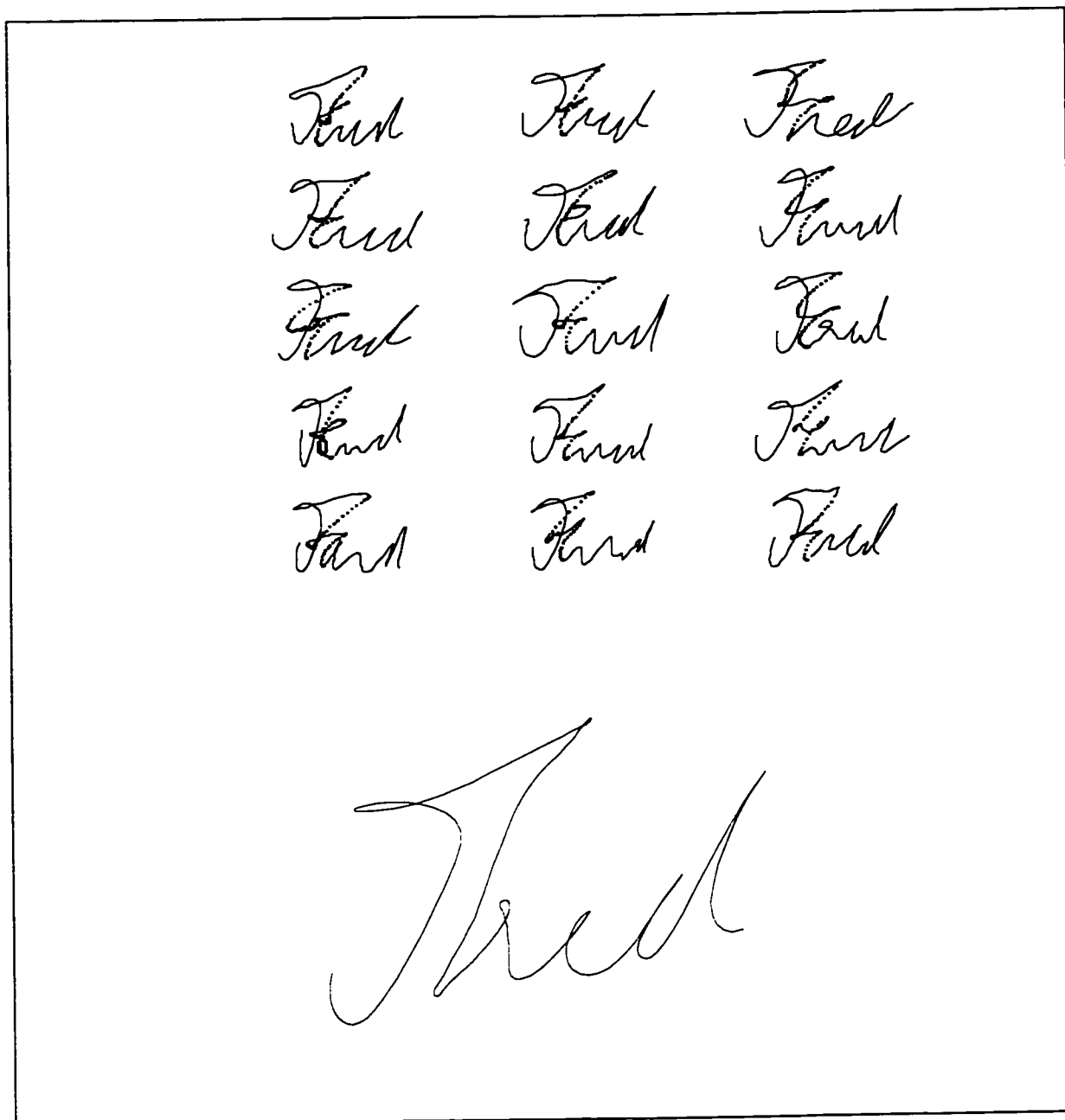


Figure 2-9 : 15 des 30 Signatures d'un scripteur instable et inconstant, et la signature moyenne qui en résulte

On peut aussi afficher sur la signature moyenne l'écart-type d'une ou plusieurs caractéristiques ponctuelles, en faisant varier l'épaisseur du tracé. On peut observer un

exemple sur la figure suivante. On discutera de l'utilité de ce type de tracé dans la conclusion.



Figure 2-10 : Signature moyenne avec écart-type représenté

2.4 Segmentation

La segmentation est un des grands problèmes de la reconnaissance de formes, et c'est peut-être l'obstacle auquel nous avons été confronté le plus souvent lors de nos recherches. Ce sujet mériterait un chapitre du mémoire à lui seul, mais la manière dont on l'a traité l'intègre naturellement dans l'ensemble de la discussion sur la synchronisation.

Nous commencerons en **sous-section 1** par une présentation générale de la segmentation, puis nous suivrons en **sous-section 2** avec la description d'outils de base de notre méthode: les minima de vitesses et le filtre gaussien. Cela nous donner les bases

pour aborder avec la sous-section 3 où l'on parlera en détail de la nature des problèmes entraînés par une présegmentation. On terminera en présentant en sous-section 4 la solution adoptée, la « post-segmentation ».

2.4.1 Présentation de la segmentation de séquences

2.4.1.1 Pourquoi segmenter une signature?

En VdS, la segmentation a pour but de faciliter l'identification des éléments manuscrits, si on considère que l'information contenue dans la signature est initialement trop complexe. Par exemple, la segmentation d'une signature de 300 points (d'une durée de 3 secondes échantillonnée à 100 Hz) pourrait produire 20 segments. Identifier 20 segments au lieu de 300 points apparaît comme une tâche beaucoup plus facile à exécuter. D'autre part, si on considère que le tracé de la signature est constitué d'une séquence de mouvements de la main correspondants aux segments (portions de la signature sans changement « majeur » de direction), il n'est pas nécessaire d'aller chercher l'information située à plus petite échelle, i.e. les points (Nous allons montrer plus loin que c'est une erreur !).

2.4.1.2 Quelques méthodes connues de segmentation

Nous présentons quelques méthodes de segmentation connues. Nous reproduisons en bonne partie la revue de littérature faite dans le mémoire de Lalonde [Lalonde 94].

- **Segmentation uniforme:** les signatures sont divisées en N segment aux de durée égale, utilisée par Herbst [Herbst 82]. Cette méthode ne tient compte ni des distorsions

ni des inconstances présentes dans les signatures, et n'offre pas grand intérêt selon nous.

- **Segmentation au extremums en y:** Utilisée par Barrière [Barrière 91], cette méthode fonctionne relativement bien pour des signatures faites d'oscillations selon l'axe y, mais très mal sinon. Cette méthode de segmentation est très primitive, et ne résout aucun des problèmes fondamentaux de la segmentation dont nous parlerons plus loin.
- **Segmentation aux minima de vitesse:** De multiples observations montrent qu'il y a un lien direct entre les changements de directions du tracé et une baisse de vitesse du tracé. Comme cette technique sera utilisée dans notre méthode de segmentation, nous en parlerons plus loin.
- **Segmentation selon un modèle de génération d'écriture:** La théorie de Plamondon sur la génération d'écriture l'a mené à construire une méthode de segmentation du tracé de signature respectant la séquence de commandes motrices issues du cerveau pour générer le tracé [Plamondon 92a]. Le problème de cette méthode est que plusieurs combinaisons de commandes peuvent mener au même tracé, entraînant un problème d'unicité de la segmentation.
- **Segmentation aux points perceptuellement importants:** Dans la méthode de Brault [Brault 87a, Brault 89b], on cherche à élaborer une fonction qui mesure "l'importance perceptuelle" de chaque point: on segmente la signature aux maxima de cette fonction. Le but visé était de segmenter une courbe sans information sur la vitesse. Fischler a consacré d'importants travaux [Fischler 86] à construire une méthode de segmentation

basée sur l'importance perceptuelle des points, dans le cas de courbes quelconques. Il dit que son algorithme a été testé sur des centaines de courbes sans produire d'erreur grossière, cependant il ajoute qu'on est encore loin du jour où on arrivera à inclure dans la segmentation des informations sur la forme dont l'humain tient compte (symétrie, symboles, problèmes d'échelle...). En jetant un œil critique sur la méthode de Fischler, on peut se demander combien de segmentations différentes d'un tracé ne produisent pas « d'erreurs grossières »..

- **Segmentation par réseaux de neurone segmenteur:** Ayant pour thème l'utilisation des réseaux de neurones pour la VdS, Lalonde a tenté une segmentation par réseau de neurones [Lalonde 93,Lalonde 94a], utilisant une validation par la technique des minima de vitesse. Cette technique, comme les autres, ne disposait pas d'information à priori sur la forme à segmenter, et il semble qu'on aurait obtenu de meilleurs résultats avec une segmentation plus stable.

Ces méthodes de segmentation ont toutes le même défaut : la segmentation se fait à l'aveugle, sans connaissance sur la forme à segmenter. Cela a pour conséquence que deux formes quasi-identiques produiront en général des ensembles de segments différents.

2.4.2 Outils de base de la méthode de segmentation

La méthode de segmentation proposée ici a beau se distinguer de plusieurs façons des méthodes usuelles, elle n'en utilise pas moins deux techniques très éprouvées en VdS et en reconnaissance de formes, qui sont décrites dans cette section.

2.4.2.1 Les minima de vitesse

Une des meilleures techniques de segmentation et en même temps une des plus simples, est de segmenter aux minima de vitesse. On peut donner l'explication biomécanique suivante : en assimilant la pointe du stylo à un mobile se déplaçant sans frottement, ayant une vitesse initiale v_0 et l'impulsion de la main à une force appliquée sur le stylo correspondant à une accélération de celui-ci, on montre que la vitesse passe par un minimum si l'angle d'application de la force par rapport à la vitesse initiale est supérieur à 90 degrés et inférieur à 270 degrés:

$$\vec{v}_0 = \begin{pmatrix} v_0 \\ 0 \end{pmatrix}$$

$$\vec{a} = \begin{pmatrix} a \cos \alpha \\ a \sin \alpha \end{pmatrix}$$

Calcul de la vitesse au temps t :

$$\vec{v}(t) = \begin{pmatrix} v_0 + at \cos \alpha \\ at \sin \alpha \end{pmatrix}$$

$$|\vec{v}(t)| = v(t) = \sqrt{v_0^2 + a^2 t^2 + 2at \cos \alpha \sin \alpha}$$

$v(t)$ passe par un minimum pour :

$$\frac{dv}{dt} = 0$$

$$\Leftrightarrow t = \frac{-\sin 2\alpha}{2a}$$

Pour avoir $t > 0$, il faut :

$$\frac{\pi}{2} < \alpha < \frac{3\pi}{2}$$

Si la main applique une force sur le stylo d'angle supérieur à 90 degrés de sa vitesse actuelle, le tracé doit en théorie passer par un minima de vitesse. En détectant les minima de vitesse, on retrouve les régions de changement majeur de direction.

Comme la signature est une séquence de données, il y a une segmentation en un point i si le module de sa vitesse $v(i)$ est inférieur aux modules de vitesse du point précédent $v(i-1)$ et du point suivant $v(i+1)$:

$$i \in [2, N-1], v(i-1) < v(i) < v(i+1)$$

Étant donné que la séquence de données se présente initialement sous forme d'une séquence de position, on calcule la vitesse vectorielle par une différence finie d'ordre 2 de la dérivée du vecteur position:

$$\bar{v}(i) \equiv \frac{\bar{p}(i+2) + 8\bar{p}(i+1) - 8\bar{p}(i-1) + \bar{p}(i-2)}{18}$$

Cette approche est largement adoptée dans le milieu de la VdS car il y a une très forte corrélation entre les changements de direction du tracé et la vitesse de la main qui exécute le tracé. Aussi cette façon de segmenter donne des points de segmentation satisfaisants dans la majorité des cas.

Cette méthode comporte deux faiblesses: elle est sensible aux bruits (on sait que la dérivation d'un signal accentue les bruits) et aux hésitations de la main qui exécute le tracé. La sensibilité au bruit est éliminée par le filtrage du signal, qui est décrit ci-après.

2.4.2.2 Le filtre gaussien

Si la méthode des minima de vitesse donne d'excellents résultats, il faut préciser toutefois qu'on ne l'utilise en général pas sur des données brutes. Sur un signal non-filtré, la majeure partie des minima locaux est due à la présence de bruits dus entre autres au frottement du crayon sur la surface et aux tremblements de la main. Comme il n'existe pas de méthode simple pour faire la distinction entre les minima locaux d'amplitude faible qui sont dus aux bruits et les minima d'amplitude plus importante correspondant à des changements importants de la direction du tracé, on ne peut utiliser la méthode des minima de vitesse sans traitement préalable du signal.

Sur la figure suivante, on peut observer un signal de vitesse absolue non filtré (a), filtré (b), puis refiltré (c). Le nombre de minima de vitesse diminue avec le filtrage.

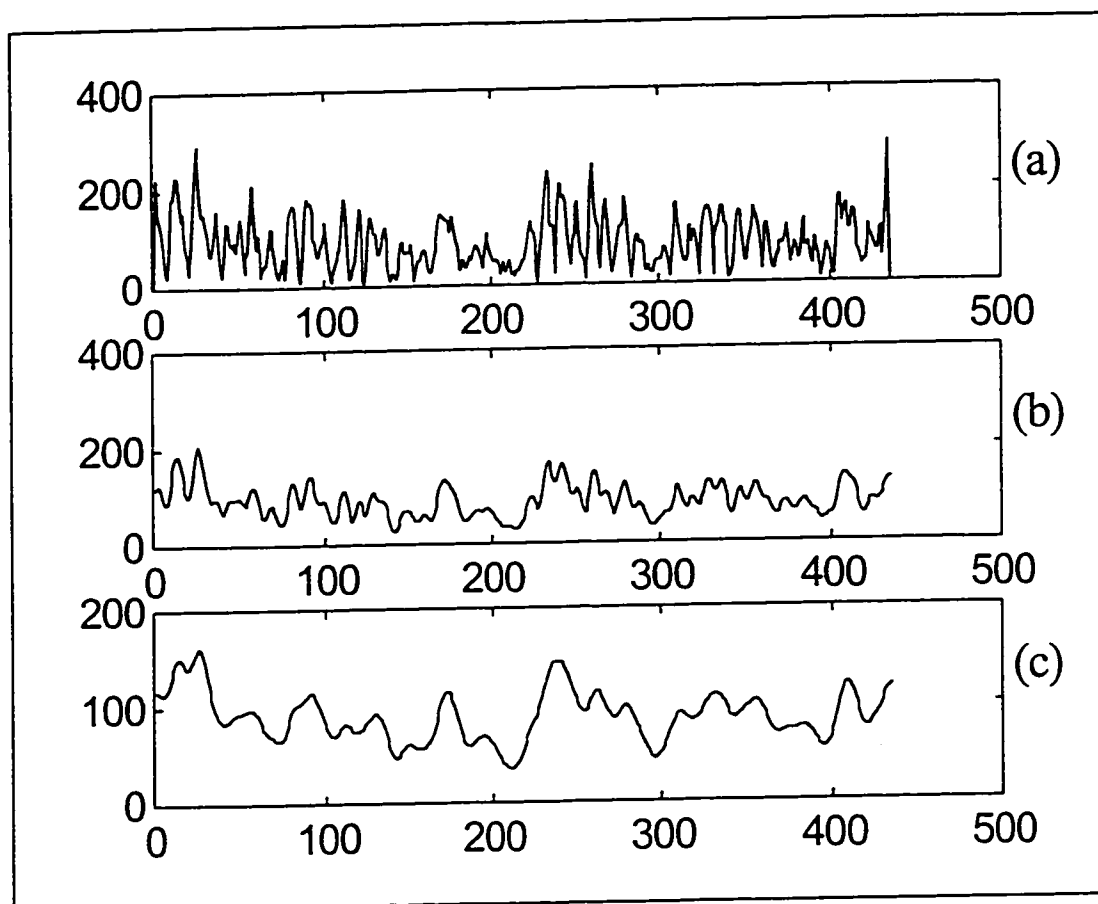


Figure 2-11 : Signal de vitesse non-filtré (a), même signal filtré (b) et (c)

Il faut donc filtrer le signal avec un filtre passe-bas pour éliminer l'énergie localisée dans les hautes fréquences, ce qui permet d'éliminer les extrêmes indésirables. Comme le résultat de la segmentation dépendra directement du choix du filtre, il faut porter une attention particulière à son choix.

Le filtre devait respecter les deux critères suivants :

- la bande-passante du filtre ne devrait pas se terminer de façon abrupte, car on ne désire pas que certaines fréquences disparaissent subitement;

- le paramètre d'échelle doit avoir la propriété de ne pas laisser apparaître d'extremum à plus large échelle qui n'aurait pas existé à une échelle plus petite.

Le filtre passe-bas parfait respecte les deux critères, mais une coupure abrupte des fréquences risque de déstabiliser la segmentation. Par contre, le filtre gaussien respecte chacune des trois conditions. En effet, Badaud [Badaud 86] a montré que le filtre gaussien est le seul filtre à moyenne mobile avec lequel il n'y a aucun risque d'apparition d'un nouvel extremum par filtrage. La coupure des fréquences est très douce et le paramètre σ permet d'ajuster le filtre. D'autre part, le filtre gaussien a été un des plus étudiés dans les années 80 dans le contexte du développement de nouvelles techniques pour le traitement d'image. Marr [Marr 82] a montré que le filtre gaussien ressemble de façon étonnante au filtrage visuel humain.

Le filtre gaussien est un filtre à moyenne mobile. Voyons maintenant comment calculer $g(x)$, résultat du filtrage de $f(x)$ avec le filtre gaussien de paramètre σ . Dans le cas, continu cela donne:

$$g(x) = \frac{1}{\sigma\sqrt{2\pi}} \int_{-\infty}^{+\infty} f(x') \exp \frac{-(x'-x)^2}{\sigma^2} dx'$$

Comme nos données sont séquentielles, nous appliquons le filtre gaussien discret.

Le calcul de $g(x)$ prend alors cette forme:

$$g(x) = \frac{1}{\sigma\sqrt{2\pi}} \sum_{i=-a}^{i=a} f(n-i) \exp \frac{-i^2}{\sigma^2}$$

La transformée de Fourier de $g(x)$ est:

$$F\{g(x)\} = \frac{\sigma}{2\pi} \exp \frac{-\sigma^2 f^2}{4}$$

Soit A la fraction de l'énergie transmise par le filtre et f_0 correspond à la fréquence d'échantillonnage (dans notre cas 100 Hz). On a:

$$A = e^{-\frac{\sigma^2 f^2}{2f_0^2}}$$

Les signatures seront filtré à $s=3$. Pour cette valeur, 83.5% de l'énergie à 20 Hz est conservée, et 98.9% de l'énergie à 5 Hz.

2.4.3 Les dangers de la pré-segmentation

L'acte de segmenter n'est bien sûr pas dangereux en soi, c'est l'utilisation du résultat qui l'est. En effet, si la segmentation a pour but la reconnaissance de la forme, elle dépend d'informations au sujet de cette forme dont elle ne dispose pas. Une méthode de segmentation est stable si pour deux signatures semblables, elle produit des segments identiques. Or la segmentation d'une séquence dépend de tant de facteurs que si elle n'est pas faite de manière concertée, il est peu probable d'arriver à un résultat cohérent pour une segmentation de plusieurs signatures.

Nous avons vu plusieurs méthodes de segmentation différentes dans la revue de littérature, et elles ont toutes le même défaut: **aucune méthode de segmentation n'est suffisamment "intelligente" pour segmenter de manière cohérente plusieurs signatures sans connaissance préalable des signatures de ces signatures.**

De manière générale, la pré-segmentation mènera à des ensembles de segments différents. A partir de là, le problème de la mise en commun des segments de même nature, de l'élimination de segment qui n'ont pas d'équivalents, la fusion de segments (à cause d'une sur-segmentation sur une signature par ex.) est très complexe à traiter, et selon nous il vaut mieux l'éviter. On peut voir sur la figure suivante un exemple de segmentation instable pour deux signatures pourtant très similaires. Il n'y a qu'un point de segmentation de différence, mais ce point de différence complexifie beaucoup les choses. La segmentation a été faite après un filtrage à $\sigma=3$ par la technique des minima de vitesse.

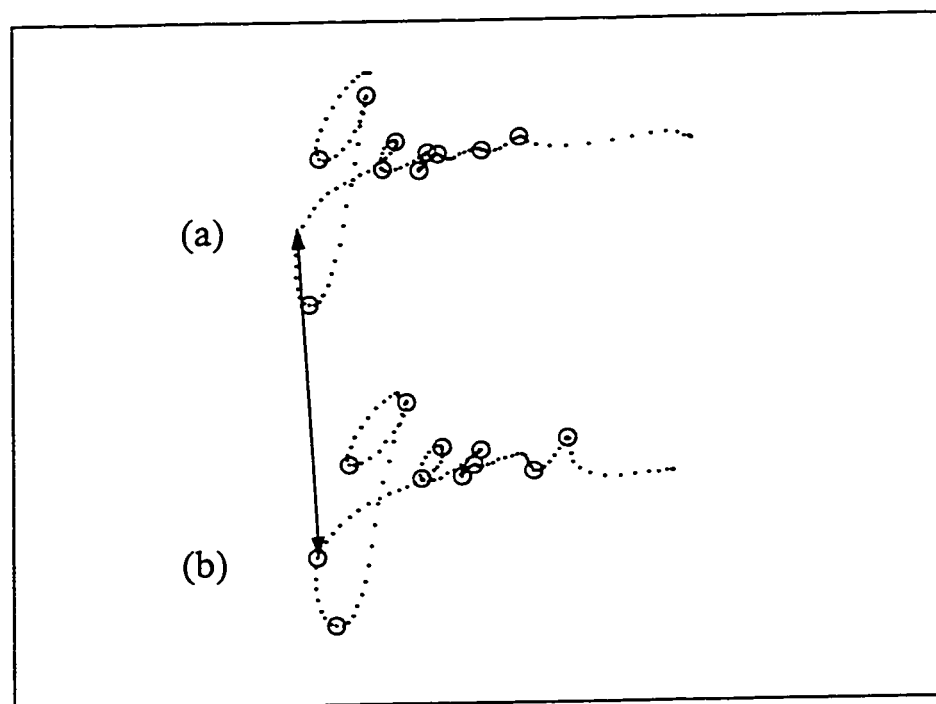


Figure 2-12 : Instabilité de la segmentation sur deux signatures très semblables. La flèche indique qu'un point de segmentation est absent sur (a) et présent sur (b).

2.4.3.1 Problème d'échelle

Le problème de l'échelle est essentiel dans la segmentation. Bien que l'on se soit rendu compte depuis longtemps que la segmentation dépendait de l'échelle [Marr 82], on n'a pas su comment incorporer de façon rigoureuse cet aspect du problème dans une méthode de segmentation de séquence. D'ailleurs, certaines méthodes ne mentionnent tout simplement pas ce problème. Les courbes utilisées pour illustrer les résultats de l'application des différentes méthodes de segmentation, et les méthodes de segmentation spécifiques aux signatures ne mettent pas toujours en évidence l'effet du choix d'une échelle pour la segmentation: puisqu'on « considère » que ces données n'ont pas une structure importante à plusieurs niveaux, on conclut qu'il y a une seule série de points de segmentation à trouver, et que s'il y en a d'autres à plus petite échelle, ils correspondent à du bruit qu'il s'agit d'éliminer, par exemple avec un filtre adéquat (passe-bas parfait, gaussien..).

Mandelbrot [Mandelbrot 82] nous démontre qu'à peu près tout objet présent dans la nature (relief terrestre, galaxie..) est de nature fractale, ce qui implique qu'il possède une structure à toute échelle. Les courbes mathématiques continues sont des exceptions que l'on ne retrouve jamais dans la nature. Ainsi, les structures détectées sur un objet quelconque dépendent entièrement de l'échelle à laquelle celui-ci est observée: si nous avions la taille d'une fourmi, nous "segmenterions" l'univers d'une façon tout à fait différente. Il est possible que pour une certaine application, on cherche à déterminer les points correspondant à une échelle ou intervalle d'échelle précis, mais les méthodes qui

considèrent implicitement que la segmentation est absolue, que le problème de l'échelle est réglé d'avance, et qui ensuite, en fonction de cela, imaginent un moyen de trouver les points "perceptuellement importants" sont selon nous arbitraires et anthropocentriques.

De toute façon, l'homme peut soit adopter un regard d'ensemble sur un objet, soit au contraire chercher à en détecter les détails les plus menus au détriment d'une vision d'ensemble. Ces modes d'observation correspondant à des échelles différentes sont complémentaires. L'homme peut donc varier les échelles d'observation à sa guise, et peut même observer plusieurs échelles à la fois.

Sur la figure suivante, on observe à 3 échelles différentes un « relief » (en réalité une marche aléatoire). Selon le point de vue, la segmentation est complètement différente.

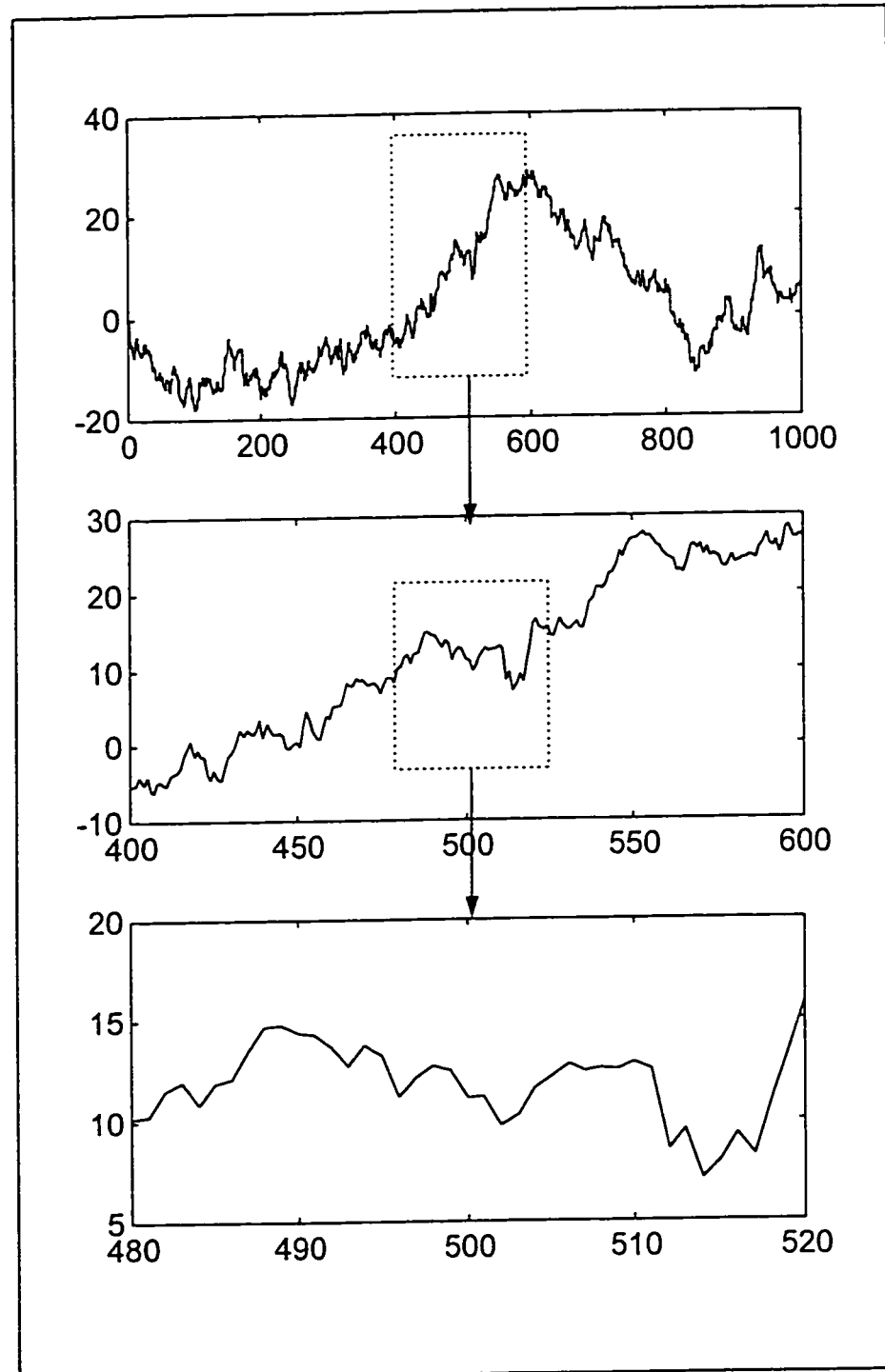


Figure 2-13 : Où segmenter ? Une question de point de vue.

Revenons à la signature. On peut représenter les points de segmentation à différentes échelles d'une méthode, selon la méthode du « space-scale filtering » ou filtrage d'échelle qu'on peut observer sur la figure suivante. Comme un segment est défini comme étant la portion de signature située entre deux points de segmentation et qu'un point de segmentation n'existe disparaît à une certaine échelle (ou valeur de σ), un segment « n'existe » qu'entre deux échelles. Lorsque pour une certaine valeur d'échelle, un point de segmentation disparaît, les deux segments adjacents fusionnent, comme on le voit sur la figure.

Deux signatures n'étant jamais identiques, les points de segmentation ne disparaîtront pas à la même échelle exactement. Les segments équivalents de deux signatures n'existeront pas entre les mêmes échelles.

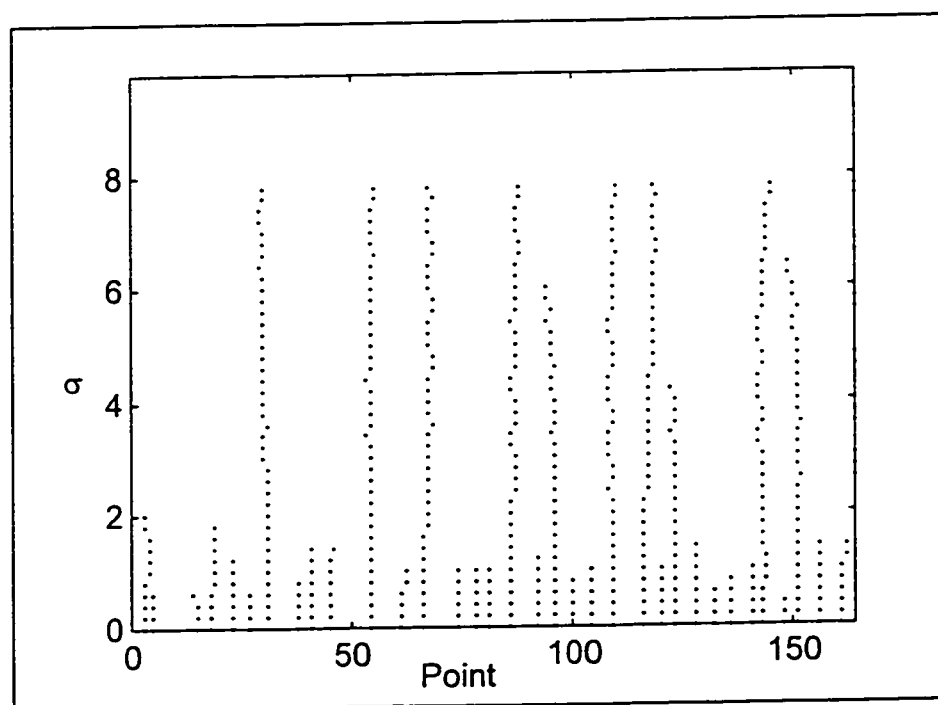


Figure 2-14 : Points de segmentation à différentes échelles

2.4.4 Approche retenue: la post-segmentation

Le filtre gaussien et la segmentation aux minima de vitesse sont des techniques éprouvées mais comme on l'a vu, elles sont impuissantes face aux problèmes impliqués par la présegmentation. Seulement dans notre cas on n'a absolument pas besoin de présegmenter (au contraire), du moins dans la phase d'identification des éléments manuscrits : avec l'alignement élastique et la méthode de synchronisation de N signatures, les points sont déjà identifiés. La segmentation ne peut donc servir dans notre cas que pour la caractérisation, si on pense que le segment représente mieux la signature que le point.

Plutôt que de présegmenter, nous allons « post-segmenter » les signatures, après la connaissance de la forme. La méthode de post-segmentation consiste à segmenter la signature moyenne en appliquant le filtre gaussien puis la méthode des minima de vitesse, comme cela est fait normalement. On détermine alors un ensemble de points de segmentation, qui vaut pour toutes les signatures. Les signatures authentiques ont pour point de segmentation les points associés aux points de segmentation de la signature moyenne. Dans le cas où le point associé d'une signature n'existe pas, on prend le point le plus proche de celui-ci.

La segmentation d'une signature-test est faite automatiquement lors de l'étape de synchronisation de N signatures, les points de segmentation des signatures étant les points identifiés aux points de segmentation de la signature moyenne.

Avec cette méthode, on a automatiquement le même nombre de segments pour chacune des signatures, et les segments sont automatiquement identifiés selon leur position dans la séquence de segments. La segmentation tient compte de la forme à segmenter : on segmente en un point que si la tendance générale des signatures authentiques est de passer par un minima de vitesse en ce point.

Sur la figure suivante, on peut voir le résultat de la segmentation concertée de quelques signatures. Notons qu'on a tenu compte de toutes les signatures de référence (30 au total) pour arriver ce résultat. Le résultat aurait été probablement plus précis en ne prenant que les six signatures en question.

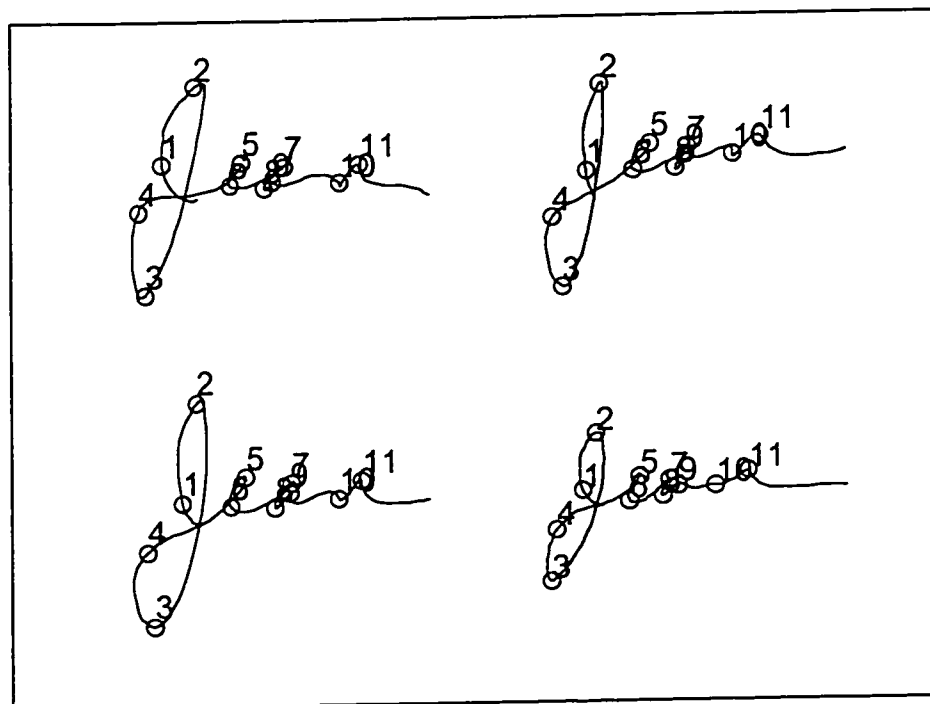


Figure 2-15 : Segmentation à $\sigma=3$ par la méthode de post-segmentation

2.5 Conclusion

Dans ce chapitre, nous avons traité du problème de l'identification des éléments manuscrits. On peut voir ci-dessous un résumé de la séquence des opérations à effectuer pour générer les classes d'éléments manuscrits et l'évaluation des paramètres descriptifs de ces classes. L'étape de segmentation n'est nécessaire que si on considère que l'élément manuscrit fondamental est le segment et non le point.

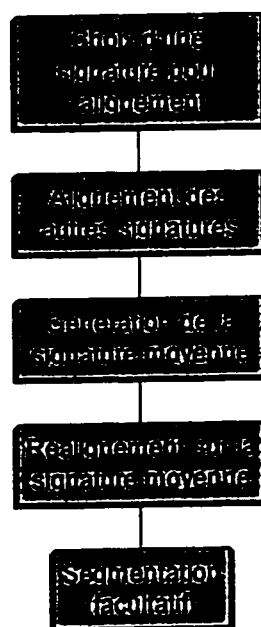


Figure 2-16 Prototypé de l'apprentissage pour l'identification des éléments manuscrits

3. MODÉLISATION STATISTIQUE DE LA VDS

3.1 Introduction

On a déjà souligné au chapitre 1 que notre ignorance devant la classe des signatures non-authentiques était un des problèmes fondamentaux de la VdS. En effet, si on cherche à classer une signature entre deux classes, une connaissance minimale de ces deux classes est indispensable.

Ce chapitre est essentiellement une tentative de modélisation des deux classes authentiques (ω_1) et non-authentiques (ω_2) qui s'appuie sur la reconnaissance de formes (RF) statistique. Dans ce cadre, le problème se pose ainsi : étant donnée une signature décrite par un vecteur de caractéristiques X , on cherche à évaluer les probabilités d'appartenance de ce vecteur aux classes ω_1 et ω_2 .

Il est nécessaire d'établir quelques hypothèses simplificatrices sur les classes ω_1 et ω_2 dans le but d'évaluer les probabilités d'appartenance. Une fois ces hypothèses établies et justifiées, on s'en servira pour dériver des résultats théoriques essentiels sur la VdS : paramètres significatifs pour l'évaluation d'un coefficient de difficulté d'imitation d'un scripteur, comparaison de systèmes de VdS basés sur des caractéristiques globales et locales.

En **section 2**, après un survol des approches possibles, l'approche de RF statistique est appliquée à la VdS. En posant certaines hypothèses, les distributions de

probabilité des classes authentiques et non-authentiques sont évaluées. On se sert ensuite de ces distributions de probabilité pour générer des faux aléatoires. En **section 3**, on se sert des résultats de la section précédente pour déterminer les paramètres essentiels d'un scripteur. Puis on montre comment les approches locales et globales (i.e. l'extraction de vecteurs de caractéristiques locales et globales) à la VdS peuvent être représentées par notre approche.

3.2 Développement du modèle

3.2.1 Approches possibles

Il existe au moins deux classes d'approche possible pour modéliser la VdS dynamique. On peut soit s'intéresser à la nature même du signal de signature, soit considérer la VdS comme un problème de reconnaissance des formes (RF).

Dans le premier cas, on considère la signature du point de vue de sa génération biomécanique. Selon cette approche, les signatures dites « authentiques » auraient des propriétés intrinsèques non possédées par les imitations. Les signatures peuvent être modélisées par une série de gaussiennes [Leclerc 92], d'où on en déduit une méthode de segmentation s'appuyant sur des bases théoriques [Plamondon 92a] dont on a discuté dans le chapitre précédent.

Dans le deuxième cas, on aborde le problème de classification de signatures par une des multiples approches de la RF : classification automatique, discrimination

fonctionnelle, méthodes statistiques, k plus proches voisins, méthodes stochastiques, connexionistes et structurelles. Une revue de ces approches est donnée dans [Belaïd 92].

Si on cherche à prédire la performance d'un système de VdS en fonction de certains paramètres fondamentaux, l'approche statistique semble toute indiquée. La suite de ce chapitre consistera à appliquer cette approche à la VdS. Pour une présentation de l'approche statistique à la RF, nous recommandons les ouvrages suivants : [Ripley 96, Fukunaga 72].

3.2.2 L'approche statistique à la VdS

Étant donnée une signature S d'origine inconnue, le système de VdS doit classer cette signature parmi une des classes ω_1 et ω_2 . Les opérations préliminaires de prétraitements et d'analyse ont pour but l'extraction d'un vecteur de caractéristiques X représentant S le mieux possible. L'étape fondamentale de reconnaissance (comparaison, décision) vient une fois ces opérations préliminaires effectuées.

Pour modéliser l'étape de reconnaissance, l'application de méthodes statistiques suppose que la représentation des formes admet un modèle probabiliste. En VdS comme dans d'autres domaines de la RF, ce point est très discutable : il est très difficile, de prouver que des signaux générés « naturellement », tel que le signal de parole ou la signature, respectent des règles statistiques. Cependant, on peut voir les statistiques comme un outil qui nous permet d'avoir une représentation simplifiée d'un problème complexe. D'ailleurs, l'application des HMM (Hidden Markov model) à la reconnaissance de la parole sont parmi les méthodes les plus répandues, et on retrouve

également des applications de méthodes statistiques en VdS [Nelson 94]. Nous allons également poser le problème de la VdS en termes statistiques.

Étant donné une signature S décrite par un vecteur X de caractéristiques, un système de VdS statistique cherche à évaluer la probabilité que la signature S soit authentique, i.e. : $P(\omega_1|S) \approx P(\omega_1|X)$. Appliquons la formule d'inversion de Bayes à $P(\omega_1|X)$:

$$P(\omega_1|X) = \frac{P(X|\omega_1)P(\omega_1)}{P(X)} \quad (3-1)$$

En décomposant $P(X)$ sur ω_1 et ω_2 , (qui forment une partition complète de l'espace E des signatures possibles), on a :

$$P(\omega_1|X) = \frac{P(X|\omega_1)P(\omega_1)}{P(X|\omega_1)P(\omega_1) + P(X|\omega_2)P(\omega_2)} \quad (3-2)$$

Pour évaluer $P(\omega_1|X)$, il faut donc connaître :

- $P(X|\omega_1)$: probabilité qu'un scripteur (= signataire authentique) produise une signature décrite par X .
- $P(X|\omega_2)$: probabilité qu'un imitateur produise une signature décrite par X .
- $P(\omega_1)$: probabilité à priori qu'une signature soit authentique. Notons que $P(\omega_2) = 1 - P(\omega_1)$.

L'équation précédente a permis de formuler le problème, qui est d'évaluer les 3 quantités ci-dessus. Les sous-sections suivantes, sont consacrées à l'évaluation de $P(\omega_1)$, $P(X|\omega_1)$ et $P(X|\omega_2)$.

3.2.3 Estimation de $P(\omega_1)$

On doit en général faire une estimation de $P(\omega_1)$, la proportion de signatures authentiques entrées dans le système de VdS, qui dépend du cadre dans lequel ce système est utilisé (carte de crédit, accès dans des locaux à haute sécurité ou éventuellement utilisation avec Internet). Bien évidemment, les imitations sont beaucoup plus rares que les signatures authentiques dans la pratique. Mais lorsqu'on teste un système de VdS, on a au contraire une proportion d'imitations largement supérieure aux proportions « réelles ». Dans notre cas, nous avons autant d'imitations que de signatures authentiques.

La règle de décision bayésienne est d'attribuer à S la classe la plus probable. Soit $d(X)$ la fonction de décision, alors $d(X)=\omega_1$ si :

$$\begin{aligned} \frac{P(w_1|X)}{P(w_2|X)} &> 1 \\ \Leftrightarrow \frac{P(X|w_1)P(w_1)}{P(X|w_2)P(w_2)} &> 1 \\ \Leftrightarrow \frac{P(X|w_1)}{P(X|w_2)} &> \frac{1 - P(w_1)}{P(w_1)} \end{aligned} \quad (3-3)$$

On voit d'après cette équation que la limite de décision, qui est X tel que $P(X|\omega_1)=P(X|\omega_2)$, dépend des proportions relatives des deux classes.

De façon plus générale, à chaque type d'erreur on rattache un coût. Le but n'est alors pas de prendre la décision minimisant le risque d'erreur, mais de prendre la décision minimisant le risque. En VdS, les coûts C_1 et C_2 rattachés aux erreurs de type I (rejeter

une signature authentique) ou de type II (accepter un faux) dépendent entièrement du champ d'application. Pour l'utilisation avec des cartes de crédit, le client est peu tolérant au rejet de sa signature (même si elle a été mal exécutée), alors qu'il n'est pas essentiel de détecter absolument tous les faussaires. Le rapport des coûts C_1 et C_2 sera alors moins élevé que dans une application à haute sécurité.

Les coûts moyens $C_1(X)$ et $C_2(X)$ sont :

$$C_1(X) = P(\omega_2 | X) \cdot C_1$$

$$C_2(X) = P(\omega_1 | X) \cdot C_2$$

On cherche la décision minimisant le coût moyen. La décision $d(X)$ est ω_1 si :

$$\begin{aligned} C_1(X) &< C_2(X) \\ \Leftrightarrow \frac{P(X|\omega_1)}{P(X|\omega_2)} &> \frac{p(\omega_2) \cdot C_1}{P(\omega_1) \cdot C_2} \end{aligned}$$

Comme nous cherchons à modéliser de façon très générale la VdS, nous ne nous embarrassons pas des estimations de probabilité à priori et de coût d'erreur. Le développement précédent n'est mis qu'à titre indicatif. Pour la suite nous allons considérer que la décision est simplement d'attribuer à X la classe la plus probable.

3.2.4 Estimation de $P(X|\omega_1)$

On cherche ici à estimer $P(X|\omega_1)$. Pour cela, nous avons à faire quelques hypothèses.

Unicité de la classe ω_1 : La première hypothèse concerne l'unicité de la classe authentique. Cette hypothèse est discutable car il est connu que les gens peuvent avoir plusieurs manières de signer. En général, ces différentes manières de signer n'affectent pas le signal au complet : un scripteur pourra signer parfois avec parfois sans prénom, par exemple. Pour l'utilisation dans le cadre d'un système de VdS, on est en droit de demander aux scripteurs de n'adopter qu'une seule façon de signer, afin de simplifier la reconnaissance. C'est d'ailleurs ce que nous avons demandé aux scripteurs, lors de la phase d'expérimentation. Nous avons cependant observé au moins un scripteur qui n'a pas respecté cette contrainte, comme on peut le voir sur la figure suivante. Le scripteur commençait sa signature de deux façons différentes, avec des proportions illustrées sur la figure.

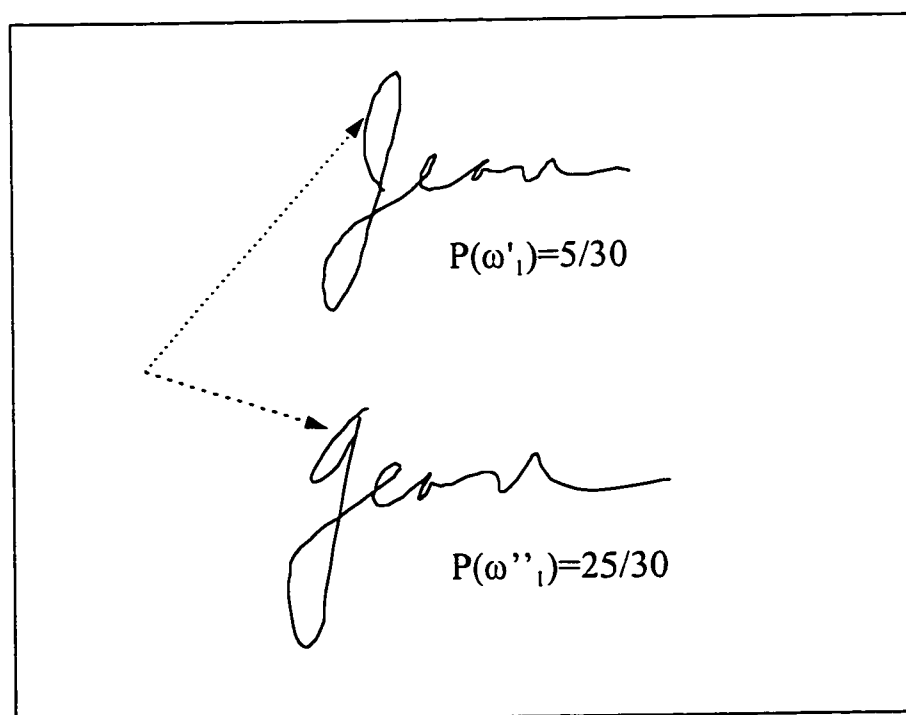


Figure 3-1 : Variation dans le début de la signature d'un scripteur

Pour ce scripteur, on aurait :

$$P(X|\omega_1) = P(X|\omega_1^*) \cdot \frac{5}{30} + P(X|\omega_1^*) \cdot \frac{25}{30}$$

Ce problème de non-unicité de la classe authentique ne concerne que certaines portions de la signature de certains scripteurs. Ne pas le considérer n'a pour conséquence qu'une variabilité plus grande des portions affectées. Pour les besoins de la modélisation, il n'est pas nécessaire de le considérer. On fait donc l'hypothèse suivante :

Hypothèse 1 : la classe ω_1 n'est pas divisible en sous-classes.

Vecteur X : Distributions des caractéristiques locales

Les caractéristiques décrivant une signature peuvent être tant globales (durée de la signature, position du centre de gravité) que locales (vitesse, accélération ponctuelle ou par segment). On a vu au chapitre précédent qu'ayant synchronisé une signature S par rapport à une signature moyenne S_M de taille n , on peut décrire S par un vecteur de caractéristiques locales $X=[x_1, x_2, \dots, x_n]^T$, chaque x_i correspondant à la caractéristique d'un élément manuscrit identifié. Comme certains points de S_R n'ont pas de points associés sur S , certaines dimensions du vecteur de caractéristique seront manquantes. Nous supposons ici pour simplifier que le vecteur de caractéristiques est complet. D'ailleurs, généralement 98% des points de la forme de référence sont identifiés sur une signature-test.

Sur un ensemble de signatures de référence, on peut, pour la $i^{\text{ème}}$ classe d'éléments manuscrits, évaluer la moyenne μ_i et l'écart type σ_i d'une certaine caractéristique (par ex. la vitesse horizontale v_x) : $M=[\mu_1, \dots, \mu_n]^T$, $\Sigma=[\sigma_1, \dots, \sigma_n]^T$.

Soit une signature S représentée par un vecteur de caractéristiques $X=[x_1, x_2, \dots, x_n]^T$.

La double hypothèse fondamentale que l'on fait est que :

- **Hypothèse 2 : Les x_i suivent une loi de probabilité gaussienne : $x_i \sim N(\mu_i, \sigma_i)$**
- **Hypothèse 3 : Les x_i sont indépendants : $\text{cov}(x_i, x_j) = 0$ pour $i \neq j$**

Nous ne chercherons pas à donner une justification biomécanique de cette loi. Il paraît naturel que chacune des parties d'une signature suive, en vitesse, position, ..., une certaine moyenne et une certaine variance, et que la distribution de probabilité va décroissant symétriquement lorsqu'on s'éloigne de la moyenne, ce qui est le cas dans le modèle gaussien.

De fait, si on prend un ensemble de signatures d'apprentissage pour évaluer les moyennes et écarts-type locaux d'une certaine caractéristique (par ex. v_x , v_y , a_x ...), et un ensemble de signatures de tests, on peut vérifier si les caractéristiques locales de ces signatures-test, une fois centrées réduites, suivent une loi normale. Sur la figure ci-dessous, nous avons, pour un certain scripteur, observé les distributions des valeurs centrées réduites pour les caractéristiques locales v_x et v_y . Quelque soit le scripteur, on observe la même distribution. Il semble bien que la loi normale est respectée.

Note : Pour une caractéristique locale x_i de moyenne μ_i et d'écart σ_i , et qui prend une valeur x_i , la valeur centrée réduite z_i est :

$$z_i = \frac{x_i - \mu_i}{\sigma_i}$$

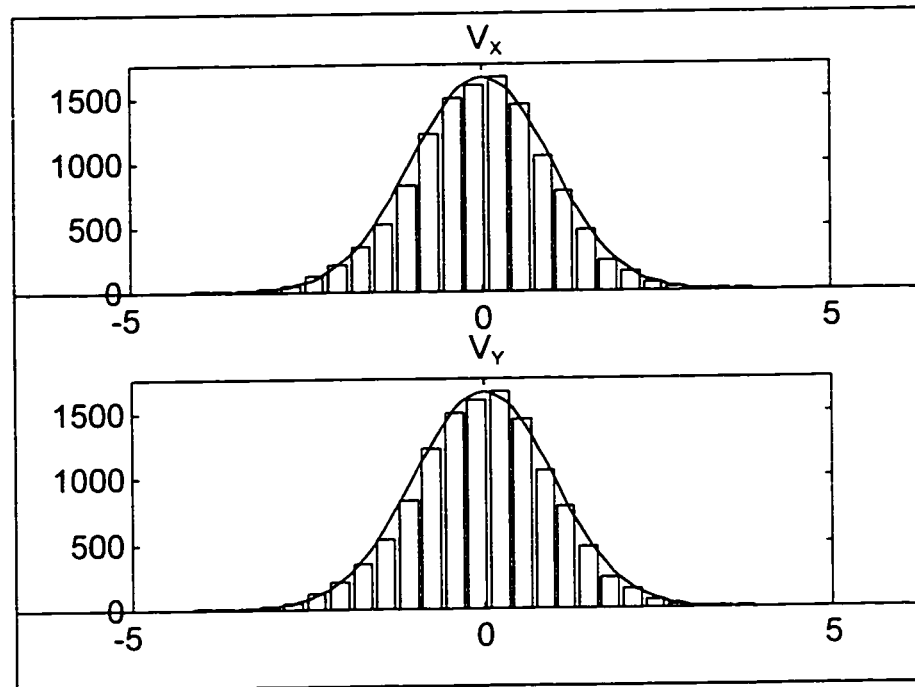


Figure 3-2 : Distribution de caractéristiques locale v_x et v_y et loi normale.

L'hypothèse 3 d'indépendance est très fréquemment faite en statistique, car elle permet de simplifier énormément les calculs. Voyons ce qu'elle signifie pour les signatures.

Pour des caractéristiques locales, l'hypothèse d'indépendance peut se justifier, si on prend des points suffisamment distants. En effet, lorsqu'on échantillonne à 100 Hz, on s'attend à ce que les caractéristiques x_i et x_{i+1} de deux points consécutifs i et $(i+1)$ soient corrélées, i.e. si $x_i > \mu_i$, alors il y a de fortes chances que $x_{i+1} > \mu_{i+1}$. Si un

segment du tracé est fait à une vitesse supérieure à la moyenne, alors l'ensemble des points du segment auront des vitesses supérieures à la moyenne.

On s'attend à ce que la corrélation diminue à mesure que la durée entre deux points augmente. Il suffit donc de choisir une fréquence d'échantillonnage suffisamment grande pour que l'indépendance soit vérifiée. Pour quel écart de durée deux points peuvent-ils être considérés comme indépendants ?

Afin d'évaluer cette distance minimale, nous avons évalué les coefficients de corrélation moyens pour un intervalle de durée entre caractéristiques allant de 1 à 50 centièmes de secondes. Soit la signature d'un scripteur décrite par n éléments manuscrits. L'estimation du coefficient de corrélation moyen pour des éléments manuscrits distants de d centièmes de secondes est :

$$corr(d) = \frac{1}{n-d} \sum_{i=1}^{i=n-d} \frac{cov(x_i, x_{i+d})}{\sqrt{var(x_i) \cdot var(x_{i+d})}} \quad (3-4)$$

On peut observer sur la figure suivante ces coefficients de corrélation, pour les caractéristiques locales v_x et v_y du scripteur no 3.

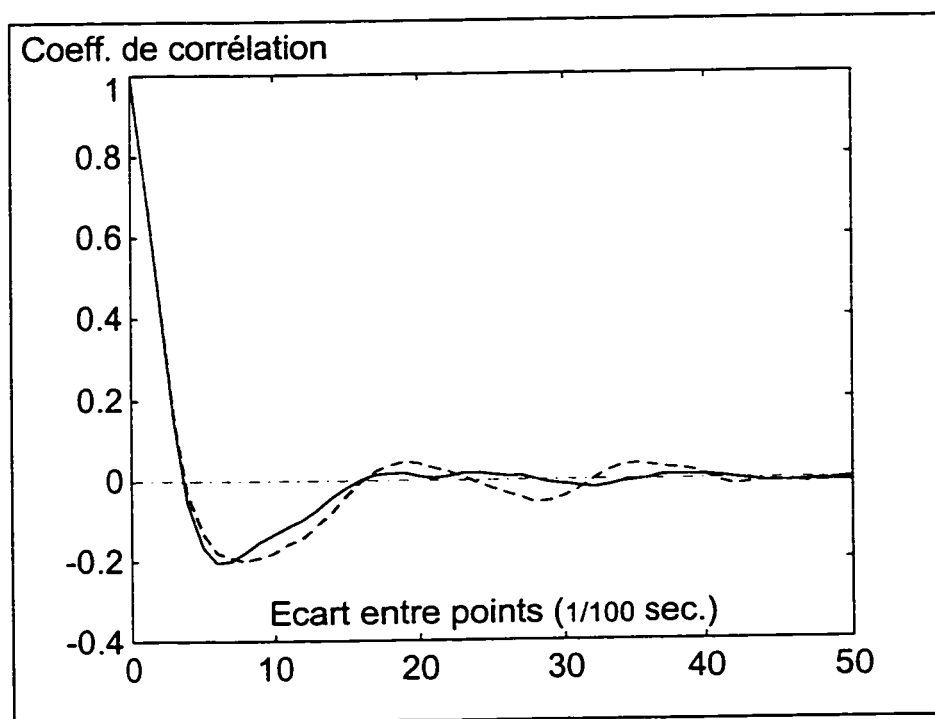


Figure 3-3 : Corrélation selon la distance entre points, pour le scripteur no 3.

(trait continu : v_x , hachuré : v_y)

Les courbes observées sur la figure précédente ont le même aspect quelque soit le scripteur. On peut observer que pour approximativement 5 centièmes de secondes d'écart, le lien entre les caractéristiques de vitesse de deux points devient à peu près nul. Il est intéressant de remarquer qu'entre 5 et 15 centièmes de secondes, il y a une phase de corrélation négative, observable chez tous les scripteurs. On peut interpréter ce phénomène ainsi : si un scripteur a eu une vitesse supérieure (inférieure) à la moyenne habituelle sur cette portion, il aura une légère tendance, entre 5 et 20 centièmes de

secondes plus loin, à avoir une vitesse inférieure à la moyenne (et inversement). Actuellement nous n'avons pas d'explication valable à fournir sur ce phénomène.

Nous avons fait une comparaison des courbes de corrélations en fonction de l'écart entre points pour les signaux de position, de vitesse et d'accélération (figure suivante). La corrélation de position reste élevée pour un laps de temps important. On peut expliquer ce phénomène ainsi : pour deux points distants de 5 centièmes de secondes, la corrélation de vitesse est à peu près nulle. On peut alors assimiler la valeur que prend la vitesse locale centrée réduite à une marche aléatoire gaussienne. La position, soit l'intégration de la vitesse, correspond à l'intégration d'une marche aléatoire. Or il y a une corrélation positive entre deux points distants de l'intégration d'une marche aléatoire.

A l'inverse, l'accélération, dérivation de la vitesse, accentue les hautes fréquences porteuses de bruit. C'est pourquoi le passage par zéro de la corrélation d'accélération est à approximativement 3 centièmes de secondes d'écart contre 5 pour la corrélation de vitesse.

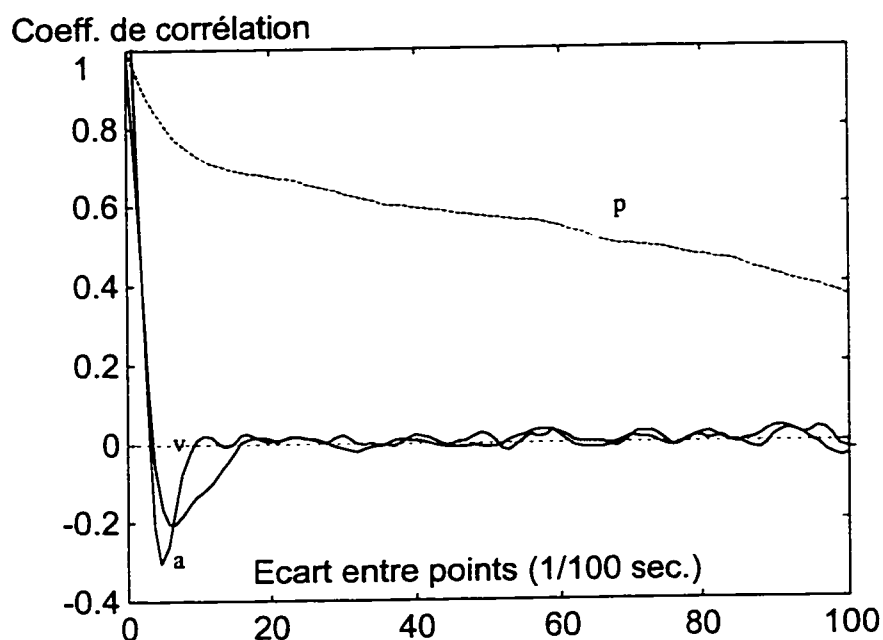


Figure 3-4 : Corrélation de position (p), vitesse (v) et accélération (a)

Les caractéristiques locales x_i suivant une loi normale et étant indépendantes (pour une fréquence d'échantillonnage suffisamment grande), $P(X|\omega_1)$ suit la distribution de probabilité suivante :

$$P(X | \omega_1) = \prod_{i=1}^n \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x_i - \mu_i)^2}{2\sigma_i^2}} \quad (3-5)$$

Afin de vérifier visuellement le modèle ci-dessus, on l'a utilisé comme méthode de génération de signatures authentiques. En évaluant la moyenne et l'écart-type des caractéristiques locales v_x et v_y (on se sert pour cela des techniques élaborées au chapitre précédent) d'un scripteur, on évalue $p(X|\omega_1)$. A partir de cette distribution de probabilité

on génère aléatoirement des vecteurs de vitesse respectant la distribution de probabilité du scripteur. Une intégration de la vitesse nous donne la séquence de positions de la signature générée. Sur la figure suivante, on donne un exemplaire de 4 signatures authentiques générées aléatoirement et de 4 signatures authentiques à titre de comparaison.

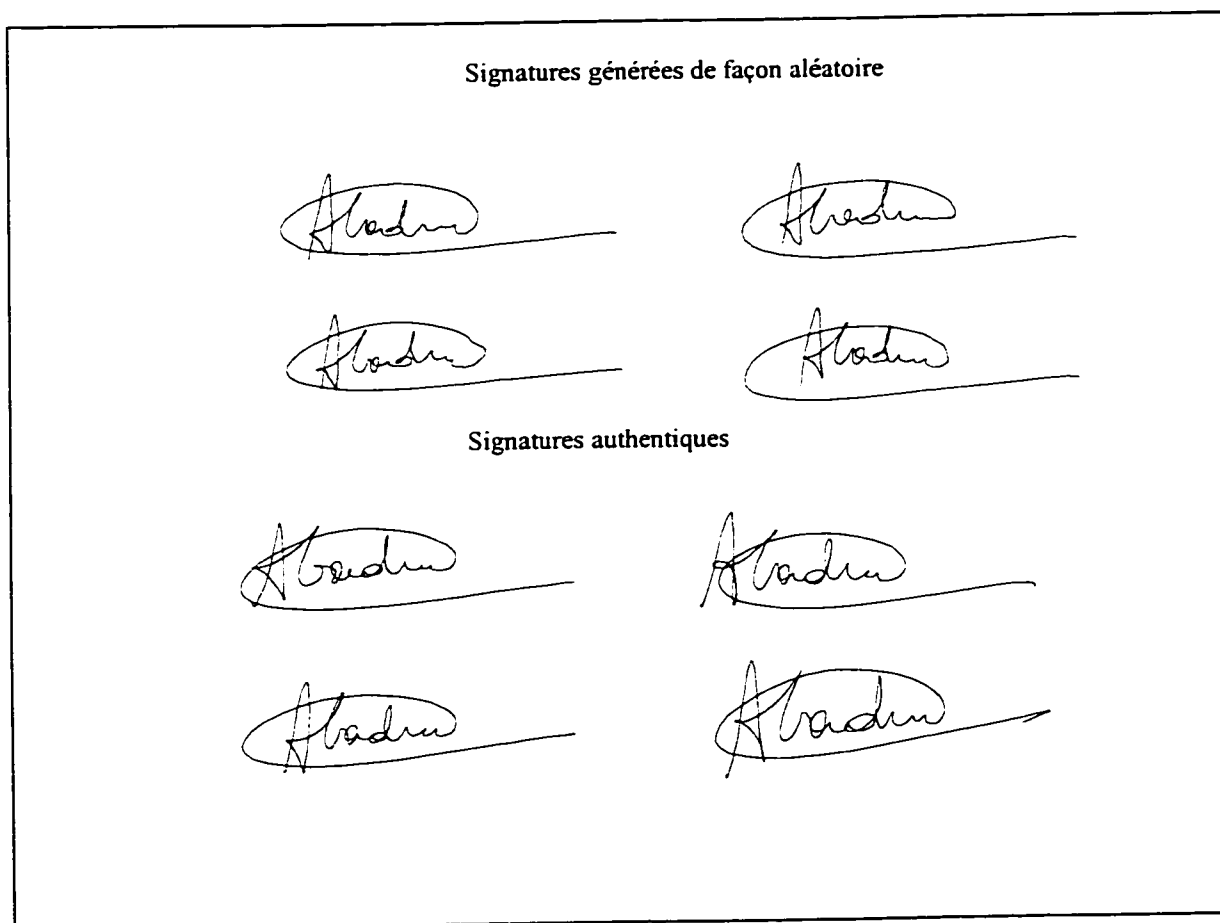


Figure 3-5 : 4 signatures générées aléatoirement et 4 signatures authentiques

3.2.5 Estimation de $P(X|\omega_2)$

Si on parvient, en faisant certaines hypothèses, à estimer $P(X|\omega_1)$, l'estimation de $P(X|\omega_2)$ paraît beaucoup plus difficile. De fait, on considère la classe non-authentique comme une classe « fantôme » en VdS, puisqu'elle est en théorie constituée d'une infinité d'imitateurs que l'on ne peut connaître, et de tous les « accidents » qui peuvent survenir : par exemple la signature d'un autre scripteur par erreur. On peut, pour une certaine expérience, recueillir un nombre statistiquement valable (disons 30) imitations de chacun des imitateurs, mais qui nous dit que ces imitateurs, qu'ils soient 3, 10, ou 30, sont représentatifs de la population des imitateurs ?

Les « accidents » engendrent des signatures très différentes et représentent selon nous un problème mineur. La grande difficulté en VdS réside dans le fait de distinguer des objets très similaires : distinguer un faux d'un vrai est un peu comme distinguer deux frères jumeaux. Par contre distinguer deux signatures de nature différente (par ex. la signature de « Marc » et celle de « Pierre ») est, du moins en VdS dynamique et avec la méthode développée ici, trivial. Nous allons donc identifier la classe ω_2 à la classe des imitations.

Avant de chercher à modéliser la distribution de la classe des imitateurs, modélisons d'abord la distribution d'un seul imitateur. Prenons un imitateur bien entraîné qui a eu le temps désiré pour se pratiquer à imiter de façon dynamique la signature d'un scripteur. Pour nos expérimentations, les 3 imitateurs disposaient d'image des signatures authentiques, pouvaient observer la signature apparaître à l'écran en temps réel, et calquer

leur mouvement dessus. Ces imitateurs pouvaient même poser une feuille de papier ou était représentée la signature authentique.

Un tel imitateur a une signature aussi stabilisée qu'un scripteur. De fait, les écarts-type moyens pour les caractéristiques locales v_x et v_y des imitateurs ne sont en général pas plus élevés que ceux des scripteurs. On peut donc appliquer le même modèle de distribution de probabilité à un imitateur qu'à un scripteur. Soit un certain imitateur I , alors :

$$P(X|I) = \prod_{i=1}^n \frac{1}{\sigma_I \sqrt{2\pi}} e^{-\frac{(x_i - \mu_{li})^2}{2\sigma_{li}^2}} \quad (3-6)$$

Sur la figure ci-dessous, nous avons, pour un certain imitateur, observé les distributions des valeurs centrées réduites pour la caractéristique locale v_x . Quelque soit l'imitateur, on observe la même distribution. Il semble bien que la loi normale est respectée.

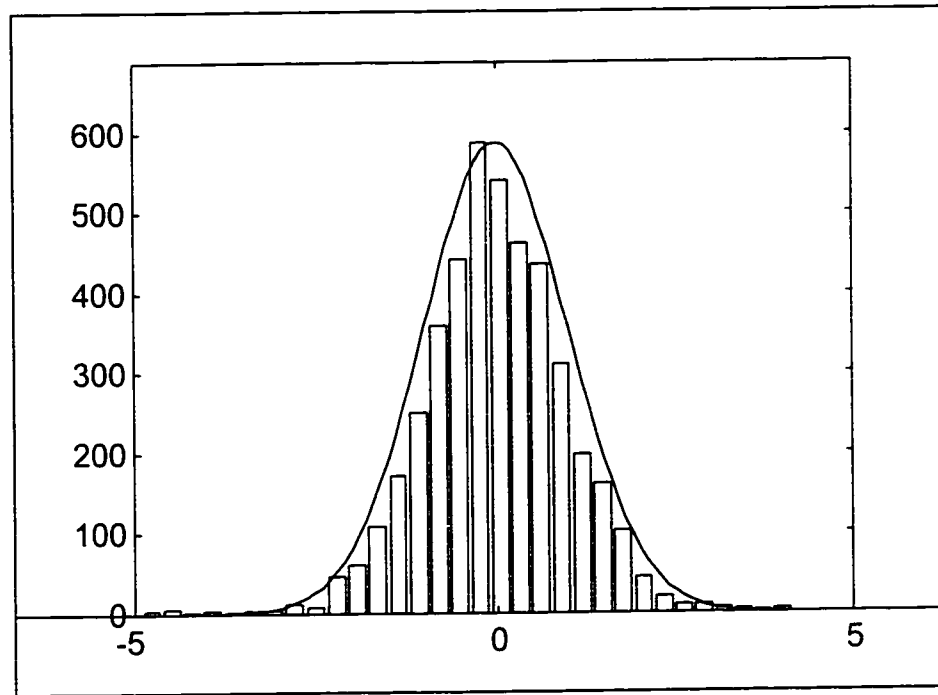


Figure 3-6 : Distribution des caractéristiques locales v_x et loi normale

Supposons une population d'imitateur de taille N tendant vers l'infini. Soit $I_1, \dots, I_N$ ces imitateurs. L'ensemble des sous-classes correspondant à ces imitateurs constituent la classe ω_2 non-authentique. Chaque classe, ou chaque imitateur, a une chance égale d'apparaître. On a alors :

$$\begin{aligned}
 P(X|\omega_2) &= \lim_{N \rightarrow \infty} \sum_{j=1}^N P(X|I_j)P(I_j) \\
 &= \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{j=1}^N P(X|I_j)
 \end{aligned}
 \tag{3-7}$$

Comme la sommation de variables normales est une variable normale, $P(X|\omega_2)$ suit une loi normale. Il nous reste à évaluer la moyenne et l'écart-type de cette loi normale. L'évaluation de l'écart-type est un problème considérable et ce mémoire ne prétend y

apporter qu'un début de réponse. Pour l'évaluation de la moyenne, on va montrer, en faisant l'hypothèse suivante, que la moyenne de $P(X|\omega_2)$ est égale à la moyenne $P(X|\omega_1)$.

Hypothèse 4 : Soit (M_j, Σ_j) les paramètres de la classe authentique. Pour tout imitateur I_j ayant une distribution de probabilités de paramètres (M_j, Σ_j) , il existe un imitateur « miroir » I_k de paramètres $(M - (M_j - M), \Sigma_j)$. Cet imitateur I_k dévie autant de la classe authentique que l'imitateur I_j , mais de façon exactement opposée.

Cette hypothèse, pour étonnante qu'elle puisse paraître, est en fait contenue implicitement dans la majorité des systèmes de VdS. Elle signifie qu'il n'y a pas de polarisation dans la population des imitateurs. Qu'une caractéristique x_i d'une signature-test dévie par « défaut » ($x_i < \mu_i$) ou par « excès » ($x_i > \mu_i$) de la moyenne, elle sera pénalisée également, avec un coût de type $|x_i - \mu_i|^N$.

Que cette absence de polarisation supposée soit « appliquée » dans les systèmes de VdS (en fait il est plus simple de l'appliquer que de ne pas l'appliquer) ne signifie pas nécessairement qu'elle soit tout à fait vraie. Il est par exemple possible que pour un scripteur particulièrement lent, la plupart des imitateurs signent avec une vitesse excessive même s'ils sont très entraînés, puisque la majorité d'entre eux signent naturellement plus vite.

Sur la figure suivante, on a représenté les deux cas possibles. Supposons les signatures représentées par un vecteur de caractéristiques 2-D. Au centre de chacune des figures, on a des signatures de classe ω_1 . A l'extérieur, on a des signatures de classe ω_2 .

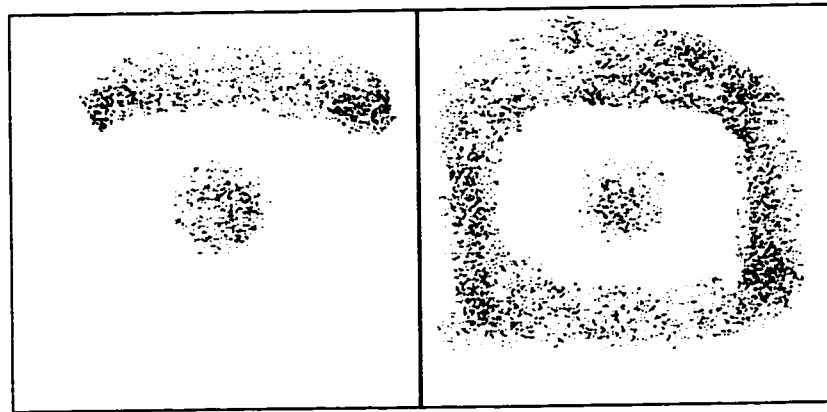


Figure 3-9 : A gauche, classe ω_2 polarisée. A droite, classe ω_2 non-polarisée.

Les imitateurs sont-ils « polarisés » ou non ? Nous n'avons pas les moyens de répondre à cette question, si tant qu'on puisse lui donner une réponse, puisque les imitations récoltées pour l'expérimentation ne proviennent que de 3 imitateurs différents. Pour les besoins du modèle, nous allons faire la supposition que la population des imitateurs n'est pas polarisée.

Pour justifier « visuellement » cette hypothèse, nous avons fait l'expérience suivante. A partir de la classe authentique de moyenne M et d'un imitateur de moyenne M' , nous avons généré des signatures hybrides de caractéristiques M_H selon l'équation suivante :

$$M_H = M + \alpha(M' - M) \quad (3-8)$$

Pour $\alpha=1$, la signature hybride M_H est égale à l'imitation « miroir », pour $\alpha=0$ elle est égale à la signature authentique, pour $\alpha=-1$, elle correspond à l'imitation M' . Lorsque α décroît sur $[1,0]$, la signature hybride générée a des caractéristiques se rapprochant d'abord de celles du scripteur. Sur $[0,-1]$, les caractéristiques se rapprochent de la

signature miroir. Pour 0.5 et -0.5, la signature hybride est de classe « totalement indéfinie », puisqu'elle comporte en proportion égale des caractéristiques d'un imitateur et du scripteur.

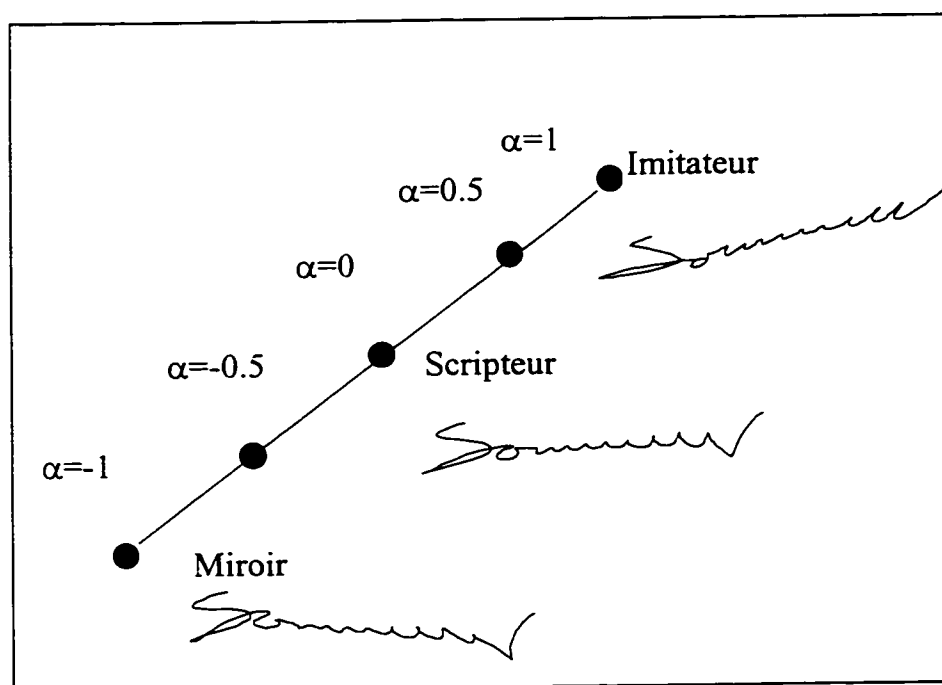


Figure 3-10 : Position des signatures hybrides par rapport au scripteur

Sur la figure suivante, on a représenté, pour deux scripteurs et deux imitateurs, les signatures hybrides pour différentes valeurs du paramètre α . Le résultat est intéressant.

Pour « Sommelet », l'imitateur a une signature qui « monte » et qui est inclinée vers la droite alors que le scripteur a une signature de hauteur stable. A mesure que α décroît, la signature hybride monte de moins en moins et est de moins en moins inclinée. Lorsque α devient négatif, cette signature se met carrément à baisser et devient inclinée vers la gauche. L'imitation miroir paraît tout à fait concevable.

En extrapolant un peu, un graphologue pourrait dire (sous réserve, l'auteur n'ayant pas de formation graphologique !) que le scripteur étant de tempérament équilibré (signature de hauteur stable), l'imitateur d'origine ($\alpha=1$) est de tempérament extraverti et emporté (signature montante, penchée vers l'avant) et que son opposé, l'imitateur « miroir », est de tempérament introverti et calme.

Pour le scripteur « Jean », on peut voir que la forme du « J » est moins prononcée pour l'imitateur que pour le scripteur. L'imitation miroir a un J beaucoup plus prononcé. On peut s'amuser à relever plusieurs détails qui se transforment lorsqu'on fait varier α . On doit cependant reconnaître qu'il y a certains aspects de la signature miroir qui choquent un peu l'œil. Cependant pour des valeurs de α jusqu'à 0.75, la signature hybride semble tout à fait concevable.

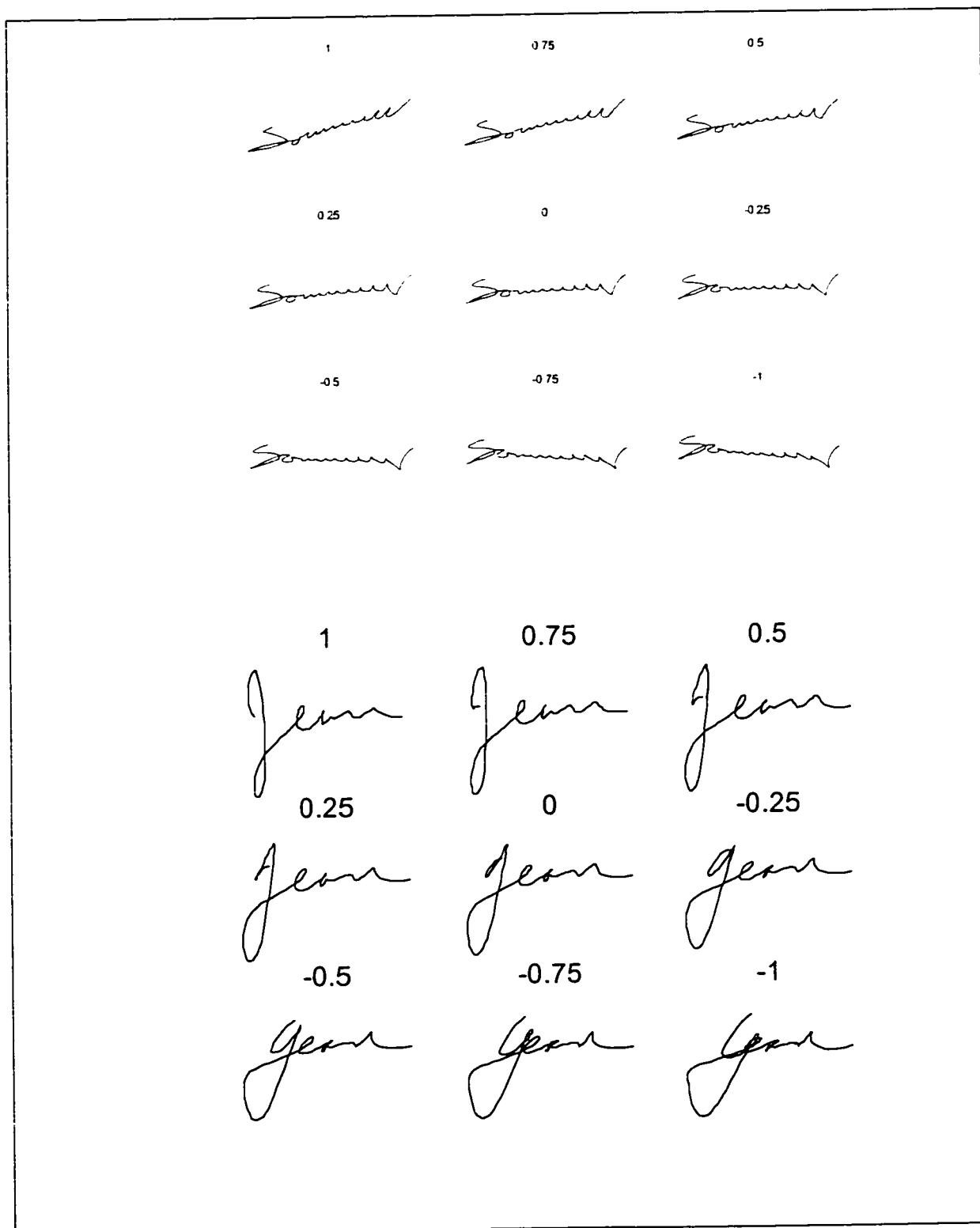


Figure 3-11 : Signatures hybrides de deux scripteurs

Avec l'hypothèse 4, on peut évaluer la moyenne de $P(X|\omega_2)$. On définit I_j comme étant l'imitateur miroir de I_j .

$$\begin{aligned} E[p(X|\omega_2)] &= E\left[\lim_{N \rightarrow \infty} \frac{1}{2N} \sum_{j=1}^N P(X|I_j) + P(X|I_{-j})\right] \\ &= \lim_{N \rightarrow \infty} \left[\frac{1}{2N} \sum M_j + M - (M_j - M)\right] \\ &= M \end{aligned}$$

La moyenne de $P(X|\omega_2)$ est égale, comme on s'y attendait, à la moyenne de $P(X|\omega_1)$. Bien que l'on ignore pour l'instant comment évaluer l'écart-type de la classe ω_2 , on sait que la distribution de cette classe prend la forme suivante :

$$P(X|\omega_2) = \prod_{i=1}^n \frac{1}{\sigma' \sqrt{2\pi}} e^{-\frac{(x_i - \mu_i)^2}{2\sigma_i'^2}} \quad (3-9)$$

De même façon que précédemment, on s'est servi de cette distribution pour générer des imitations de façon aléatoire. Pour évaluer l'écart-type de la classe ω_2 , on a considéré considérant que cette classe est formée des 3 imitateurs à disposition (dont on peut calculer l'écart local moyen à la classe authentique), ainsi que de leurs 3 imitateurs « miroirs ».

Sur la figure suivante, on voit un exemplaire de ces signatures générées.

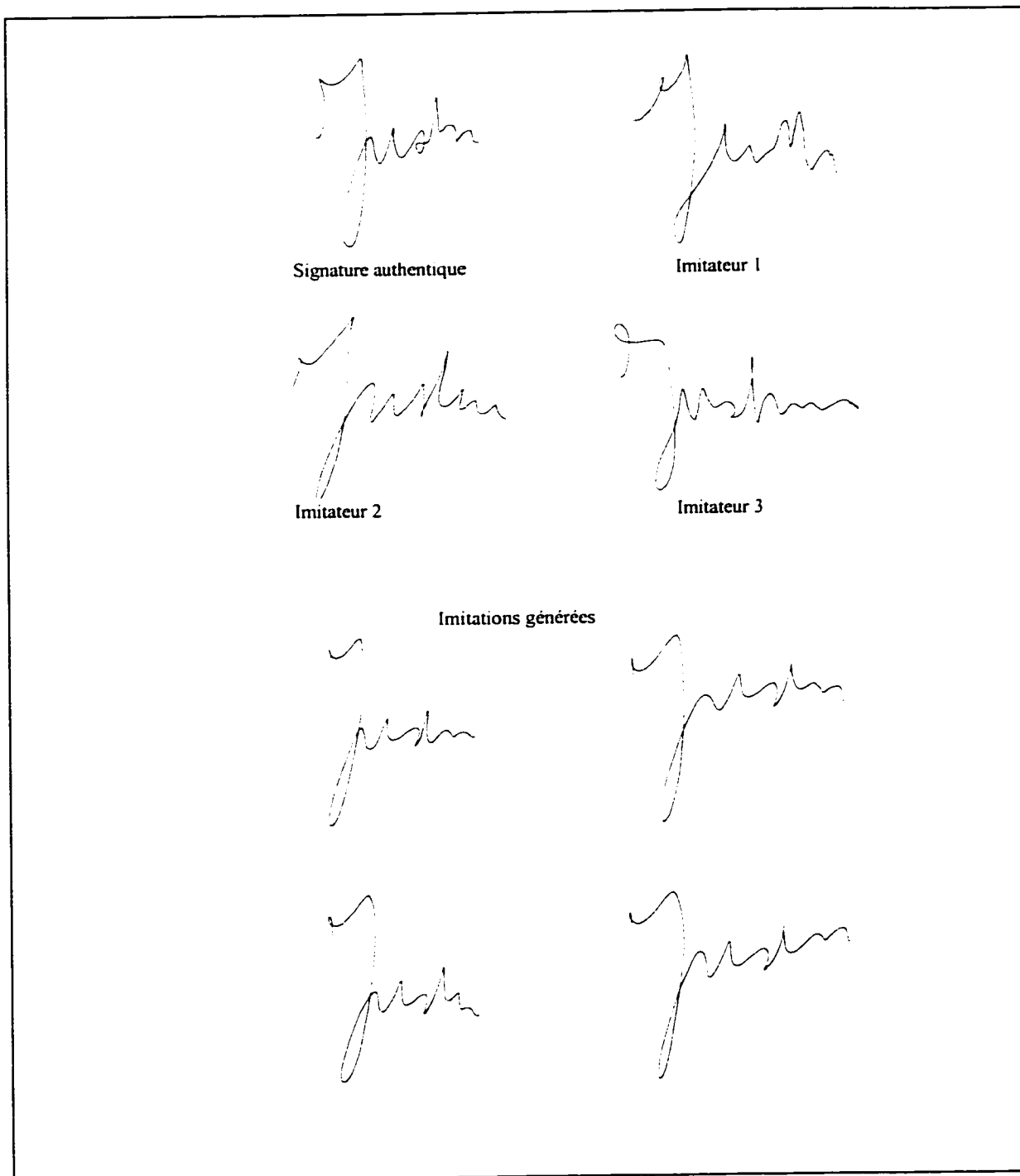


Figure 3-12 : Imitations « originales » et imitations générées

3.3 Résultats théoriques tirés du modèle

3.3.1 Estimation du taux d'erreur de reconnaissance d'un scripteur par le système

On a vu à la section 3.2.3 comment estimer le seuil en fonction de la théorie de la décision bayésienne. Ce seuil détermine la région où la signature est classée authentique et la région où elle est classée non-authentique. On peut estimer le taux d'erreur de classification pour chacune des régions, le taux d'erreur total étant la somme des taux d'erreurs de type I et II.

Pour simplifier, on pose que les probabilités à priori des classes ω_1 et ω_2 sont égales. Soit $X=[x_1, \dots, x_n]^T$ le vecteur de caractéristiques d'une signature. Sans perte de généralité, si X est de classe ω_1 , x_i suit une normale $N(0,1)$, et si X est de classe ω_2 , x_i suit une normale $N(0,\alpha)$. En effet :

$$x_i \propto N(\mu_i, \sigma_i) \Rightarrow \frac{x_i - \mu_i}{\sigma_i} \propto N(0,1)$$

$$x_i \propto N(\mu_i, \sigma_i') \Rightarrow \frac{x_i - \mu_i}{\sigma_i} \propto N(0, \frac{\sigma_i'}{\sigma_i})$$

On peut calculer la surface de seuil de décision :

$$\begin{aligned}
\frac{P(\omega_1 | X)}{P(\omega_2 | X)} &= 1 \\
\Leftrightarrow \frac{P(X | \omega_1)P(\omega_1)}{P(X | \omega_2)P(\omega_2)} &= 1 \\
\Leftrightarrow \frac{P(X | \omega_1)}{P(X | \omega_2)} &= 1 \\
\Leftrightarrow \log(P(X | \omega_1)) &= \log(P(X | \omega_2)) \\
\Leftrightarrow \sum_{i=1}^n \frac{x_i^2}{2} - \frac{x_i^2}{2\alpha_i^2} &= \sum_{i=1}^n \log \alpha_i \\
\Leftrightarrow \sum_{i=1}^n \left(\frac{\alpha_i^2 - 1}{\alpha_i^2} \right) x_i^2 &= 2 \sum_{i=1}^n \log \alpha_i
\end{aligned}$$

Ceci détermine l'équation d'une ellipse dans un espace à n dimensions. Si on suppose que chacune des caractéristiques locales est imitée avec la même dextérité par la population des imitateurs, on a, $\forall i, \alpha_i = \alpha$. D'où :

$$\sum_{i=1}^n x_i^2 = 2n \frac{\alpha^2 \log \alpha}{\alpha^2 - 1} \quad (3-10)$$

R_1 est la région dans laquelle la décision est de classer la signature comme authentique ($d(X) = \omega_1$) est l'hypersphère de rayon $c = \alpha \sqrt{\frac{2n \log \alpha}{\alpha^2 - 1}}$. R_2 est la région contenue à l'extérieur de cette hypersphère de décision $d(X) = \omega_2$ (classer la signature non-authentique).

L'espérance du taux d'erreur est :

$$\begin{aligned}
\text{erreur} &= \int_{R_1}^{R_2} P(\omega_2 | X) dX + \int_{R_2}^{R_1} P(\omega_1 | X) dX \\
&= \frac{1}{2} \int_{R_1}^{R_2} P(X | \omega_2) dX + \frac{1}{2} \int_{R_2}^{R_1} P(X | \omega_1) dX \\
&= \frac{\varepsilon + \beta}{2}
\end{aligned} \tag{3-11}$$

où ε et β sont la probabilité d'accepter un faux et de rejeter un vrai.

Calcul de β :

Soit X de classe ω_1 , alors $x_i \sim N(0,1)$. Par conséquent, x_i^2 suit une distribution du Khi-Deux à un degré de liberté ($x_i^2 \sim \chi^2_1$), et $\sum x_i^2$ suit une loi du Khi-Deux à n degré de liberté ($\sum x_i^2 \sim \chi^2_n$). La probabilité de rejeter X (ou de classer X dans la classe ω_2) est égale à la probabilité que $\sum x_i^2 > c^2$. Ceci correspond à $P(\chi^2_n > c^2)$.

Pour n grand, on peut approximer la loi du Khi-Deux à n degré de liberté, de moyenne et de variance $(n, 2n)$, à une distribution normale $N(n, 2n)$. D'où :

$$\begin{aligned}
\beta &\approx P(N(n, 2n) > c^2) \\
&= P(N(0,1) > \frac{c^2 - n}{\sqrt{2n}}) \\
&= P(N(0,1) > \sqrt{n} \left(\frac{2\alpha^2 \log \alpha}{\alpha^2 - 1} - 1 \right)) \\
&= 1 - \Phi \left(\sqrt{n} \left(\frac{2\alpha^2 \log \alpha}{\alpha^2 - 1} - 1 \right) \right) \\
&= 1 - \Phi(\sqrt{n} f(\alpha))
\end{aligned} \tag{3-12}$$

Avec $f(\alpha) = \frac{2\alpha^2 \log \alpha}{\alpha^2 - 1} - 1$

On montre [Moran 69] que la fonction $1 - \Phi(x)$ est bornée par :

$$(2\pi)^{-1/2} e^{-x^2/2} \left(\frac{1}{x} - \frac{1}{x^3} \right) < 1 - \Phi(x) < (2\pi)^{-1/2} e^{-x^2/2} \frac{1}{x}$$

Pour x grand, $1 - \Phi(x)$ tend asymptotiquement vers :

$$1 - \Phi(x) \approx \frac{(2\pi)^{-1/2} e^{-x^2/2}}{x}$$

D'où on a :

$$\beta = \frac{e^{-nf^2(\alpha)/2}}{(2\pi n)^{1/2} f(\alpha)} \quad (3-13)$$

Calcul de ε :

Soit X de classe ω_2 , alors $x_i \sim N(0, \alpha^2)$. Par conséquent, (x_i^2/α^2) suit une distribution du khi-deux à un degré de liberté, et $(\sum x_i^2/\alpha^2)$ suit une loi du khi-deux à n degré de liberté. La probabilité d'accepter X (ou de classer X dans la classe ω_1) correspond à la probabilité que $\sum x_i^2 < (c/\alpha)^2$. Ceci correspond à $P(\chi_n^2 < (c/\alpha)^2)$. En prenant encore une fois l'approximation de la Khi-Deux par la loi normale :

$$\begin{aligned}\varepsilon &= P(N(n, 2n) < (\frac{c}{\alpha})^2) \\ &= P(N(0, 1) < \sqrt{n}(\frac{2 \log \alpha}{\alpha^2 - 1} - 1)) \\ &= \Phi\left(\sqrt{n}\left(\frac{2 \log \alpha}{\alpha^2 - 1} - 1\right)\right) \\ &= 1 - \Phi\left(\sqrt{n}\left(1 - \frac{2 \log \alpha}{\alpha^2 - 1}\right)\right) \\ &= 1 - \Phi(\sqrt{n}g(\alpha))\end{aligned}$$

$$\text{Avec } g(\alpha) = 1 - \frac{2 \log \alpha}{\alpha^2 - 1}.$$

En utilisant la même approximation que précédemment, on a :

$$\varepsilon = \frac{e^{-ng^2(\alpha)/2}}{(2\pi n)^{1/2} g(\alpha)} \quad (3-14)$$

On en déduit une approximation du taux d'erreur :

$$\begin{aligned}\text{erreur} &= \frac{\varepsilon + \beta}{2} \\ &\approx \frac{1}{(2\pi n)^{1/2}} \left(\frac{e^{-ng^2(\alpha)/2}}{g(\alpha)} + \frac{e^{-nf^2(\alpha)/2}}{f(\alpha)} \right) \quad (3-15)\end{aligned}$$

Sur la figure suivante, on représente, pour différentes valeurs de α entre 1.2 et 2 (qui correspondent à différents écarts entre la classe des imitateurs et la classe authentique) et pour n compris entre 30 et 100 (pour n plus petit l'approximation d'une Khi-Deux par une normale n'est pas valable), le taux d'erreur théorique espéré. On observe que le taux d'erreur décroît avec l'augmentation du nombre de caractéristiques et l'augmentation du rapport des écarts-type.

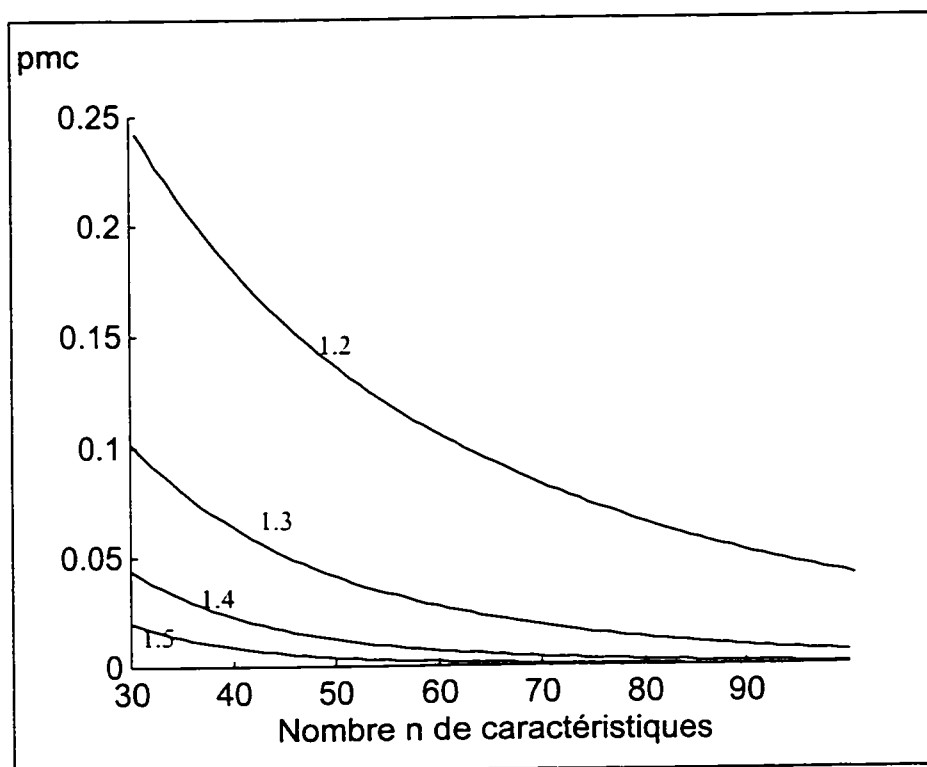


Figure 3-13 : Évolution de pmc en fonction de n pour différentes valeurs de α

3.3.2 Difficulté d'imitation d'un scripteur

Le problème de l'estimation d'un coefficient de difficulté d'imitation d'un scripteur a déjà été soulevé dans [Brault 88 , Brault 91]. Il serait très intéressant de prédire

les capacités de classification d'un système de VdS en fonction des paramètres intrinsèques à un scripteur. On pourrait notamment détecter les scripteurs problématiques, c'est-à-dire imitables facilement.

Selon la modélisation faite ici, les deux paramètres cruciaux sont la durée de la signature qui est environ égale au nombre n de dimensions du vecteur de caractéristiques, et le rapport des écarts-type locaux α des classes authentiques et non-authentiques. La formule établie à la sous-section précédente permet d'évaluer un taux d'erreur théorique en fonction de ces deux valeurs.

On objectera que le rapport des écarts-type α est précisément ce qui est difficile à évaluer, puisqu'en général, on ne dispose pas d'imitateurs. En effet, $\alpha = \frac{\bar{\sigma}'}{\bar{\sigma}}$.

Si on fait l'approximation que σ' est constant, soit que les imitateurs imitent « aussi bien » les différents types de tracés manuscrits, le rapport α dépend seulement de σ , la variabilité d'un scripteur, qui peut être évaluée sur un échantillon de référence.

Cependant, pour plausible qu'elle soit, n'est pas tout à fait vraie : une signature rapide avec un tracé complexe est plus difficile à imiter qu'une signature lente avec peu de changements de direction.

Le rapport α est donc une fonction de σ , mais aussi de différents paramètres du scripteur tel que la densité de changements de direction, la vitesse, etc. Étant donné les outils d'identification des éléments manuscrits à notre disposition, il serait intéressant d'étudier, sur une population d'imitateurs et de scripteurs assez grande, le lien entre α et différents paramètres d'un scripteur. Il serait peut être possible d'évaluer cette fonction

$\alpha(\sigma, a_1, \dots, a_m)$ où les $a_1, \dots, a_m$ sont m paramètres représentatifs de la difficulté à imiter un scripteur.

3.3.3 Qualité d'un système de VdS

On peut se servir de la modélisation pour comparer les systèmes de VdS qui extraient des caractéristiques locales et globales des signatures. Les caractéristiques globales sont en général des sommes ou des moyennes de caractéristiques locales sur l'ensemble de la signature. Soit X une caractéristique globale, et x_i la $i^{\text{ème}}$ caractéristique locale associée. Alors:

$$X = \sum_{i=1}^n x_i$$

ou

$$X = \frac{1}{n} \sum_{i=1}^n x_i$$

On peut donner quelques exemples de caractéristiques globales :

- Vitesse moyenne : $\frac{1}{n} \sum_{i=1}^n \sqrt{v_x(i)^2 + v_y(i)^2} = \frac{1}{n} \sum_{i=1}^n |v(i)|$
- Centre de gravité horizontal: $\frac{1}{n} \sum_{i=1}^n p_x(i)$
- Durée totale : $\frac{1}{n} \sum_{i=1}^n t_i$ (où t_i est la durée du $i^{\text{ème}}$ élément manuscrit,

estimée par les insertions et suppressions environnantes).

On montre simplement que si X est une caractéristique globale et que x_i est la $i^{\text{ème}}$ caractéristique locale d'un scripteur, alors X suit une loi normale :

$$\begin{aligned}
x_i &\propto N(\mu_i, \sigma_i^2) \\
\Rightarrow \sum_{i=1}^n x_i &\propto N(\sum \mu_i, \sum \sigma_i^2) \\
\Leftrightarrow \frac{1}{n} \sum_{i=1}^n x_i &\propto N(\bar{\mu}, \bar{\sigma}^2) \\
(\bar{\mu} = \frac{\sum \mu_i}{n}, \bar{\sigma}^2 = \frac{\sum \sigma_i^2}{n}) &
\end{aligned}
\tag{3-169}$$

La classe non-authentique, dont les caractéristiques locales sont de même moyenne que celle de la classe authentique, auront des caractéristiques globales de même moyenne que la classe authentique. Cela est logique, dans la mesure où il n'y a pas de polarisation parmi les imitateurs. En faisant l'hypothèse que le rapport des écarts types locaux α_i est constant, on vérifie également que le rapport des écarts-type de caractéristiques globales est égal à celui des caractéristiques locales.

$$\begin{aligned}
\sum \mu_i' &= \sum \mu_i \\
\Rightarrow \bar{\mu}' &= \bar{\mu} \\
\sum \sigma_i'^2 &= \sum \alpha^2 \sigma_i^2 \\
\Rightarrow \bar{\sigma}'^2 &= \alpha^2 \bar{\sigma}^2 \\
\Rightarrow \frac{\bar{\sigma}'}{\bar{\sigma}} &= \alpha
\end{aligned}$$

On arrive à la conclusion suivante: **une caractéristique globale n'est pas plus discriminante qu'une caractéristique locale.**

Le nombre de caractéristiques locales indépendantes d'une signature est beaucoup plus grand que le nombre de caractéristiques globales. Nelson [Nelson 94] extrait 20 caractéristiques locales d'une signature, mais toutes ne sont pas indépendantes. Pour une

signature de durée n , on a approximativement $n/4$ caractéristiques locales indépendantes (voir section précédente). Soit n et m le nombre de caractéristiques locales et globales et supposons un même rapport des écarts-type α . On peut calculer les rapport C_1 et C_2 des erreurs de type I et II . Pour C_1 :

$$\begin{aligned}
 C_1 &= \frac{\alpha_{GLOBAL}}{\alpha_{LOCAL}} \\
 &= \frac{\frac{\exp(-mg^2(\alpha)/2)}{(2\pi m)^{1/2} g(\alpha)}}{\frac{\exp(-(n/4)g^2(\alpha)/2)}{(2\pi(n/4))^{1/2} g(\alpha)}} \\
 &= \sqrt{\frac{n}{4m}} \exp\left(\frac{g^2(\alpha)}{2} \left(\frac{n}{4} - m\right)\right)
 \end{aligned} \tag{3-17}$$

Comme $n/4$ est à peu près toujours supérieur à m , le rapport C_1 des taux d'erreur est supérieur à 1, ce qui démontre bien qu'une caractérisation globale est moins efficace qu'une caractérisation locale.

Cas d'une caractérisation par segments

On peut considérer un segment comme une petite signature dont on extrait les caractéristiques globales. Ainsi, si on divise une signature en m segments dont on extrait deux caractéristiques chacun, on a un vecteur de caractéristiques de taille $2m$. La taille d'un segment varie bien sur, mais on peut établir une moyenne de 15 à 20 points par segment, d'où $m \approx n/8$.

Comme on a environ $n/4$ caractéristiques de points indépendants, on peut évaluer le rapport des taux d'erreurs :

$$\begin{aligned}
C_1 &= \frac{\alpha_{SEGMENTS}}{\alpha_{LOCAL}} \\
&\approx \sqrt{\frac{n/4}{n/8}} \exp\left(g^2(\alpha)\left(\frac{n}{4} - \frac{n}{8}\right)\right) \\
&\approx \sqrt{2} \exp\left(\frac{ng^2(\alpha)}{8}\right)
\end{aligned} \tag{3-18}$$

Pour une signature de 2 secondes ($n=200$) et un rapport des écarts-type $\alpha=1.4$ (valeur trouvée pour une caractéristique discriminante comme la vitesse horizontale), on a $C_1 \approx 14.3$, donc 14.3 fois plus d'erreurs, ce qui est un résultat très pessimiste pour la segmentation !

3.4 Conclusion

La méthode d'identification des éléments manuscrits établie au chapitre 2 a rendu possible une caractérisation locale des signatures et, sur un échantillon de taille suffisante, l'évaluation de la moyenne et de l'écart-type des caractéristiques locales de la signature d'un scripteur. A partir de là, on a pu établir un modèle statistique de la VdS. Ce modèle donne des résultats intéressants sur le taux d'erreur théorique d'un système en fonction du type de système de VdS (représentations globales, ponctuelles ou par segment des signatures) et en fonction des caractéristiques d'un scripteur (durée et rapport des écarts-type).

On peut débattre des hypothèses sur lesquelles le modèle s'appuie. Seulement les questions soulevées grâce à ces hypothèses sont fondamentales. En effet :

- s'il existe une polarité (éventuellement locale) parmi les imitateurs et qu'on arrive à évaluer ou prédire cette polarité, on pourrait élaborer des mesures de distances plus discriminantes entre signature-test et classe authentiques ;
- si on arrive à prédire le rapport des écarts-type pour un certain scripteur, on pourrait avoir un coefficient de difficulté d'imitation d'un scripteur très précis ; on aurait également une évaluation fiable du seuil de décision.

Ces deux questions ne seront que soulevées dans ce mémoire. Des recherches ultérieures sur une population d'imitateurs et de scripteurs assez grande sont nécessaires pour répondre à ces questions.

4. CONCEPTION GÉNÉRALE DU SYSTÈME DE VDS ET MÉTHODOLOGIE D'ÉVALUATION

4.1 Introduction

On a vu au premier chapitre que le système de VdS effectue un ensemble d'opérations, et que ces opérations peuvent être classifiées en opérations d'acquisition, prétraitement, analyse, comparaison, décision et apprentissage. Au deuxième chapitre on a décrit en détail l'identification des éléments manuscrits, opération que l'on peut classer parmi les prétraitements puisqu'elle vise à préparer l'information contenue dans une signature en vue de son analyse. Ce prétraitement étant de nature particulière, puisqu'il est un véritable système de RF en lui-même, nous y avons consacré un chapitre entier.

La **section 2** décrit l'ensemble des opérations, dans l'ordre. A chacune des étapes, on peut faire varier différents paramètres dont certains se révèlent cruciaux pour les performances du système. C'est notamment le cas à l'étape d'analyse, dans laquelle sont choisies les caractéristiques locales qui décriront la signature (doit-on extraire l'information de vitesse et/ou d'accélération, quel poids affecter à chaque type d'information, etc.). Dans quelques cas, un choix est fait parmi plusieurs alternatives, choix qui est justifié par une discussion. Dans la majorité des cas, une expérimentation complète est nécessaire pour prendre une décision sur la pertinence d'une certaine opération. A chaque étape, les diverses alternatives sont présentées. Le lecteur doit se

référer au chapitre 5, s'il désire connaître les résultats expérimentaux de chacune des alternatives.

Qui dit expérimentation dit méthodologie expérimentale, objet de la **section 3**. Il est nécessaire d'établir un schéma expérimental rigoureux, sinon, les résultats risquent d'être sans valeur. De quelle manière évaluer les performances d'un système ? A quel point peut-on faire confiance aux résultats ?

Nous croyons qu'en VdS, on fait parfois l'erreur d'accorder trop d'importance aux performances d'un système, dont on ne retient finalement qu'une chose: les taux d'erreurs de type I et II. Ces taux d'erreurs dépendent dans une large mesure de la qualité des données et des conditions expérimentales. Notamment, l'acquisition des imitations peut être faite en étant peu exigeant sur leur qualité, ce qui donne des taux d'erreurs artificiellement bas. Quelque soit la rigueur avec laquelle l'évaluation d'un système a été faite, les résultats dépendront toujours de la qualité des signatures servant à l'expérimentation.

Les taux d'erreurs en RF dépendent entièrement des données utilisées, ici de la qualité des signatures authentiques et des imitations. Dans ce but, les seuls types de conclusions que nous cherchons à établir sont de nature comparative: quel type d'élément manuscrit (point ou segment), quel type d'information (position, vitesse, accélération), quel axe (horizontal, vertical...) caractérisent au mieux les signatures (i.e. permettent la meilleure distinction des classes authentiques et imitées)? Ce sont ce type de questions auxquelles nous tenterons de répondre lors de cette phase d'expérimentation.

4.2 Conception du système de VdS

4.2.1 Plan de la section

Cette section suit l'ordre dicté par les opérations habituelles appliquées en RF, comme on peut le voir sur la figure suivante.

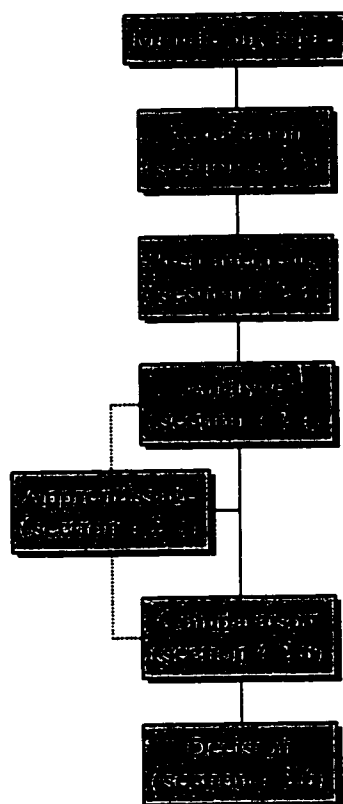


Figure 4-1 : Étapes du système de VdS et plan de la section

La **sous-section 2** décrit les spécifications techniques de la tablette graphique d'acquisition des signatures. La **sous-section 3** décrit les prétraitements appliqués aux signatures, soit le filtrage, la normalisation, la synchronisation et la segmentation. Comme la normalisation n'a pas encore fait l'objet d'une discussion, une plus longue

partie lui est réservée. La **sous-section 4** est consacrée à l'étude des différents types de caractéristiques que l'on peut extraire d'un élément manuscrit, utilisées pour représenter les signatures. La **sous-section 5** présente l'apprentissage des paramètres représentatifs de la classe authentique. La **sous-section 6** étudie les différents types de distances que l'on peut calculer pour des fins de comparaison entre le vecteur de caractéristiques d'une signature-test et celui de la classe authentique, ainsi qu'à l'étape cruciale de décision : à quelle classe appartient la signature-test ?

4.2.2 Acquisition des signatures

L'acquisition des signatures est faite sur un ordinateur portable de marque GRID permettant l'enregistrement de données manuscrites. Cet ordinateur est équipé d'un crayon électromagnétique émetteur et d'un écran fonctionnant également comme tablette graphique réceptrice. La résolution de la surface d'enregistrement est de 100 points/pouce en coordonnées x et de même en coordonnée y.

La fréquence d'échantillonnage f_{ech} est de 100 Hz, ce qui respecte largement le critère de Nyquist ($2*f_{max} < f_{ech}$), car l'énergie d'une signature est localisée à 98% entre 5 et 20 Hz [Brault 88].

Les signatures sont enregistrées par un logiciel de vérification de signature breveté, ce qui fait que les signatures sont pré-testées. La signature se fait dans un cadre de dimension 10*15 cm qui permet de guider le signataire au niveau des axes horizontaux et verticaux, et pour la taille de sa signature. Un fichier est créé pour chaque signature enregistrée, qui se présente sous la forme d'un tableau $\{x(nT), y(nT), p(nT)\}$, n étant

discret, et $T=0.01$ s. $p(nT)$ est un signal binaire exprimant l'état du contact entre le crayon et la tablette, et n'est utilisé que pour la représentation visuelle des données.

Notons que lorsque le crayon électromagnétique n'est pas en contact avec la tablette et ne s'en élève pas à plus d'approximativement un centimètre, le signal continue à être enregistré. On peut donc permettre aux signataires de soulever le crayon lumineux, si cette élévation n'est pas trop élevée et ne correspond pas à une hésitation importante de la main.

4.2.3 Prétraitements des données

Cette section décrit les 4 prétraitements effectués aux données : filtrage, normalisation, synchronisation et segmentation.

4.2.3.1 Filtrage des signatures

On applique un filtre passe-bas à moyenne mobile gaussien aux signatures pour éliminer les bruits résiduels situés dans les fréquences élevées. Ces bruits correspondent au frottement du stylo sur la table électronique où à certains tremblements de la main. Leur présence pourrait nuire à la synchronisation des signatures, mais surtout peut affecter la caractérisation des éléments manuscrits. Pour voir en quoi consiste le filtre gaussien, on suggère au lecteur de se référer au chapitre 2 où une description en est donnée. Nous filtrons les signatures à $\sigma=3$, afin de conserver 83.5% des fréquences à 20 Hz.

Après la construction d'une signature moyenne, celle-ci sera refiltrée à $\sigma=3$ pour éliminer certains bruits ou déformations qui apparaissent parfois aux régions instables du scripteur. En effet pour ces régions, il peut y avoir des erreurs dans l'identification des éléments manuscrits, ou encore un nombre très faible d'éléments manuscrits sont identifiés sur les signatures de référence ce qui entraîne parfois une évaluation bruitée des caractéristiques locales.

4.2.3.2 Normalisation des signatures

La normalisation des signatures a généralement pour but de réguler les variations de certaines caractéristiques globales causées par des conditions expérimentales inégales ou qui ne permettent pas au signataire de se repérer. Par exemple, les variations de taille sur différentes signatures, essentiellement dues à l'absence d'une délimitation claire de l'espace de signature, sont considérées non-significatives.

On distingue au moins trois types de normalisation des signatures, la normalisation selon l'angle, selon la durée et selon la dimension:

- **Normalisation selon l'angle:** On cherche parfois à « annuler » l'effet d'inclinaison d'une signature, cette inclinaison pouvant être due à une mauvaise position d'un scripteur plutôt qu'à qu'aux caractéristiques intrinsèques de sa signature. Comme le scripteur signait dans un cadre lui donnant comme repère les axes horizontaux et verticaux, nous ne jugeons pas nécessaire de normaliser selon l'angle moyen que fait la signature avec l'horizontale, d'autant plus que certains scripteurs ont des orientations angulaires spécifiques à leur signature (par ex. le scripteur no.7).

- **Normalisation selon la durée:** On a vu que les variations de durée étaient dues à des distorsions temporelles locales, et qu'une dilatation pouvait très bien suivre une compression. De plus dans l'étape de synchronisation on s'est chargé d'éliminer l'effet des distorsions. Nous considérons donc complètement inadéquat d'effectuer une normalisation selon la durée.
- **Normalisation selon la dimension:** Ce type de normalisation est le seul que nous avons eu à effectuer, aussi allons nous y réserver la discussion qui suit.

A la base, nous ne pensions pas avoir à normaliser selon la dimension des signatures: en effet le cadre imposé aux scripteurs leur permettait de respecter facilement la taille habituelle de leur signature. Cependant, de la manière dont était conçue l'expérience, les imitateurs n'avaient pas accès à l'information sur la dimension des signatures, ce qui risquait de biaiser les résultats à la hausse favorablement. Nous avons donc normalisé les signatures selon la taille, afin de ne pas nuire aux imitateurs.

Il y a deux façons de normaliser selon la taille. On peut le faire selon la longueur L du tracé cursif de la signature contenant N points

$$L = \sum_{i=2}^N \sqrt{(x(i) - x(i-1))^2 + (y(i) - y(i-1))^2}$$

On peut le faire également selon la surface de l'enveloppe de la signature:

$$S = (x_{\max} - x_{\min}) \times (y_{\max} - y_{\min})$$

Soit F le facteur de normalisation. Que F soit égal à S ou L on normalise les coordonnées par une division par F . Nous avons choisi le facteur de normalisation L car celui-ci nous semble plus représentatif de la taille de la signature que des valeurs

extrêmes (dans le cas de S) qui peuvent être sujettes à d'assez grandes variations, notamment pour des signatures comportant des tracés cursifs très amples, et d'amplitude variable.

Le problème de la normalisation est complexe, et il est difficile de prouver qu'une méthode est supérieure à une autre, ou même que la normalisation est nécessaire. Tout comme pour la segmentation, la normalisation nécessite la connaissance de la structure de la forme pour être appliquée adéquatement. Or la normalisation est un prétraitement « aveugle » .

Une méthode de normalisation « idéale » devrait être appliquée localement. Cette normalisation ne pourrait être faite qu'une fois l'identification des éléments manuscrits faite, ce qui nous ramène au cercle vicieux de la segmentation: une opération est censée faciliter la reconnaissance de la forme, mais pour être appliquée correctement, elle demande de déjà connaître la forme.

4.2.3.3 Synchronisation

Cette étape est d'abord effectuée lors de l'apprentissage par le système des éléments manuscrits descriptifs de la signature d'un scripteur, dont la méthode a été largement décrite au chapitre 2. Si on dispose de N signatures de références, on en synchronise $(N-1)$ par rapport à une signature choisie pour sa durée proche de la durée moyenne. Après la construction d'une signature moyenne, une resynchronisation est possible dans le but d'améliorer la description des classes d'éléments manuscrits.

Afin de vérifier si la resynchronisation est utile, une expérimentation sera faite pour comparer les performances des systèmes avec et sans resynchronisation.

Un signature-test est synchronisée sur la signature moyenne, permettant d'identifier ses éléments manuscrits.

4.2.3.4 Segmentation des signatures

La segmentation des signatures est facultative, et se fait après toutes les étapes précédentes du prétraitement. Une expérimentation est faite pour comparer les performances avec et sans segmentation. On calcule une seconde signature moyenne, c'est-à-dire calculée à partir des signatures synchronisées avec la première signature moyenne, et on trouve les minima de vitesse de cette signature après avoir filtré à $\sigma=3$, comme déterminé au chapitre 2. On a alors une segmentation de la signature moyenne, segmentation qui est alors « projetée » sur les signatures synchronisées.

4.2.4 Analyse

L'analyse est l'étape d'extraction des caractéristiques décrivant les éléments manuscrits d'une signature. On distingue deux types d'éléments manuscrits, les points et les segments. Le nombre de caractéristiques, ou la dimension du vecteur de caractéristiques, est fonction du nombre d'élément manuscrits. Les caractéristiques dépendent du type de coordonnées. On peut également chercher à représenter une signature par plusieurs caractéristiques locales à la fois.

Type d'information :

Une signature est initialement décrite par sa séquence de positions horizontales $\{p_x(1), \dots, p_x(n)\}$ et verticales $\{p_y(1), \dots, p_y(n)\}$. On ajuste la position initiale au point (0,0).

Par dérivation, on peut déduire la séquence de vitesse et d'accélération :

- **l'accélération** $\bar{a}(i)$ au point i est calculée à partir de la formule de différence finie du 2ème ordre suivante:

$$\bar{a}(i) = \frac{-\bar{p}(i+2) + 16\bar{p}(i+1) - 30\bar{p}(i) + 16\bar{p}(i-1) - \bar{p}(i-2)}{64}$$

- **la vitesse** $\bar{v}(i)$ au point i est calculée par la formule de différence finie vue au chapitre 2.

L'expérimentation nous dira quelle est l'information la plus représentative d'une signature.

Système de coordonnées :

La vitesse, l'accélération et la position peuvent être représentés selon différents systèmes de coordonnées. En 2-D, on utilise habituellement les coordonnées cartésiennes. Pour la vitesse, l'accélération et la position, on a $v_x, v_y, a_x, a_y, p_x, p_y$. On peut également utiliser les coordonnées polaires (r, α) : on a alors $v_r, v_\alpha, a_r, a_\alpha, p_r, p_\alpha$.

Lors de l'expérimentation, nous chercherons à identifier les types de coordonnées permettant la meilleure discrimination des classes de signatures.

Cas d'une représentation par segments :

On peut avoir pour les segments les mêmes caractéristiques que pour les points. La vitesse, l'accélération, la position peuvent être évaluées en moyenne sur l'ensemble des points composant un segment. Par exemple, la vitesse moyenne horizontale d'un segment est $\frac{1}{m} \sum v_x(i)$, où m est le nombre de points contenus dans le segment.

Visuellement, on décrit un segment par sa longueur, sa durée, sa taille, son orientation. Mais toutes ces quantités sont liées à la vitesse : taille = durée * vitesse moyenne, longueur = durée * vitesse absolue, orientation = angle α de la vitesse moyenne. Considérant que la vitesse est la caractéristique fondamentale d'une signature (et l'analyse des résultats nous le prouvera), les segments seront caractérisés seulement par leur vitesse moyenne.

On comparera alors les performances avec et sans segmentation des signatures.

Combinaison de caractéristiques :

Une signature S caractérisée localement par la vitesse selon l'axe x aura un vecteur de caractéristiques de cette forme : $V_x = [v_{x1}, \dots, v_{xn}]^T$. De même façon, S pourrait être caractérisée par V_y , A_x , etc... Serait-il pertinent de combiner ces caractéristiques de façon à avoir une meilleure représentation de la signature, permettant ainsi une meilleure discrimination entre classes ?

On peut penser que toute information supplémentaire devrait réduire l'incertitude sur la classe d'une signature-test, comme cela est d'ailleurs montré au chapitre 3. Mais on peut également se dire que l'ajout d'une information plus bruitée, de moins bonne qualité, pourrait au contraire rajouter de l'incertitude.

Les combinaisons étant quasi infinies, on se limitera en combinant la caractéristique locale la plus représentative (ou discriminante) d'une signature qui, comme on le verra, est V_x , avec une caractéristique un peu moins représentative, V_y et V_α . On obtient alors un vecteur de caractéristiques $[V_x, V_y] = [v_{x1}, \dots, v_{xm}, v_{y1}, \dots, v_{yn}]^T$.

Différents facteurs de pondérations pour V_y et V_α seront essayés : voir pour cela la section 4.2.6.

4.2.5 Apprentissage

L'apprentissage se situe à deux niveaux dans notre système de VdS. L'identification des éléments manuscrits (génération de la signature moyenne, segmentation,...) est une étape de prétraitements nécessitant un apprentissage dont on a déjà discuté.

Un deuxième apprentissage est nécessaire pour caractériser la classe authentique. Cet apprentissage consiste en l'estimation sur un échantillon de N signatures authentiques de la moyenne et de l'écart-type des caractéristiques locales.

Pour l'estimation des μ et σ locaux, on dispose d'un ensemble de N signatures d'apprentissage, synchronisées sur une signature moyenne de taille n . Pour chacun des n éléments manuscrits de la signature moyenne, on a identifié les éléments manuscrits correspondants sur les N signatures d'apprentissage. Sur certaines signatures, la méthode de synchronisation ne reconnaît pas d'élément manuscrit de même classe. Par conséquent, pour l'estimation des paramètres μ et σ d'une classe d'élément manuscrit, on dispose d'un échantillon de M' éléments manuscrits de même classe ($M' \leq M$).

Pour une classe d'éléments manuscrits, identifié sur M' signatures, on a les estimateurs $\bar{x}$ et s^2 non biaisés des paramètres μ et σ :

$$\bar{x} = \frac{\sum_{i=1}^{M'} x_i}{M'}$$

$$s^2 = \frac{\sum_{i=1}^{M'} (x_i - \bar{x})^2}{M' - 1}$$

Il est intéressant de remarquer que ces deux estimateurs sont les estimateurs de maximum de vraisemblance des paramètres d'une distribution normale (à l'exception que s^2 est alors divisé par M' au lieu de M), qui est la distribution de la classe authentique selon le modèle vu au chapitre précédent.

Il serait intéressant d'utiliser des estimateurs plus robustes. En effet les estimateurs précédents sont très sensibles à des données aberrantes. Des caractéristiques locales anormales, c'est-à-dire situées à plusieurs écarts-type de la moyenne et donc extrêmement

peu probables en théorie, arrivent relativement souvent en VdS, en raison de signatures localement « ratées » du scripteur et d'éléments manuscrits mal identifiés.

La statistique robuste est une branche de la statistique qui traite de ce genre de problèmes, où les estimateurs classiques ne sont pas adéquats. Les deux estimateurs de la moyenne et de l'écart-type les plus connus sont la médiane $\hat{x}$ et le « M.A.D » $\hat{\sigma}$ soit « median absolute deviation » :

$$\hat{\sigma} = \frac{\text{med}|x_i - \hat{x}|}{0,67}$$

(où « med » signifie médiane)

Une comparaison rapide des résultats obtenus avec les deux types d'estimateur n'a pas donné de différence significative sur les résultats obtenus.

4.2.6 Comparaison et décision

4.2.6.1 Comparaison par mesure d'une distance

L'étape de comparaison et décision est basée sur le principe de classification par la distance. D'après Belaïd [Belaïd 92] : « La classification par la distance constitue une des approches les plus simples et les plus intuitives de la RF. La motivation première de cette approche est qu'il est naturel de considérer qu'un élément appartienne à une classe s'il est plus proche de cette classe que toutes les autres. Il est donc nécessaire de définir la distance entre un point et une classe. ».

Choix d'une mesure de distance :

On définit une mesure de distance générale entre $X=[x_1, \dots, x_n]^T$ et $Y=[y_1, \dots, y_n]^T$

ainsi:

$$d_N(X, Y) = \left(\sum_{i=1}^n |x_i - y_i|^N \right)^{1/N} \quad (4-1)$$

En effet, pour différentes valeurs de N, on retrouve un certain nombre de distances classiques :

- Distance de Hamming (N=1) : $d_1(X, Y) = \sum_{i=1}^n |x_i - y_i|$
- Distance euclidienne (N=2) : $d_2(X, Y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2}$

On retrouve également, pour $N=\infty$, la distance du maximum. Parmi ces mesures de distance, nous allons conserver seulement la distance de Hamming. En effet, pour des valeurs de N supérieures, trop d'importance est accordée aux écarts éloignés de la moyenne, puisqu'ils sont élevés à une puissance. Or des écarts éloignés de la moyenne arrivent régulièrement sur les signatures authentique. Des essais avec la distance euclidienne nous ont montré qu'il était nuisible d'utiliser ce type de distance.

D'une manière plus générale encore, on peut affecter un paramètre de pondération à chaque dimension. En reprenant l'équation générale de mesure d'une distance, on a :

$$d_N(X, Y) = \left(\sum_{i=1}^n \frac{|x_i - y_i|^N}{\omega_i^N} \right)^{1/N}$$

Nous allons essayer les deux mesures de distances suivantes :

- La distance de Hamming normalisée par l'écart-type : $d_n(X, Y) = \sum_{i=1}^n \left| \frac{x_i - y_i}{\sigma_i} \right|$
- La distance de Hamming normalisée par la racine de l'écart-type :

$$d_{n2}(X, Y) = \sum_{i=1}^n \left| \frac{x_i - y_i}{\sqrt{\sigma_i}} \right|$$

Les essais faits avec la distance de Malhanobis notamment (distance euclidienne normalisée par l'écart-type) se sont révélés très décevants. En effet si l'écart-type d'un élément manuscrit x_i est 10 fois plus petit que celle d'un écart type x_j (ce qui est très courant), la pondération en élevant l'inverse cet écart-type au carré sera 100 fois plus importante : la moindre déviation légèrement importante sera donc très coûteuse sur x_i , alors qu'elle aura des répercussions négligeables sur x_j .

La distance de Hamming normalisée par la racine de l'écart-type a pour but de réduire l'influence des valeurs extrêmes (très petites) de σ . Cette distance de Hamming est plus proche de la distance de Hamming classique.

Donc les 3 mesures de distance d_1 , d_{n1} , d_{n2} seront testées, afin de déterminer laquelle est la plus appropriée au problème.

Distance d'un point à la classe ω_1 :

Soit une signature-test décrite par $X=[x_1, \dots, x_n]^T$, et la classe authentique ω_1 composée initialement d'un échantillon de N signatures de référence $Y_1, \dots, Y_N$ de même dimension. On a vu qu'on pouvait «résumer» l'information contenue dans ces N signatures de référence en $M=[\mu_1, \dots, \mu_n]^T$ et $\Sigma=[\sigma_1, \dots, \sigma_n]^T$. M contient les valeurs moyennes prises par les caractéristiques locales du scripteur, et Σ contient la variabilité locale de ces caractéristiques. Pour calculer la distance entre X et ω_1 , on distingue deux approches :

- l'approche « par prototypes » : La distance $d(X, \omega_1)$ est donnée par la distance entre X et la signature de référence Y_i la plus proche de X :

$$d(X, \omega_1) = \min_i d(X, Y_i)$$

- l'approche statistique : La distance $d(X, \omega_1)$ est donnée par la distance de la forme « moyenne » résumant w_1 , soit M :

$$d(X, \omega_1) = d(X, M)$$

Tout au long de ce mémoire, nous avons privilégié une approche statistique, nous avons donc choisi d'appliquer cette approche ici, d'autant plus que l'approche par prototype est une méthode relativement primitive. Il serait toutefois intéressant de comparer les résultats obtenus avec une approche « par prototype » lors d'une recherche ultérieure.

Cas des éléments manuscrits non-reconnus :

Nous avons considéré jusqu'à maintenant que le vecteur de caractéristiques X était « plein », i.e. que chaque classe d'éléments manuscrits avait été identifiée lors de la synchronisation (cela signifie que tous les points de la signature-test ont subi des substitutions). D'une manière générale ce n'est pas le cas, bien que le nombre moyen d'éléments manuscrits non-reconnus soit faible sur une signature-test : il varie de 3.05% à 7.80% pour le scripteur le plus instable. Il est intéressant de constater que le pourcentage moyen d'éléments manuscrits non-reconnus des imitations test n'est d'une manière générale pas plus élevé que pour un authentique test : ce n'est pas une information pertinente pour classer les signatures.

Comme la non-identification d'un élément manuscrit ne semble pas être une indication de non-authenticité d'une signature-test, nous avons décidé d'attribuer le coût moyen attribué aux éléments manuscrits de cette classe (coût évidemment relatif à la mesure de distance utilisée).

Combinaison de caractéristiques :

On peut représenter par plusieurs caractéristiques locales à la fois. Soit A et B deux types de caractéristiques locale (par exemple V_x et V_y). On peut représenter la signature par $X_A = [a_1, \dots, a_n]^T$ ou $X_B = [b_1, \dots, b_n]^T$. En combinant les caractéristiques A et B , la signature est alors représentée par $X = [a_1, \dots, a_n, b_1, \dots, b_n]^T$.

Supposons que la caractéristique A prise individuellement est plus discriminante que la caractéristique B. Il serait judicieux de pondérer davantage cette caractéristique lors de la mesure de distance. Soit β le paramètre de pondération, on évalue la mesure de distance de Hamming normalisée par l'écart-type:

$$\begin{aligned}
 d(X, \omega_1) &= \sum_{i=1}^n \left(\frac{|a_i - \mu_{Ai}|}{\sigma_{Ai}} + \beta \cdot \frac{|b_i - \mu_{Bi}|}{\sigma_{Bi}} \right) \\
 &= \sum_{i=1}^n \frac{|a_i - \mu_{Ai}|}{\sigma_{Ai}} + \beta \cdot \sum_{i=1}^n \frac{|b_i - \mu_{Bi}|}{\sigma_{Bi}} \\
 &= d(X_A, \omega_1) + \beta \cdot d(X_B, \omega_1)
 \end{aligned} \tag{4-2}$$

Cette distance correspond à la somme pondérée par β des distances obtenues avec les caractéristiques individuelles A et B

Lors de la phase d'expérimentation, on cherchera à combiner les caractéristiques locales les plus performantes pour différentes valeurs de β , en espérant obtenir un système de VdS plus performant.

4.2.6.2 Prise de décision à partir de la mesure de distance à la classe ω_1

Le lecteur pressent peut-être le retour de l'éternel problème de la VdS : comme on ne connaît pas la classe ω_2 (non-authentique), on ne dispose que d'une mesure de distance entre la classe authentique ω_1 et la signature. Or la méthode généralement appliquée est d'attribuer à la forme de classe inconnue X la classe la plus proche :

$$X \in \omega_i \Leftrightarrow \omega_i = \min_{\omega_i} d(X, \omega_i) \tag{4-3}$$

Telle qu'elle, cette méthode n'est pas utilisable en VdS car $d(X, \omega_2) = ?$.

La RF statistique apporte une solution partielle à ce problème.

Détermination d'une distance et d'un seuil optimaux par la RF statistique :

On a vu que normalement en RF, on attribue à un vecteur X la classe la plus proche selon une certaine métrique. La RF statistique peut se ramener à une démarche similaire. En effet, selon le critère de décision de Bayes, pour un problème à deux classes de distribution normale, on attribue à X la classe ω_1 si :

$$\begin{aligned}
 & \frac{P(\omega_1|X)}{P(\omega_2|X)} > 1 \\
 \Leftrightarrow & \frac{P(X|\omega_1)P(\omega_1)}{P(X|\omega_2)P(\omega_2)} > 1 \\
 \Leftrightarrow & \log(P(X|\omega_1)) - \log(P(X|\omega_2)) > \log \frac{P(\omega_2)}{P(\omega_1)} \\
 \Leftrightarrow & -\sum \left(\frac{x_i - \mu_{1i}}{\sigma_{1i}} \right)^2 + \sum \log \sigma_{1i} + \sum \left(\frac{x_i - \mu_{2i}}{\sigma_{2i}} \right)^2 - \sum \log \sigma_{2i} > \log \frac{P(\omega_2)}{P(\omega_1)} \\
 \Leftrightarrow & \sum \left(\frac{x_i - \mu_{1i}}{\sigma_{1i}} \right)^2 < \sum \left(\frac{x_i - \mu_{2i}}{\sigma_{2i}} \right)^2 + \sum \log \frac{\sigma_{2i}}{\sigma_{1i}} - \log \frac{P(\omega_2)}{P(\omega_1)}
 \end{aligned}$$

(4-4)

Seulement en pratique on ne connaît pas les paramètres d'écart-type de la classe ω_2 , bien que la discussion du chapitre 3 ait souligné que des recherches pourraient être faites dans le but de trouver des méthodes pour les estimer. Mais en appliquant le principe de moyennes égales pour les deux classes, on a (avec $\mu_1 = \mu_{1i} = \mu_{2i}$) :

$$\begin{aligned}
& \sum \left(\frac{x_i - \mu_i}{\sigma_{1i}} \right)^2 - \sum \left(\frac{x_i - \mu_i}{\sigma_{2i}} \right)^2 < \sum \log \frac{\sigma_{2i}}{\sigma_{1i}} \\
& \Leftrightarrow \sum \frac{\sigma_{2i}^2 - \sigma_{1i}^2}{\sigma_{1i}^2 \sigma_{2i}^2} (x_i - \mu_i)^2 < \sum \log \frac{\sigma_{2i}}{\sigma_{1i}} \quad (4-5) \\
& \Leftrightarrow \sum \frac{\alpha_i^2 - 1}{\alpha_i^2} \left(\frac{x_i - \mu_i}{\sigma_{1i}} \right)^2 < \sum \log \alpha_i
\end{aligned}$$

On se ramène alors à une notion de seuil, la région où une signature est classée authentique étant la région comprise à l'intérieur d'une surface de seuil elliptique. La pondération optimale est une division par l'écart-type multipliée par un paramètre $\frac{\alpha^2 - 1}{\alpha^2}$ qui tient compte du rapport local des écarts-types.

Si on pose que ce rapport α_i est constant ($\alpha_i = \alpha$), soit que toute région d'une signature possède la même difficulté à être imitée par l'ensemble des imitateurs, on a :

$$\begin{aligned}
& \sum \left(\frac{x_i - \mu_i}{\sigma_i} \right)^2 < \frac{\alpha^2}{\alpha^2 - 1} \sum \log \alpha \\
& \Leftrightarrow \sum \left(\frac{x_i - \mu_i}{\sigma_i} \right)^2 < \frac{n \alpha^2 \log \alpha}{\alpha^2 - 1} \quad (4-6)
\end{aligned}$$

La pondération optimale est alors de diviser par l'écart-type. Si on pouvait connaître le rapport α , on pourrait déterminer le seuil.

Le développement ci-dessus a montré comment sont liés le seuil et la mesure de distance utilisée. Par contre, il semble difficile d'utiliser le résultat obtenu, même en acceptant les deux hypothèses de moyenne égale et d'écart-type constant, puisqu'on n'a pas développé de méthode pour évaluer ce rapport α en fonction de certaines

caractéristiques du scripteur. A partir des données récoltées (voir section suivante), les rapports moyen d'écart-type variaient pour 8 scripteurs différents entre 1.2 et 2 pour une caractérisation par la vitesse horizontale locale. La valeur 1.2 correspondait effectivement au scripteur « visuellement » le moins stable.

On espère que des recherches ultérieures permettront de trouver des méthodes pour prédire dans une certaine mesure, le rapport α , ou mieux, les rapports α_i .

La méthodologie expérimentale ayant été conçue à des fins de comparaison entre différents systèmes de VdS, nous avons décidé d'évaluer les taux d'erreurs pour un seuil optimal, afin de comparer les systèmes sur une base égale (note : on est conscient qu'un seuil bien choisi biaise favorablement les résultats, mais l'expérimentation est axée sur la comparaison). Ainsi le problème du seuil n'est pas résolu, mais du moins il ne bloque pas l'expérimentation. On discutera de ce point en prochaine section.

4.3 Méthodologie d'évaluation

4.3.1 Introduction

Dans l'expérimentation, on cherche parmi un ensemble de systèmes de VdS (ou « classificateurs »), celui qui fournit les meilleures performances. Il faut donc comparer les performances du système général de VdS avec les différentes paramétrisations possibles, au niveau des prétraitements, de l'analyse, etc.

On a souvent tendance à prendre le design d'une expérience à la légère. Ripley [Ripley 96] qui présente la RF selon un point de vue de statisticien, dit :

« We will often want to choose between different candidate classifiers, and it will be usual to check that the performance targets are likely to be met. This needs an experimental test on some unseen examples. Such experiments are often (usually ?) very poorly designed, and slanted towards a favorite method. ».

Et plus loin, il ajoute :

« Prechelt surveyed two leading neural networks journal for 1993 and half of 1994. He deemed an evaluation of an algorithm acceptable if it used two or more realistic or more real problems and compared at least one alternative algorithm. Only 18% passed -in his words 'sad but true'. »

L'avis de Ripley et Prechelt est pour le moins négatif ! Mais au moins nous voilà avertis sur cette question. Le mot d'ordre sera donc de suivre un schéma expérimental rigoureux, pour l'acquisition des signatures (4.3.2), l'évaluation de la performance d'un système (4.3.3) et la validité des résultats obtenus (4.3.4).

4.3.2 Acquisition des signatures

4.3.2.1 Principes respectés lors de l'acquisition des données

L'acquisition des signatures doit être faite de façon très minutieuse si on désire avoir un échantillon de qualité. Par exemple, si on teste les performances d'un système de VdS dans l'absolu (ce qui est fréquemment fait malgré l'avertissement de Prechelt), un échantillon d'imitations de mauvaise qualité sera évidemment avantageux.

Nous citons ici les 3 principes généraux que devait suivre la phase d'acquisition des données, et qui seront discutés dans les sections suivantes:

- l'échantillon de scripteurs doit être autant que possible représentatif de la population des scripteurs ;
- pour chaque scripteur, les échantillons de signature authentiques et imitées doivent être de taille statistiquement valable (typiquement de taille 30);
- les imitations doivent être d'une qualité jugée "suffisante". Une façon de faire est de ne retenir que les imitations ayant été acceptée par un système de VdS pouvant classifier la signature lors de son acquisition

Seulement, en VdS la collecte d'un échantillon représentatif de la population des scripteurs est pratiquement impossible. Dans l'absolu, un échantillon représentatif serait composé des « strates » existantes dans la population mondiale, soit les différentes écritures, âges, sexe, niveau d'éducation. Cela est évidemment complètement irréaliste, et pas nécessaire dans la mesure où l'utilisateur d'un système de VdS n'est pas le « terrien moyen », mais possède un profil particulier qui dépend du domaine d'application du système.

4.3.2.2 Établissement de classes de signatures

Afin de limiter le risque d'avoir un échantillon trop particulier, nous avons établi une méthode de classification simple d'un scripteur. En utilisons cette méthode nous aurons probablement une meilleure chance d'avoir un échantillon représentatifs de l'ensemble de la population des signataires.

Nous avons vu dans le chapitre précédent qu'il semble possible de prédire la difficulté d'une signature à être imitée en fonction de certains paramètres, tels que la durée ou encore la grandeur de l'écart-type des caractéristiques. Les scripteurs choisis ici ayant autant que possible des caractéristiques différentes, il sera intéressant d'observer selon quels paramètres la facilité d'une signature à être imitée varie.

Cette typologie des signatures est cependant limitée, comme d'ailleurs toute typologie: on ne peut rendre compte de l'infinie diversité des signatures par une classification, d'autant plus que la notre n'est basée que sur trois propriétés. Et il est difficile d'évaluer par des chiffres la difficulté d'une signature à être imitée, d'autant plus que cette difficulté est loin d'être un absolu : par exemple, un faussaire imite mieux en général les signatures présentant des similitudes avec la sienne. De plus, nous aurions bien aimé pouvoir classifier initialement les signatures selon la stabilité du scripteur, qui est certainement un des meilleurs indices de la difficulté d'imitation. Cela demande par contre une étape de synchronisation, qui rendrait la typologie plus complexe. Nous avons donc du nous contenter de caractéristiques globales facilement et rapidement calculables.

Afin d'avoir une typologie présentant un nombre de classes compatible avec la taille de l'expérience, les signatures étaient évaluées selon trois propriétés, qu'elles pouvaient ou non avoir, ce qui donne un total de $2^3=8$ types de signatures. Les trois propriétés ont été choisies selon deux critères: être aussi peu que possible corrélées entre elles, et avoir une influence sur la difficulté d'imitation.

30 scripteurs ont donné un exemplaire chacun de leur signature. Les caractéristiques choisies sont :

- la vitesse absolue moyenne ;
- la durée ;
- la densité moyenne d'extremums.

Ces caractéristiques ont une corrélation assez faible entre elles: 0.2 entre la durée et la densité ainsi qu'entre la vitesse et la densité, et -0.5 entre la vitesse et la durée. La mesure de corrélation entre deux variables X et Y utilisée est simplement :

$$\frac{\text{cov}(X, Y)}{\sqrt{\text{var}(X) \cdot \text{var}(Y)}}$$

Comme on l'a vu au chapitre précédent, la durée a une influence sur la difficulté d'imitation. Une densité temporelle d'extremums élevée correspond à de fréquents changements de direction, donc à une séquence de mouvements plus complexe à apprendre. Pour une vitesse moyenne élevée, le signataire doit maîtriser avec beaucoup de précision la séquence de mouvements.

La nomenclature des classes est établie ainsi en respectant la numérotation binaire, avec (+) pour au dessus de la moyenne, et (-) pour au dessous de la moyenne.

Tableau 4-1 : Typologie des scripteurs

cla	vit	du	de
0	-	-	-
1	-	-	+
2	-	+	-
3	-	+	+
4	+	-	-
5	+	-	+
6	+	+	-
7	+	+	+

On trouve un exemplaire de chacun des scripteurs choisis (et donc de chacune des classes) sur la figure ci-dessous.

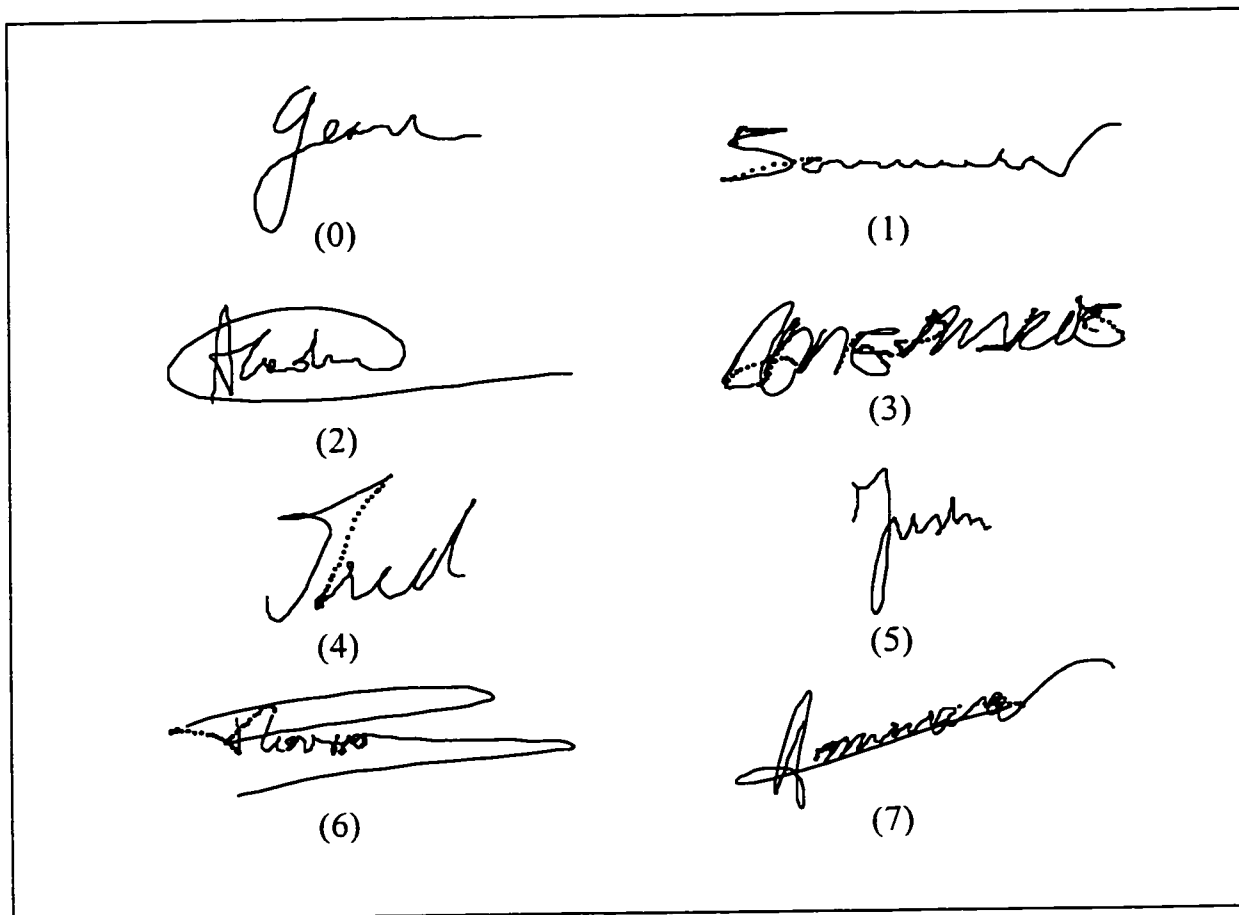


Figure 4-2 : Exemple de la signature de chacun des 8 scripteurs (qui sont identifiés par un numéro).

4.3.2.3 Acquisition des signatures authentiques

Les scripteurs ont été choisis de telle façon que chacune des 8 classes de notre typologie soit représentée. Chacun d'eux devait fournir un total de 30 signatures, afin d'avoir un échantillon statistiquement valable.

Voici les règles qui ont été respectées lors de l'acquisition des signatures:

- Si un scripteur était habitué à signer son nom et son prénom, il était contraint de choisir l'un des deux: en effet nous ne voulions pas rendre la tâche trop difficile aux imitateurs. Parvenir à automatiser un mouvement que l'on n'a jamais fait tout en respectant une forme est une tâche complexe. Plus la séquence est longue, plus il sera difficile pour l'imitateur de l'apprendre.
- L'acquisition des signatures devait être faite en une seule séance, mais des repos étaient accordés au scripteur, afin que sa signature ne soit pas déformée par la fatigue.
- Avant de commencer l'enregistrement, le scripteur se pratiquait sur papier puis sur la table électronique, jusqu'à ce que sa signature se soit régularisée: mouvement fluide, pas ou peu d'hésitations.
- Le scripteur avait la possibilité de rejeter une signature qu'il ne jugeait pas satisfaisante.

4.3.2.4 Acquisition des imitations

Nous considérons cette étape comme cruciale car il peut y avoir une différence de qualité immense dans les faux, selon la motivation et l'habileté de l'imitateur. Les conditions expérimentales habituelles, soit des prix en argent, la possibilité de se pratiquer sur papier, la possibilité de rejeter des faux dont le volontaire n'est pas satisfait sont utiles, mais ne nous ont pas apparus comme étant suffisants pour nous assurer de la qualité des imitations.

Nous pensons qu'il faut tester un système dans ses limites, et pour cela il faut donner toutes les conditions favorables à l'imitateur. Il faut supposer que l'imitateur ait

toutes les données statistiques décrivant une signature en main, et qu'il a eu tout le temps de se pratiquer pour corriger petit à petit ses erreurs, par exemple à l'aide d'un logiciel de vérification de signatures qui lui dit si son imitation a été acceptée ou non. La signature de l'imitateur est donc stabilisée au moment où l'acquisition commence.

Il vaut mieux tester un système sur un petit nombre d'imitateurs doués et motivés que sur un grand nombre d'imitateurs quelconques. Trois imitateurs ont été choisis: un étudiant en art spécialisé en dessin, une designer graphique et une étudiante motivée et douée. Chaque imitateur devait fournir 10 imitations par signature, pour un total de 80.

Au départ, l'imitateur avait une copie manuscrite de quelques exemplaires de chacune des signatures, avec leur temps d'exécution. Puis on lui prêtait pendant quelques jours un ordinateur portable équipé d'un crayon. Il pouvait ainsi pratiquer à sa guise chacune des signatures en s'habituant à la surface particulière de l'écran. Sur cet ordinateur, il pouvait également observer à l'écran les différentes signatures authentiques et les siennes, de façon à la fois dynamique et statique, pour maîtriser les particularités d'une signature et éliminer les défauts de ses imitations. Il pouvait aussi se pratiquer tant qu'il voulait sur le logiciel de vérification automatique de signature qui lui disait si sa signature était acceptée, rejetée ou considérée comme "imprévisible". De toute façon, seules les signatures acceptées ou considérées "imprévisibles" par ce logiciel étaient enregistrées, à l'exception d'une signature particulièrement difficile (scripteur no 3). Voici un exemplaire de cette signature. On peut voir comment la séquence biomécanique de

mouvements est complexe : nombreux levers de crayon, écriture en lettres détachées. Chaque lettre est presque une signature en soi.

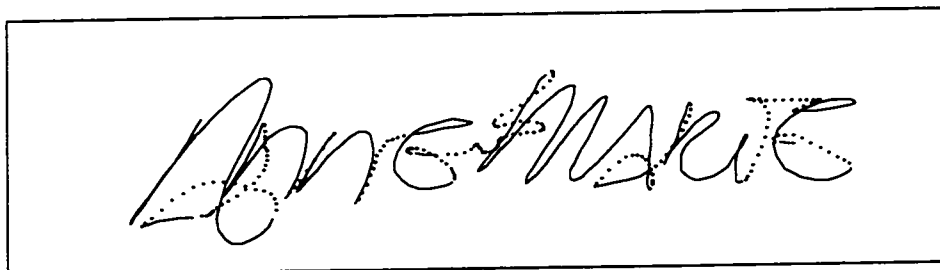


Figure 4-3: Exemple d'un scripteur difficile à imiter

En général, un imitateur pratiquait une signature à la fois, d'abord sur le papier puis sur l'écran. Ensuite, il lui fallait d'une dizaine jusqu'à une centaine d'essais avant que la première imitation soit acceptée. A partir de ce moment, cela lui prenait peu de temps (en général de dix à quinze minutes) pour fournir les dix imitations.

On dispose donc, pour chacune des 8 catégories de signatures, de 30 signatures authentiques et de 30 imitations faites par trois imitateurs, ce qui fait un total de $8 \times (30 + 30) = 480$ signatures.

4.3.2.5 Analyses humaines

Étant donné que nous analysons les résultats sur une base comparative, nous pensons qu'il serait intéressant de confronter les résultats obtenus expérimentalement aux performances obtenues par des analyses humaines.

Nous avons construit le test suivant: pour chacun des 8 scripteurs, nous avons mélangé 25 signatures authentiques et 25 imitations, pour un total de $8 \times (25 + 25) = 400$ signatures. Les personnes qui passaient l'examen (les analystes) avaient à leur disposition

un échantillon de 5 signatures authentiques (différentes de celles servant au test) leur servant de référence. On imposait aux analystes de donner leur avis (vrai ou faux) pour chacune des signatures.

4 personnes ont passé le test, ce qui demandait de 2 à 3 heures en moyenne. Les conditions du test leur étaient expliquées très clairement. Pour s'assurer de leur motivation, un montant de 100\$ était décerné à celui qui ferait le moins d'erreurs dans la classification des signatures. Les résultats du test seront donnés dans l'analyse des résultats.

4.3.3 Évaluation des performances d'un système

Détermination d'un seuil optimal :

Comme le problème du seuil n'a pas été résolu, le taux d'erreur d'un système est estimé pour un seuil optimal. On rappelle que deux erreurs sont possibles :

- une signature-test authentique dont la mesure de distance est supérieure au seuil est refusée. On note ce type d'erreur α , ou de type I.
- un signature-test non-authentique dont la mesure de distance est inférieure au seuil est acceptée. Ce type d'erreur est notée β , ou de type II.

On comprend que l'ajustement du seuil a une influence déterminante sur les deux types d'erreurs. Si le seuil est trop élevé, les signatures-test sont plus facilement acceptées. On obtient alors des taux d'erreurs α bas et β élevé (imitations classées authentiques). A l'inverse, pour un seuil trop petit, les signatures-test sont difficilement acceptées, avec pour conséquence des taux d'erreurs α élevé et β petit.

Selon le problème, on attribue des importances variables à α et β . Lorsqu'un système de VdS est utilisé en association avec une carte de crédit, le client accepte mal de voir sa signature refusée. Il vaut donc mieux que quelques faussaires passent à travers les mailles du filet, plutôt que risquer d'avoir des milliers de clients mécontents. Par conséquent, le seuil de décision est assez élevé. A l'inverse, lorsque le système de VdS est une barrière de sécurité à l'entrée d'un système informatique comportant des données de haute importance, il vaut mieux avoir un seuil peu élevé, afin de minimiser les risques d'intrusion.

Si on cherche à minimiser le risque de Bayes plutôt que l'erreur, en reprenant les équations de la page 123, avec des coûts C_1 et C_2 rattachés aux erreurs de type I et II :

$$\sum \left(\frac{x_i - \mu_{1i}}{\sigma_{1i}} \right)^2 < \sum \left(\frac{x_i - \mu_{2i}}{\sigma_{2i}} \right)^2 + \sum \log \frac{\sigma_{2i}}{\sigma_{1i}} - \log \left(\frac{P(\omega_2)}{P(\omega_1)} \cdot \frac{C_1}{C_2} \right)$$

Ainsi plus le rapport des coûts C_1/C_2 correspondant à l'acceptation d'un faux est grand, plus la distance entre le seuil de décision et le centre de la classe ω_1 diminue. Cela est logique, puisqu'on devient moins tolérant avec la signature-test.

Le seuil optimal dépend donc du problème traité et de l'importance relative accordée aux deux types d'erreurs. Comme ce travail n'a pas pour objet le développement d'une application particulière, nous avons choisi d'accorder une importance égale aux deux types d'erreurs, soit $C_1=C_2$. On veut alors minimiser $\frac{\alpha + \beta}{2}$.

Comme on ne connaît pas (ou mal) la classe ω_2 , l'évaluation d'un seuil est problématique. Sur la figure suivante, on peut voir l'influence de la valeur du seuil pour les taux d'erreurs α et β et le taux d'erreur moyen.

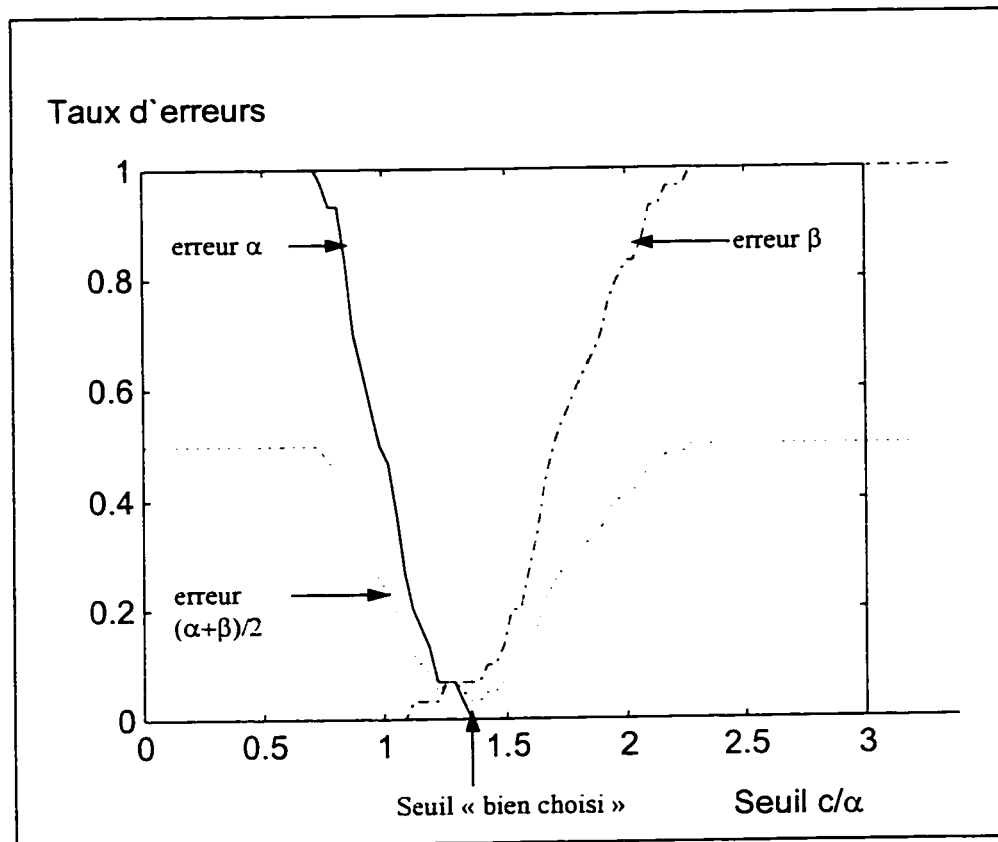


Figure 4-4 : Évolution des taux d'erreur en fonction du seuil

Comme on le voit sur la figure, le taux d'erreur d'un système est évalué au niveau du seuil « bien choisi », qui minimise $(\alpha + \beta)/2$. On remarque que le seuil est exprimé en unité d'écart-type.

Validation croisée :

Nous rappelons au lecteur que pour l'expérience, nous disposons pour chaque scripteur de 30 signatures authentiques et 30 imitations. On désire avoir une estimation du

taux d'erreur la plus précise qui soit, relativement à un certain système et un certain scripteur. Les signatures authentiques doivent servir à la fois d'ensemble d'entraînement (pour déterminer la forme qui représente la classe des signatures authentiques) et d'ensemble-test. Bien entendu, les ensembles d'entraînement et tests ne peuvent être confondus car cela avantagerait fortement le système, qui disposerait implicitement d'information préalable sur la forme à reconnaître.

La **validation croisée** s'intéresse au problème de la détermination des ensembles d'entraînements et de test. La méthode classiquement utilisée est de séparer l'échantillon total de taille N en ensemble d'entraînement de taille $N(V-1)/V$ et en ensemble-test de taille N/V . On évalue le taux d'erreur pour les V différents ensembles d'entraînement. L'estimation du taux d'erreur est la moyenne des taux d'erreurs de chacun des ensembles-test. Concrètement, si on a un échantillon de 30 signatures authentiques et que l'ensemble-test est de taille 5 ($V=6$) (donc l'ensemble d'apprentissage de taille $(6-1)/6 \cdot 30 = 25$), on doit répéter $V=6$ fois pour estimer le taux d'erreur total. En général, plus V est grand, plus l'ensemble d'entraînement est grand, donc plus l'estimation du taux d'erreur est précise (bien que cela ne soit pas tout à fait vrai, voir [Ripley 96, page 70]). Par contre le temps de calcul est linéaire $O(V)$.

Nous avons utilisé la technique du « leave-one-out », soit d'avoir un échantillon d'entraînement de taille $(N-1)$ et un échantillon test de taille 1. Mais le cas de la VdS est particulier, puisque les imitations ne peuvent servir à l'entraînement. La solution appliquée est donc :

- Lorsque les signatures-test sont les imitations, l'ensemble d'entraînement est composé des $N=30$ signatures authentiques.
- Lorsque la signature-test est la $i^{\text{ème}}$ signature authentique, l'ensemble d'entraînement est composé $(N-1)=29$ autres signatures authentiques

Ainsi on doit calculer 31 fois les statistiques locales au lieu d'une fois! Cependant cela nous permet d'utiliser au maximum les données dont nous disposons.

Performance globale du système :

Un système est évalué sur 8 scripteurs différents. Pour chacun des scripteurs, un taux d'erreur est obtenu. Le taux d'erreur d'un système qui servira de comparaison avec d'autres systèmes est la moyenne des taux d'erreurs obtenus sur les 8 scripteurs.

Cette approche peut paraître évidente. Toutefois il arrive qu'un système A soit plus performant globalement qu'un système B, mais produise de mauvais résultat sur un scripteur particulier. Par conséquent, l'approche choisie est simplificatrice dans la mesure où deux systèmes de VdS devraient idéalement être comparés relativement à un certain scripteur. On tiendra compte du fait que les conclusions dégagées ne sont pas nécessairement absolues.

4.3.4 Validité des résultats obtenus et comparaison de systèmes

Intervalle de confiance sur le taux d'erreur réel :

Pour évaluer un système, on fait l'estimation d'un taux d'erreur. Ce taux d'erreur estimé est la moyenne des taux d'erreurs pour les 8 échantillons de 60 signatures correspondant aux 8 scripteurs. Quelle confiance peut-on faire à ce taux d'erreur ?

La démarche la plus simple est de modéliser un classificateur par un « vrai » taux d'erreur que l'on appellera p . La probabilité que le classificateur fasse une erreur pour une forme-test est donc de p . Le nombre d'erreurs faites par le classificateur pour la classification d'un échantillon de taille n suit une loi binomiale $B(n,p)$, de moyenne np et de variance $np(1-p)$. Pour p petit, on peut approximer $(np(1-p))$ par (np) .

Pour n grand, on peut approximer la binomiale du nombre d'erreurs par la loi normale $N(np,np)$. Par conséquent, l'estimateur $\hat{p}$ du taux d'erreur suit une loi $N(p,p/n)$.

Un intervalle à 95% pour $\hat{p}$ est :

$$p - 1,96\sqrt{\frac{p}{n}} < \hat{p} < p + 1,96\sqrt{\frac{p}{n}}$$

Par exemple, si le taux d'erreur estimé $\hat{p}$ est de 5% sur 480 signatures, en estimant $\sqrt{\frac{p}{n}}$ par $\sqrt{\frac{\hat{p}}{n}}$, on a :

$$\begin{aligned} p - 0,02 &< \hat{p} < p + 0,02 \\ \Leftrightarrow \hat{p} - 0,02 &< p < \hat{p} + 0,02 \\ \Leftrightarrow 3\% &< p < 7\% \end{aligned}$$

Ainsi l'estimation du taux d'erreur est relativement peu précise pour 480 signatures.

Comparaison de deux systèmes (ou « classificateurs ») :

Si on désire comparer deux classificateurs de taux d'erreurs p_1 et p_2 , l'estimation de $(p_1 - p_2)$ suit une normale $N(p_1 - p_2, (p_1 + p_2)/n)$, d'après ce qu'on vient de voir. En appliquant cette méthode, il sera rare que l'on obtienne des différences de taux d'erreur

significatives entre deux systèmes. Seulement cette comparaison de deux classificateurs n'est pas adéquate parce que les deux classificateurs utilisent le même ensemble test. Des méthodes plus appropriées sont disponibles, tel que le test de McNemar [Fleiss 81]. Soit n_A les erreurs faites par la méthode A mais pas par la méthode B, et n_B les erreurs faites par B mais pas par A. Alors :

$$C_{AB} = \frac{|n_A - n_B| - 1}{\sqrt{n_A + n_B}} \propto N(0,1)$$

Il faut donc des grands ensembles-test pour distinguer des classificateurs à faible taux d'erreurs et fonctionnant de manière similaire, .i.e. qui font à peu près les mêmes erreurs. Par exemple, si sur 480 signatures un système de VdS A fait 5 erreurs et que le système B fait toutes les erreurs que A fait plus deux autres, on a $n_A = 5$, $n_B = 7$, et $C_{AB} = (1/2)^{1/2}$. La différence n'est pas significative (cela nous empêchera au chapitre suivant de conclure définitivement sur le système aux meilleures performances).

D'après Ripley, ce test suggère qu'il faut $n_A > 4$ pour avoir une différence significative (mais seulement si $n_B = 0$). Donc pour détecter une différence de 1% cela prend un échantillon de taille 500.

Comme notre échantillon est de taille 480, nous allons considérer **qu'une méthode A est significativement plus performante qu'une méthode B si le taux d'erreur de B est supérieur de 1% au taux d'erreur de A.**

Selon Prechelt (voir plus haut), un taux d'erreur doit être estimé par un algorithme de nature différente. Nos résultats seront donc comparés avec les résultats de l'algorithme

de calcul d'une distance par alignement élastique, décrit au chapitre 2 et que nous souhaiterions améliorer.

4.4 Conclusion

Les chapitres 2 et 3 ont présenté les idées fondamentales de la recherche. Ce chapitre a décrit d'un point de vue plus pratique et explicite le système de VdS élaboré, en soulignant les aspects de ce système qui ne peuvent être établis qu'après l'expérimentation, dont on vient de décrire la méthode en détail.

Il ne reste plus qu'à compiler des résultats et à en faire l'analyse. Ceci est l'objet du prochain chapitre.

5. ANALYSE DES RÉSULTATS

5.1 Introduction

Même si la structure générale du système de VdS a été fixée au chapitre précédent, il y a encore un nombre important de degrés de liberté à fixer. Par exemple lors de l'analyse, si on se restreint à ne représenter les signatures que par un type de caractéristiques locales, on peut extraire des caractéristiques de position, vitesse, ou accélération (3 choix), en choisissant des coordonnées cartésiennes (x,y) ou polaires (r,α) (4 choix). Lors de l'étape de comparaison, on a sélectionné (en se restreignant) 3 mesures de distance. Cela donne déjà $4 \times 3 \times 3 = 36$ possibilités à expérimenter, sur 480 signatures à chaque fois ! En fait, si on a n degrés de liberté dans notre système et que pour chaque degré de liberté, il y a m choix possibles, une étude exhaustive nécessitera n^m paramétrisations différentes. Il n'est donc pas question d'étudier toutes les possibilités.

On trouve les résultats de chacune des expériences en annexe. Mais chaque résultat ne sera pas présenté analysé individuellement: des regroupements seront faits, par exemple selon le type de coordonnées. On ne pourra pas toujours tirer des conclusions du type « le système A est meilleur que le système B », la discussion de la méthodologie d'évaluation ayant souligné que les résultats doivent être interprétés en respectant les bornes que nous impose la statistique.

On a déjà dit que nous préférons mettre l'emphase non pas sur les résultats eux-mêmes, mais sur les comparaisons entre les résultats obtenus de différentes façons. Aussi on commencera par la **section 2** réservée aux taux d'erreur dits "de référence", soit ceux trouvés lors de l'alignement élastique classique et ceux des analyses humaines. Cela nous donnera une échelle d'évaluation des résultats obtenus par notre approche « synchronisation/statistique locale ».

En **section 3**, on fait une comparaison des taux d'erreur obtenus pour les différents types de prétraitements, analyses et comparaisons possibles. On en tire des conclusions sur le choix des opérations à effectuer, sur les informations discriminantes.

La **section 4** sera réservée à une tentative d'analyse plus approfondie on y vérifiera notamment si certaines hypothèses de départ du modèle sont valides. On cherchera également à voir s'il est possible de prédire, partiellement du moins, la difficulté à laquelle un scripteur est imité : en ceci on poursuit l'étude commencée dans la thèse de Brault [Brault 88] sur la détermination d'un coefficient de difficulté d'imitation des signatures manuscrites.

5.2 Résultats de référence

5.2.1 Alignement élastique classique

Le chapitre 2 a présenté l'algorithme d'alignement élastique, qui calcule une distance entre deux séquences en déterminant la séquence de comparaisons optimale des éléments des deux séquences. Nous avons calculé une distance par cet algorithme entre la

signature de référence choisie pour l'identification des éléments manuscrits et chacune des autres signatures. Pour un seuil optimal (voir méthodologie expérimentale), on a obtenu les taux d'erreurs suivants pour chacun des scripteurs :

Tableau 5-1 : Taux d'erreur de la méthode d'alignement élastique

Signature	Taux (%)
0	0.00
1	6.67
2	0.00
3	10.00
4	0.00
5	0.00
6	0.00
7	0.00
TOTAL	2.08

L'erreur moyenne n'est que de 2.08%, soit 10 erreurs pour 480 signatures. De plus ces 10 erreurs sont réparties sur 2 scripteurs, le 1 et le 4. Ce résultat nous laisse une faible marge de manœuvre : en effet, on cherche à faire mieux que l'alignement élastique, donc à concevoir un système produisant moins de 10 erreurs sur les 480 signatures. Pour que la différence soit statistiquement significative, le système devra produire 1% d'erreurs en moins (voir méthodologie d'évaluation), soit au plus 6 erreurs.

L'alignement élastique ne produit aucune erreur sur 6 des 8 scripteurs. Les 2 scripteurs problématiques ont des signatures plus courtes, donc plus facilement imitables

d'un point de vue statistique comme démontré au chapitre 3. Un de ces deux scripteurs (« Fred », le no 3), est très instable, alors que l'autre (« Sommelet », le no 1) possède une séquence de mouvements faciles à exécuter : la partie centrale de la signature n'est constitué que d'oscillations de bas en haut. Cela nous donne une indication pour la détermination du rapport des écarts-type α'/α des classes non-authentiques et authentiques : un tracé essentiellement constitué d'oscillations verticales semble plus facile à exécuter. A l'inverse, on pourrait dire qu'un tracé constitué de nombreuses oscillations horizontales serait difficile à exécuter. Le lecteur peut d'ailleurs vérifier par lui-même qu'il est probablement moins précis pour des oscillations horizontales. Une recherche ultérieure serait nécessaire pour confirmer cette idée, par exemple en vérifiant que le rapport local des α_i'/α_i est plus élevé pour des portions de signatures horizontales.

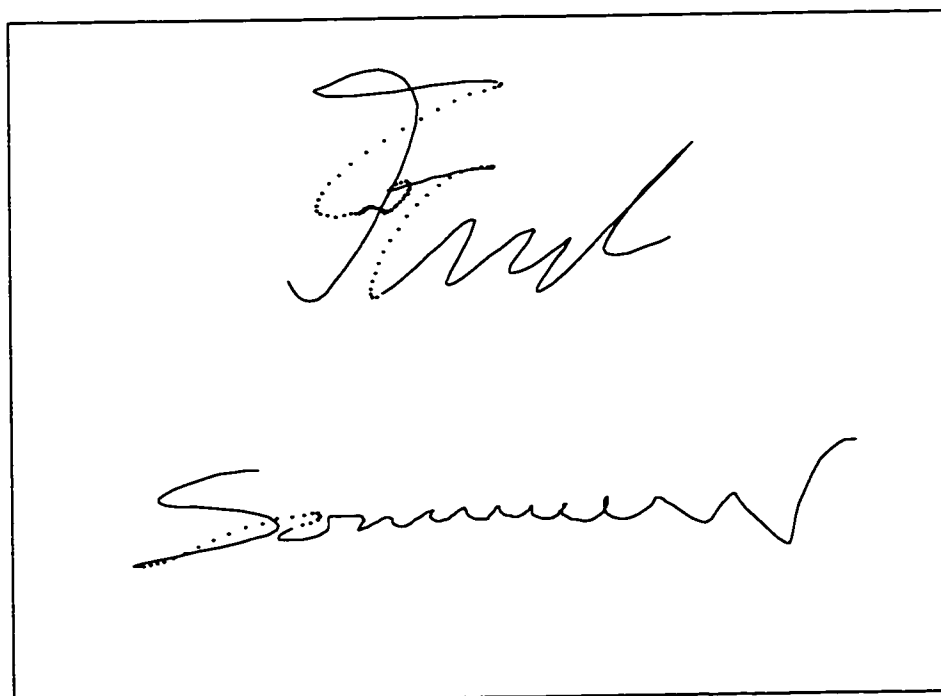


Figure 5-1 : Deux scripteurs problématiques

Malgré tous les efforts des imitateurs et les conditions sévères de sélection des imitations, il est très difficile de produire un faux de qualité suffisamment grande pour qu'il déjoue l'alignement élastique classique.

5.2.2 Résultats des analystes humains

Les résultats obtenus par des analystes humains sont rassurants pour les systèmes de VdS car les humains font beaucoup plus d'erreurs que l'alignement élastique. Il semble en effet très difficile pour un humain d'atteindre un taux d'erreur comparable à l'alignement élastique, du moins s'il n'a pas reçu une formation préalable en analyse de signatures.

Le test, comportant 400 signatures (50 par scripteur) à classer, a été soumis à 4 personnes, qui mettaient de 2 à 3 heures à le compléter. Les résultats vont de 10% à 32.5% d'erreurs. Des taux d'erreurs de reconnaissance ont avoisiné les 50% pour certaines combinaison analyste/scripteur, ce qui est pratiquement équivalent à faire des choix au hasard. L'auteur de ces lignes doit lui-même reconnaître la difficulté du test, car il a fait plusieurs erreurs en cherchant à reconnaître sa propre signature !

Ce test démontre qu'un humain n'ayant pas de formation en VdS ne peut rivaliser avec un système automatique de VdS. Évidemment, il serait intéressant de soumettre le même test à des experts en analyse de signatures, test ultime pour les systèmes de VdS.

5.3 Comparaison des résultats pour différents classificateurs

Des résultats ont été compilés pour différents types de systèmes correspondant à des combinaisons prétraitement/analyse/comparaison. On en fait l'analyse dans cette section.

La **sous-section 1** porte sur l'influence des différents prétraitements sur les taux d'erreurs : resynchronisation et segmentation. La **sous-section 2** compare les résultats selon les caractéristiques locales extraites d'une signature. La **sous-section 3** compare les résultats selon la mesure de distance choisie.

Un certain système de VdS qui extrait une seule caractéristique locale est généralement décrit ainsi : on note V, A ou P le type d'information extraite (vitesse, accélération ou position), on note x, y, r ou α le type de représentation de l'information (horizontal, vertical, valeur absolue, ou angle (orientation)), et on note d_H , d_{N1} et d_{N2} le type de mesure de distance utilisée (Hamming, Hamming normalisé par l'écart-type, Hamming normalisé par la racine de l'écart-type). Par exemple, le système $[V\alpha, d_H]$ correspond à un système qui extrait l'angle de la vitesse locale et mesure une distance de Hamming entre la signature-test et la classe authentique.

On rappelle que les taux d'erreurs sont évalués pour un seuil « optimal ». On trouve tous les résultats en annexe.

5.3.1 Influence des prétraitements

5.3.1.1 Effet d'une resynchronisation

Dans cette section on compare les résultats obtenus avec et sans resynchronisation par la signature moyenne (voir chapitre 2). La première synchronisation dépend d'une signature choisie en partie arbitrairement, malgré les efforts pour choisir une signature aussi représentative que possible. Si cette signature présente de plus des parties de taille significative (les « inconstances ») qui ne sont pas présentes sur d'autres signatures du même scripteur, la synchronisation risque d'être entachée d'un certain bruit. La signature moyenne issue de cette première synchronisation, si elle est de même taille que la signature choisie initialement, a perdu beaucoup des particularités de la signature initiale. On espère que cette signature moyenne, plus représentative de la classe authentique, devrait synchroniser de façon plus précise l'ensemble des signatures. La conséquence indirecte serait qu'avec une meilleure synchronisation, la mesure de distance d'une signature-test serait plus précise, favorisant ainsi les signatures-test authentiques.

Dans le tableau suivant, on compare les résultats obtenus avec et sans resynchronisation pour des systèmes ayant extrait des caractéristiques locales de vitesse en coordonnées cartésiennes, et pour les mesures de distance d_H et d_{N1} .

Tableau 5-2 : Comparaison des taux d'erreurs avec et sans resynchronisation

Système	Taux d'erreur sans resynchronisation	Taux d'erreur avec resynchronisation	Différence
$[Vx, d_H]$	2.71	1.25	-1.46
$[Vx, d_{N1}]$	2.71	1.04	-1.67
$[Vy, d_H]$	5.21	4.58	-0.63
$[Vx, d_{N1}]$	4.38	4.38	0
Moyenne	3.75	2.81	-0.96

Dans tous cas sauf un, les taux d'erreurs sont moins élevés lorsqu'on resynchronise, toute chose égale par ailleurs. L'alignement initial est donc bel et bien entaché de bruit, et ce bruit nuit à la reconnaissance. Pour la suite de l'expérimentation, l'opération de resynchronisation sera automatiquement appliquée.

5.3.1.2 Segments et points

En VdS, on est parfois porté à croire que le segment (portion de signature séparée par des maxima de courbure ou minima de vitesse) possède des caractéristiques très significatives de la signature d'une personne. Nous avons montré au chapitre 3 que selon la modélisation statistique, les segments sont en quelque sorte de petites signatures que

l'on analyse globalement. Cela a pour conséquence une perte d'information utile pour la discrimination.

Grâce à notre méthode de segmentation très fiable, on peut comparer de façon juste les deux types d'éléments manuscrits, soit les segments et les points. En effet, il eut été injuste d'utiliser des segments mal identifiés par une méthode de pré-segmentation, et d'en tirer une conclusion à leur désavantage.

Nous avons segmenté les signatures pour la valeur $\sigma=3$ par la méthode de post-segmentation. On extrait les caractéristiques locales ($\Delta x, \Delta y$) (voir figure).

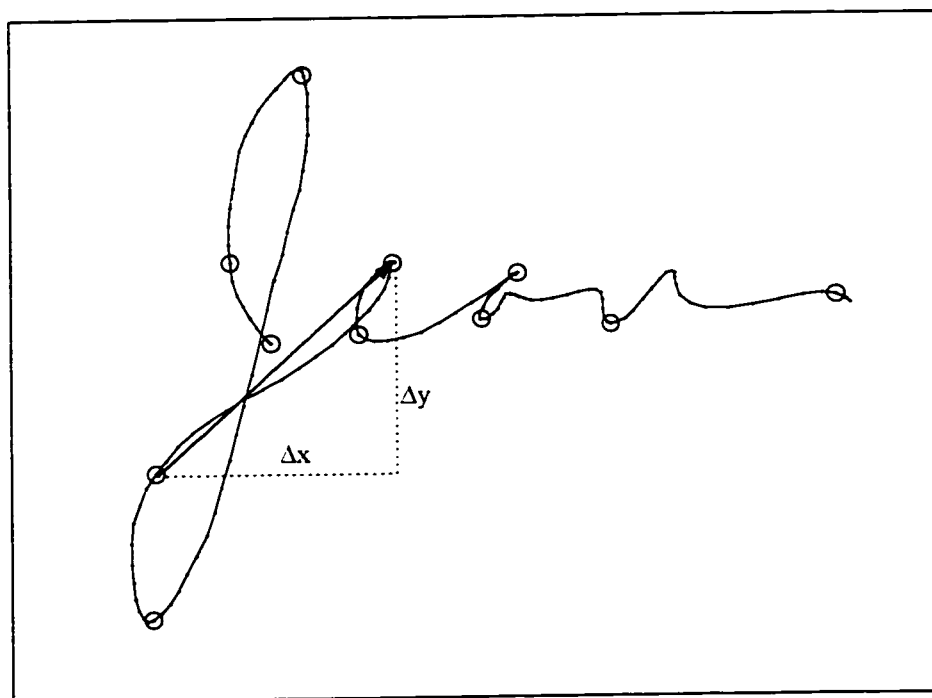


Figure 5-2 : Caractéristiques extraites des segments

Nous avons choisi 3 modes de représentation pour les signatures : Δx , Δy , $[\Delta x, \Delta y]$, couplés aux mesures de distance de Hamming et de Hamming normalisée. On trouve

les résultats en annexe. Le tableau suivant présente les taux d'erreur moyens exprimés en %.

Tableau 5-3 : Taux d'erreur obtenue par la représentation par segments

Δx	4.38	11.67
Δy	5.62	13.12
$[\Delta x, \Delta y]$	3.12	9.58

Le meilleur taux d'erreur, obtenu en représentant les segments par des vecteurs $(\Delta x, \Delta y)$ avec une distance de Hamming, est de 3.12%. Ce taux d'erreur est significativement supérieur aux 1.04% obtenus en représentant les signatures par le système $[Vx, d_{N2}]$, soit le meilleur système à notre connaissance.

Un défenseur de la segmentation pourrait dire que la représentation des segments adoptée n'est pas complète : on pourrait ajouter les informations de durée, taille, position, etc. Mais ces informations supplémentaires sont corrélées aux premières et pourraient très bien être combinées dans une représentation ponctuelle des signatures. De plus, la suite prouvera que lorsqu'on représente les signatures en combinant deux informations locales (par ex. $[Vx, Vy]$), il faut un ajustement précis du paramètre β représentant l'importance relative des deux types d'information, pour obtenir un meilleur résultat qu'en utilisant la plus discriminante des deux informations locales. Un ajustement adéquat sur un ensemble de données ne sera pas nécessairement adéquat pour de nouvelles données.

On peut donner une explication au résultat trouvé : Δx et Δy peuvent être assimilés aux vitesses moyennes sur le segment. En effet :

$$\begin{aligned}
 \Delta x &= p_x(n) - p_x(1) \\
 &= p_x(n) - p_x(n-1) + p_x(n-1) - p_x(n-2) + \dots - p_x(n-2) + p_x(n-2) - p_x(1) \\
 &= (p_x(n) - p_x(n-1)) + \dots + (p_x(2) - p_x(1)) \\
 &= \sum_{i=1}^{n-1} p_x(i+1) - p_x(i)
 \end{aligned}$$

Or ceci correspond à la sommation des vitesses des points du segment estimées par une différence finie d'ordre 0 (pour le point i , tronquée au point $i+1/2$). On peut faire l'approximation :

$$\begin{aligned}
 \sum_{i=1}^{n-1} p_x(i+1) - p_x(i) &\approx \sum_{i=1}^n v_x \left(i + \frac{1}{2} \right) \\
 &\approx \sum_{i=1}^n v_x(i)
 \end{aligned}$$

Au chapitre 3, on a montré qu'une caractéristique globale, soit la sommation de caractéristiques locales, suivait une loi normale et ne permettait pas une meilleure discrimination entre classes. Les segments ne sont donc pas plus discriminants que les points selon le modèle. En pratique ce n'est peut-être pas tout à fait vrai, seulement ils sont beaucoup moins nombreux que les points : une signature représentée par des segments est représentée par un plus petit nombre de caractéristiques. Comme on a montré qu'en théorie le taux d'erreur est proportionnel à $(n)^{-1/2} \exp^{-n \cdot g(a)}$ (où n est le nombre de caractéristiques représentant les signatures), il est normal que la représentation par segment soit moins adéquate que la représentation ponctuelle.

Cette expérience est une forte indication que la segmentation n'est pas une opération adéquate, si on peut s'en passer pour l'identification des éléments manuscrits.

5.3.2 Choix des caractéristiques

5.3.2.1 Comparaison des différentes caractéristiques locales

Sur le tableau qui suit, on a concentré au maximum les résultats pour comparer en un seul tableau les résultats obtenus selon la caractéristique locale extraite par le système. Sur l'axe horizontal on trouve le type d'information locale choisie (vitesse, accélération, position) pour chacune des 3 mesures de distance, sur l'axe vertical la coordonnée sur laquelle est représentée l'information. Par exemple, si on choisit la colonne accélération puis la sous-colonne d_{N1} , au niveau de la ligne x , on lit le taux d'erreur obtenu avec le système $[Ax, d_{N1}]$.

Tableau 5-4 : Taux d'erreur (%) par coordonnée et type d'information

	Vitesse			Accélération			Position			Information
	d_H	d_{N1}	d_{N2}	d_H	d_{N1}	d_{N2}	d_H	d_{N1}	d_{N2}	?
Vitesse	1.25	2.08	1.04	6.88	6.04	6.04	9.17	9.38	9.17	V
	4.58	4.38	4.38	6.88	6.04	6.46	9.17	10.83	10.00	V
	4.38	3.54	2.92	5.83	4.17	4.38	19.79	13.12	13.54	V
	4.58	3.54	3.75	8.54	5.42	7.08	7.71	8.12	7.92	V
	x	x	x	α	α	α	r	r	r	$[Vx, d_{N2}]$

Analyse des résultats obtenus par type d'information :

Sans aucun doute possible, le meilleur type d'information pour représenter une signature localement est la vitesse. Ce résultat est en accord avec l'étude de [Parizeau 90], qui en était arrivé à la même conclusion. Il semble que l'information « enregistrée » dans le cerveau est l'information de vitesse. En effet les mouvements de la main produisant la signature sont très réguliers d'une signature à l'autre.

L'accélération correspond à la dérivation du signal de vitesse, or la dérivation d'un signal laisse passer en priorité les hautes fréquences. Les hautes fréquences d'une signature sont porteuses de bruit, et ce bruit nuit à la reconnaissance. D'autre part l'accélération est très élevée et peut être très variable au niveau des changements de

direction du tracé. Pour ces deux raisons, l'accélération n'est pas une information très discriminante.

La position correspond à l'intégration du signal de vitesse. Or on a vu au chapitre 3 que cette intégration entraîne une corrélation entre caractéristiques locales très éloignée. Plus des caractéristiques sont corrélées, moins elles sont discriminantes. La piètre performance obtenue avec ce type d'information montre combien les systèmes statiques sont désavantagés, puisqu'ils n'ont accès qu'à l'information de position (en plus sans l'information apportée par la séquence).

Analyse des résultats obtenus par coordonnée :

La meilleure coordonnée de représentation varie selon le type d'information : en effet, x est la meilleure coordonnée pour la vitesse, α pour l'accélération et r pour la position.

La position correspond à la différence entre le point courant et le premier point de la signature. Or cette différence est de moins en moins significative à mesure que l'on s'éloigne en durée du début de la signature. La coordonnée r , qui correspond à la distance du point courant au premier point de la signature contient davantage d'information que les distances horizontale et verticale (x et y), et est donc plus discriminante.

L'angle α de l'accélération (et non l'accélération angulaire !) est le type d'accélération la plus discriminante, alors qu'elle est la moins discriminante pour la position. L'accélération sur la portion centrale d'un segment est essentiellement parallèle

au tracé, alors qu'elle est perpendiculaire au tracé lors d'un changement de direction (voir figure suivante).

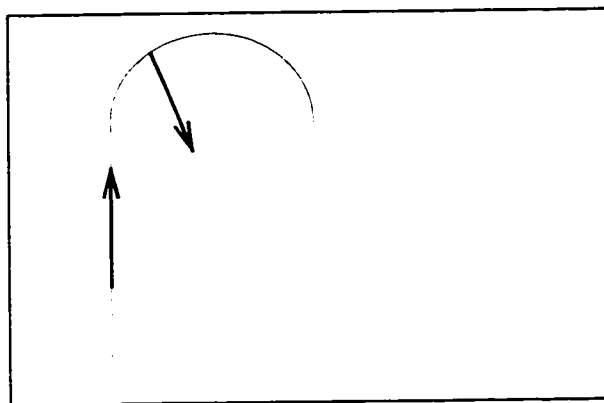


Figure 5-3 : Direction de l'accélération sur le tracé manuscrit

Par conséquent, sur les régions entre les changements de direction, l'angle de l'accélération suit à peu près l'angle de la direction du tracé. Cette information est probablement plus représentative du tracé que les accélération absolue, horizontale, verticale. Cela pourrait expliquer les meilleures performances obtenues avec l'angle de l'accélération.

Les taux d'erreurs obtenus avec la vitesse sont les moins grands, puisque cette information est la plus discriminante, donc la plus utile dans la définition du système de VdS. En particulier, la comparaison de V_x et V_y laissant croire que V_x est l'information la plus importante d'une signature, nécessite une discussion approfondie. Selon le test de McNemar, cette différence est significative.

- **Comparaison x et y :**

En VdS, on considère presque comme une évidence que le signal en y permet une meilleure discrimination des signatures authentiques et imitées que le signal en x. En effet, l'énergie du signal est en majorité distribuée sur l'axe y, ce qui est logique car la séquence spasmodique de mouvements lors de la signature est faite essentiellement de haut en bas. Parizeau, dans son étude comparative de différents algorithmes de VdS [Parizeau 90], en arrivait à la conclusion qu'il valait mieux utiliser le signal en y.

Comment se fait-il que nous soyons en contradiction avec la pensée courante ? Bien que ce résultat soit étonnant, il semble bien que la séquence de mouvements ne soit pas ce qu'il y a de plus difficile à capter pour l'imitateur. Selon nous, les imitateurs arrivaient à relativement bien capter la séquence de mouvements (selon y), ce qui suffisait en général pour que leurs imitations soient acceptées. Seulement, ils n'ont pas toujours réussi à reproduire la forme fidèlement, par exemples les arrondis étaient faits différemment. Et la forme est pour beaucoup décidée par de subtiles déviations du mouvement qui se trouvent selon l'axe x. L'apprentissage de la séquence de mouvements est une chose, le respect de la forme en est une autre. Il semble difficile pour un imitateur de se concentrer sur ces deux aspects à la fois.

Il est possible également qu'il y ait une sorte de « résidu inconscient » chez l'imitateur, qui applique sans s'en rendre compte les caractéristiques propres à sa signature et à son écriture lorsqu'il cherche à imiter un scripteur. Sur la figure suivante, on observe sur la gauche deux signatures authentiques et sur la droite des imitations du même imitateur. Il est frappant de constater que l'imitateur conserve son tracé anguleux

sur les deux imitations ainsi qu'une certaine inclinaison du tracé lorsque les mouvements sont ascendants, qui augmente la vitesse horizontale sur les parties droites du tracé.

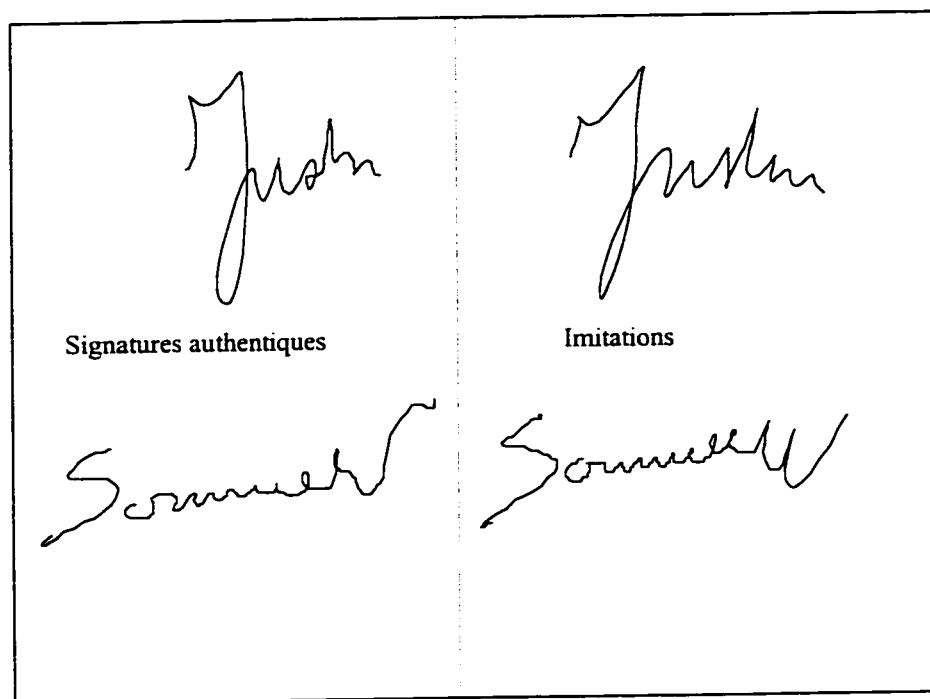


Figure 5-4 : Le résidu inconscient

Lorsqu'un imitateur fait une imitation sur un système dynamique, il cherche d'abord à respecter la séquence verticale de mouvements. Comme il ne peut se concentrer sur tous les aspects de la signature à imiter et qu'il ne peut « oublier » sa propre manière d'écrire, il laisse quelques traces propres à son écriture qui sont essentiellement horizontales.

5.3.2.2 Combinaison de caractéristiques

Des résultats précédents, il apparaît que V_x est l'information locale la plus discriminante. S'il fallait choisir un seul type de caractéristique pour décrire les signatures, on choisirait V_x .

Si l'information V_x est la plus révélatrice de la classe d'une signature-test, l'information contenue dans des caractéristiques comme V_y et V_α n'est pas négligeable pour autant (4.38% et 2.92% d'erreur). Une combinaison adéquate de V_x et de V_y , ou de V_x et de V_α pourrait permettre de représenter les signatures de façon plus précise.

Avec la formule de distance élaborée au chapitre précédent, on a évalué les taux d'erreur pour différentes valeurs de b . On peut évaluer les résultats sur la figure suivante.

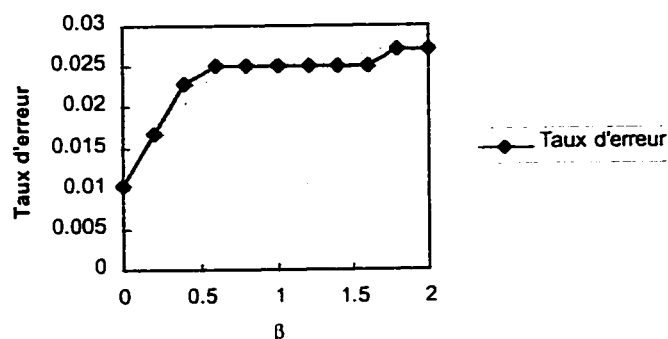


Figure 5-5 : Taux d'erreur pour la représentation $[V_x, V_y]$ des signatures

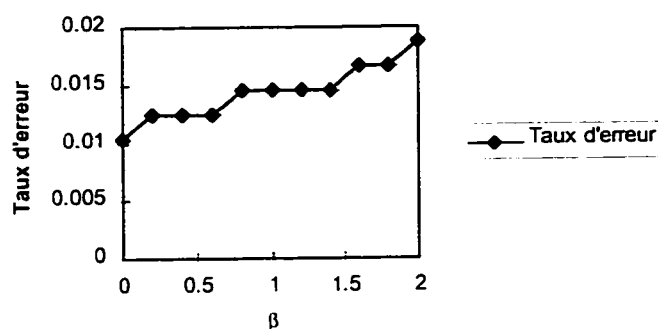


Figure 5-6 : Taux d'erreur pour la représentation $[V_x, V_\alpha]$ des signatures

Les résultats sont négatifs. De fait le taux d'erreur de la représentation combinée, égal à celui de la représentation par V_x pour $\beta=0$, augmente progressivement jusqu'au taux d'erreur de la représentation par V_y lorsque β augmente (de même pour V_α).

Ces informations ne sont pas indépendantes : lorsqu'une imitation est réussie, les différentes informations locales sont aussi bien imitées. Par conséquent, une signature

étant à l'intérieur du seuil d'acceptation pour l'information locale la plus discriminante (V_x) a de fortes chances être à l'intérieur du seuil pour d'autres informations locales (V_y , V_α , ...).

5.3.3 Mesure de distance

Au chapitre précédent, on avait proposé une équation générale de mesure de distance, et on s'était restreint à trois cas: distance de Hamming, distance normalisée par l'écart-type et par la racine de l'écart-type. On peut voir sur le **Erreur! Source du renvoi introuvable.** que les taux d'erreurs obtenus pour les différents types de distance ne sont pas sensiblement différents. En fait, la moyenne des taux d'erreurs pour les douze systèmes différents (accélération, vitesse, position selon x , y et α) est de 7.40% pour d_H , 6.39% pour d_{N1} et 6.39% pour d_{N2} . Il y a une différence entre les distances normalisées et non-normalisées qui semble statistiquement significative.

De fait, la distance D_H donne un taux d'erreur plus élevé que d_{N2} dans 10 systèmes sur 12, et que d_{N1} dans 8 systèmes sur 12.

On a montré au chapitre 4 que sous certaines conditions, la pondération par l'écart-type donne une mesure de distance optimale, du moins lorsqu'on élève la distance au carré. Les résultats obtenus confirment cette idée mais pas de façon décisive.

On peut invoquer l'incertitude sur la valeur exacte sur l'écart-type : soit s^2 l'estimation de la variance σ^2 sur un échantillon de taille 30. Alors s^2 suit une loi du Khi-deux à 29 degré de liberté. L'intervalle de confiance à 80% pour s est alors:

$$0.82 \leq \frac{s}{\sigma} \leq 1.16$$

Autrement dit, pour un échantillon de taille 30 la marge d'erreur de s est suffisamment grande pour que l'écart-type estimé soit une approximation moyennement fiable de l'écart-type réel.

Cette explication est possible, mais les valeurs aberrantes (voir la section sur l'apprentissage au chapitre 4) sont probablement une cause d'incertitude plus importante. On a proposé la médiane et « M.A.D. » comme estimateurs plus robustes. Cependant une évaluation rapide des taux d'erreurs obtenus avec ce type d'estimateur et la caractéristique Vx n'a pas donné d'amélioration.

La cause véritable se situe probablement au niveau des valeurs aberrantes de la signature testée. Certaines caractéristiques locales prennent parfois des valeurs très en dehors de la moyenne, notamment aux régions de changement de direction, où les écarts-type sont plus faibles (et la pondération plus élevée).

5.3.4 Conclusion

Le meilleur résultat obtenu, avec $[Vx, d_{N2}]$, est de 1.04% d'erreurs. Mais comme le système $[Vx, dH]$ obtient toutefois 1.25% d'erreurs, on ne peut différencier ces deux systèmes d'après ces résultats. L'alignement élastique obtient 2.08% d'erreurs. La différence avec le meilleur système est de 5 erreurs sur 500 signatures, mais l'alignement élastique ne répète pas toutes les erreurs de la meilleure méthode. D'après le test de

McNemar, la différence n'est pas significative. Cependant les résultats obtenus avec une fixation automatique du seuil seront significatifs.

5.4 Vérification expérimentale des hypothèses du modèle de système de VdS

5.4.1 Traitement des portions non-identifiées

Le système de VdS élaboré identifie les éléments manuscrits d'une signature-test qui correspondent aux éléments manuscrits de référence. Les éléments manuscrits de la signature-test qui ne sont pas identifiés sont mis à l'écart lors de l'étape de caractérisation. Or la méthode d'alignement élastique utilise ces éléments manuscrits dans le calcul d'une distance entre deux signatures. Et à première vue, il semble évident que le travail d'édition (suppression, insertion et comparaison) devrait être plus important pour une imitation que pour une signature authentique.

On peut calculer, pour une signature-test, le nombre d'éléments qui ne sont pas identifiés sur cette signature. On peut également calculer le coût total rattaché à ces éléments (suppression+insertion). Nous avons utilisé ces deux caractéristiques globales pour représenter les signatures, et calculé les taux d'erreurs obtenus avec une distance d_H :

Tableau 5-5 : Taux d'erreurs pour des signatures caractérisées par leurs éléments non-identifiés

Scripteur	Nombre d'éléments non-reconnus(%)	Coût des éléments non-reconnus(%)
0	3.33	1.67
1	10.00	1.50
2	1.67	3.33
3	0.00	0.00
4	33.33	31.67
5	3.33	0.00
6	1.67	5.00
7	8.33	1.67
Moyenne	7.71	7.29

Les taux d'erreurs obtenus sont de 7.71% et 7.29%. Cette information permet donc une certaine discrimination entre les deux classes. Mais les deux taux d'erreur les plus élevés ont été obtenus avec les scripteurs les plus problématiques (1 et 4). Le scripteur no 4 est particulièrement instable, sa signature comportant de nombreuses inconstances. L'utilisation de ces caractéristiques est nuisible pour les scripteurs problématiques, et risque d'augmenter leur taux d'erreur.

D'après ces résultats, il semble préférable de ne pas tenir compte des portions de signatures non-identifiées par l'alignement élastique.

5.4.2 Détermination du seuil

Au chapitre 4, on a montré que le seuil dépendait du rapport des écarts-type locaux des classes authentiques et imitées. A cause de la grande difficulté à prédire l'écart-type local de la classe non-authentique, on ne peut évaluer le seuil par cette méthode. Aussi la méthode employée pour obtenir les taux d'erreurs jusqu'à maintenant consiste à trouver le seuil pour lequel le taux d'erreur est minimal.

A défaut d'avoir un seuil personnalisé pour chaque scripteur, on peut se contenter d'un seuil fixe qui minimise le taux d'erreur. Bien sûr, un seuil minimisant le taux d'erreur sur un ensemble de signatures ne sera pas nécessairement optimal sur un autre ensemble de signatures. Mais si le taux d'erreur du système est peu sensible aux variations du seuil dans un voisinage du seuil optimal, on est en droit de penser que le seuil trouvé sur un ensemble de signatures d'apprentissage produira un taux d'erreurs acceptable sur un ensemble de signatures-test.

Afin d'avoir une échelle commune pour l'ensemble des scripteurs, on normalise les distances obtenues par les signatures authentiques d'un scripteur par la moyenne des distances obtenues pour ce scripteur : soient $d(X_1, \omega_1), \dots, d(X_N, \omega_1)$ les N distances des N signatures authentiques à la classe w_1 . Alors on normalise ces distances ainsi :

$$d'(X_i, \omega_1) = \frac{d(X_i, \omega_1)}{\left(\frac{\sum_{i=1}^N d(X_i, \omega_1)}{N} \right)}$$

Pour le système choisi à l'issue de la section précédente $[Vx, d_{N2}]$, on a calculé le taux d'erreur sur les 480 signatures pour différentes valeurs de seuil. On peut observer le résultat sur les deux figures suivantes :

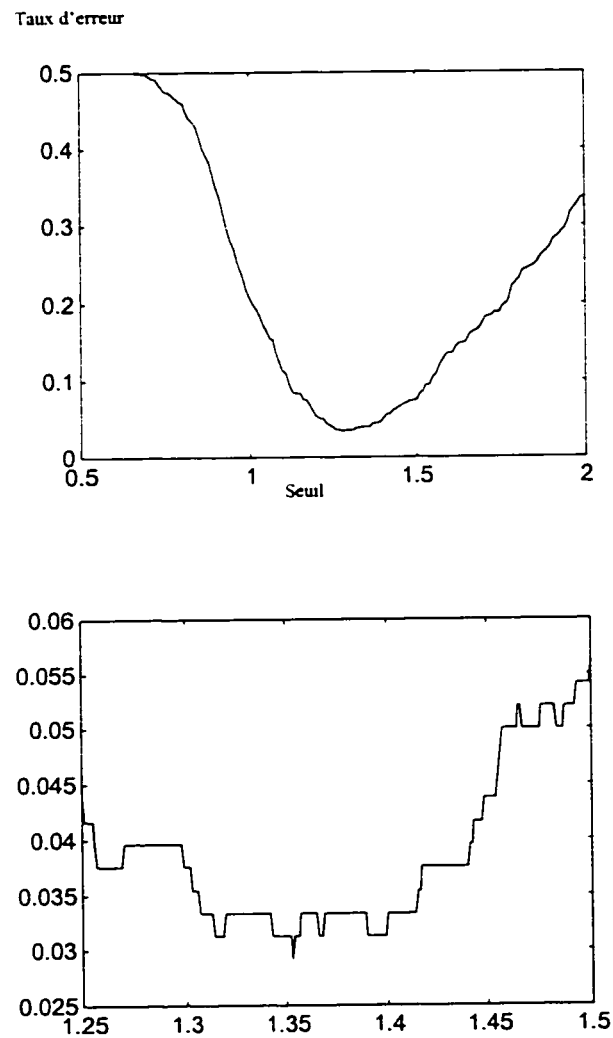


Figure 5-7 : Variation du taux d'erreur en fonction du seuil

Le taux d'erreur minimum obtenu est de 2.92%, pour un seuil de 1.353. Sur la figure du haut, on peut observer que le taux d'erreur est relativement stable autour du minimum. Sur la figure du bas, on voit que lorsque le seuil varie de 1.25 et 1.45, le taux d'erreur reste compris entre 2.92 et 4.25%. L'estimation exacte du seuil optimal n'est donc pas crucial (à 1.25% d'erreurs pers).

La méthode d'alignement élastique, qui avait obtenu 2.08% d'erreur pour un seuil optimal par scripteur, n'obtient plus que 7.29% pour un seuil de 1.41.

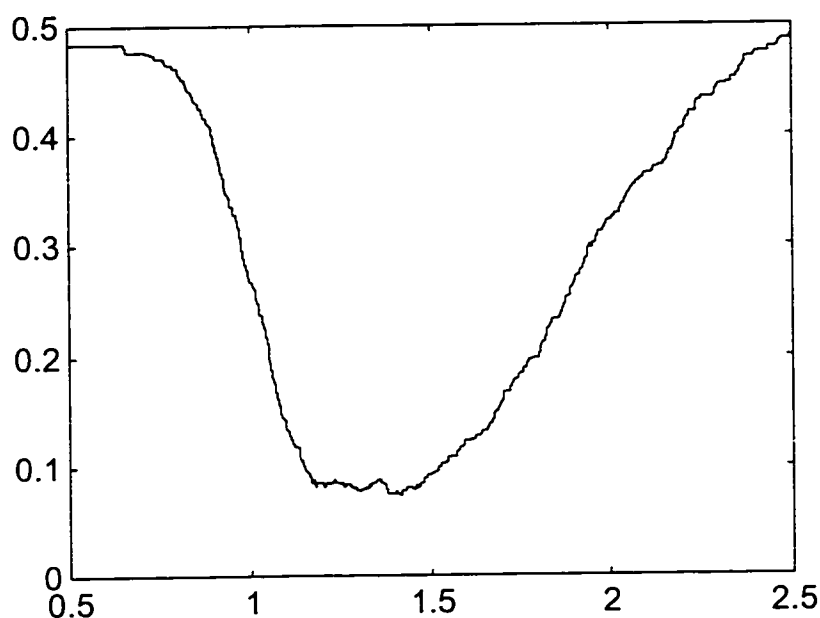


Figure 5-8: Variation du taux d'erreur en fonction du seuil pour l'alignement élastique

5.4.3 Détermination du taux d'erreur en utilisation normale :

Le taux d'erreur obtenu en ajustant le seuil à une valeur optimale pour chacun des scripteurs est fortement biaisé par rapport au taux d'erreur réel. Le taux d'erreur obtenu en ajustant le seuil à une valeur optimale pour l'ensemble des 8 scripteurs est biaisé mais moins fortement. Pour l'évaluation d'un taux d'erreur non-biaisé, il faut séparer l'ensemble des scripteurs servant à évaluer le seuil des ou du scripteur pour lequel le taux d'erreur avec ce seuil est calculé.

Nous avons fait une évaluation du taux d'erreur pour le système $[Vx, d_{N_2}]$ et pour l'alignement élastique par la méthode « leave-one-out » : on trouve le seuil optimal pour un groupe de 7 scripteurs, puis on évalue le taux d'erreur obtenu avec le scripteur n'ayant pas servi à l'apprentissage. On répète ce processus pour les 8 scripteurs, et on déduit le taux d'erreur moyen. Notons que pour cette expérience, il n'y a aucune confusion entre les ensembles d'apprentissage et test.

Le taux d'erreur obtenu pour le système $[Vx, d_{N_1}]$ est de 4.17%, contre 10.21% pour l'alignement élastique. Il est frappant de constater que si la méthode d'alignement élastique produit d'excellent résultats (2.08%) pour un seuil parfaitement choisi, le seuil est très variable en fonction du scripteur, ce qui rend l'application de la méthode problématique

5.5 Conclusion

Les résultats de l'expérimentation permettent de déduire qu'un système de VdS très performant aurait les propriétés suivantes :

- prétraitements : resynchronisation, pas de segmentation
- analyse : extraction de caractéristiques de vitesse horizontale seulement
- comparaison : mesure de distance de Hamming normalisée par la racine de l'écart-type
- décision : le rapport du seuil sur la distance moyenne des signatures de référence est d'approximativement 1.36

Pour cette paramétrisation, le système obtient 4.17% d'erreurs contre 10.21% pour la méthode d'alignement élastique.

Si les conclusions qui ont été tirées sont confirmées par l'application du test de McNemar, on aimerait faire une expérience à plus large échelle pour confirmer ces résultats.

6. CONCLUSION

Dans ce mémoire, on a proposé un nouveau système de VdS basé sur l'identification locale des éléments manuscrits. On a également proposé une modélisation statistique de la VdS. Dans cette conclusion, on fait une revue des contributions majeures de ce mémoire, suivie des directions les plus prometteuses dans lesquelles la recherche pourrait être continuée.

6.1 Contributions

Pour évaluer l'apport de ce mémoire à la recherche en VdS, on revient sur les 4 objectifs que l'on s'était fixé au début du mémoire :

1. **Améliorer les performances de l'alignement élastique** : L'algorithme d'alignement élastique constitue selon nous une des méthodes les plus performantes en VdS. Le système de VdS proposé ici avait pour but initial d'affiner cet algorithme par une évaluation des statistiques locales. Au cours des recherches, l'algorithme d'alignement élastique a progressivement perdu sa place centrale pour devenir un outil d'identification des éléments manuscrits. Le système avec la paramétrisation apparemment la plus adéquat obtient un taux d'erreur de 1.04% contre 2.08% pour l'alignement élastique pour un seuil optimal, et 4.17% contre 10.21% en utilisation « réelle ». **Les performances de l'alignement élastique sont significativement supérieures avec le système élaboré dans cette recherche.**

2. **Élaborer une méthode de détection et d'identification des éléments manuscrits** : La méthode élaborée s'appuie sur l'algorithme d'alignement élastique, et on montre qu'elle est la solution optimale à une mesure d'association des éléments manuscrits qui tient compte du contexte. Cette méthode fonctionne très bien, y compris pour les scripteurs instables. Elle permet d'associer les éléments manuscrits de plusieurs signatures à la fois, et peut servir à une segmentation stable des signatures, appelée « post-segmentation ».
3. **Modéliser statistiquement la VdS** : La modélisation, qui s'appuie sur des distributions gaussiennes des classes authentiques et non-authentiques, permet d'approfondir le problème de la VdS, notamment en montrant que de nombreux modèles font implicitement des hypothèses qui sont discutables (par ex. absence de polarisation pour la classe ω_2). L'évaluation d'un taux d'erreur théorique permet de comparer différents types de systèmes et différents scripteurs. Le modèle ouvre à la voie à une meilleure compréhension du comportement des imitateurs.
4. **Déterminer la meilleure représentation locale des signatures** : Une grande quantité de systèmes de VdS se rattachant à une certaine représentation locale des signatures ont été testés. Les résultats montrent notamment que l'information de vitesse horizontale est la plus discriminante, que la segmentation est une opération nuisible, et que les portions non-identifiées d'une signature ne devraient pas être utilisées pour représenter les signatures.

Cette recherche représente une étape vers des systèmes de VdS qui pourraient dépasser les experts humains en analyse de signature. Elle constitue également un progrès dans un domaine qui fait partie de la RF, en proposant une méthode d'identification des éléments d'une séquence et une modélisation d'un problème de RF à deux classe dont une est inconnue.

6.2 Problèmes non-résolus et recherche future

Bien que la plupart des objectifs initiaux ont été accomplis, plusieurs problèmes ont été partiellement résolus. Des recherches additionnelles sont nécessaires pour résoudre les problèmes suivants :

- Le choix d'une signature de référence pour l'identification des éléments manuscrits est arbitraire, bien qu'il soit atténué par l'étape de resynchronisation. Si le choix se porte par hasard sur une signature « ratée », la forme servant à l'identification des éléments manuscrits sera fausse, et l'identification sera mal faite. Il serait nécessaire de développer une méthode de détection des signatures inadéquates pour la synchronisation
- La classe authentique est parfois représentée par deux ou plusieurs sous-classes de signatures. Bien qu'il soit demandé aux scripteurs de ne signer que d'une seule façon, on ne peut exclure la possibilité qu'un scripteur ait plusieurs façons de signer. Le scripteur no 0 (« Jean ») commençait sa signature de deux façons différentes, sans être

conscient qu'il ne respectait pas la procédure. Il faudrait développer une méthode de « clustering » qui permette de détecter et distinguer des sous-classes d de signatures.

- Les 30 signatures demandées aux scripteurs leur demandait un effort assez grand. Il est probablement irréaliste d'exiger 30 signatures à l'utilisateur d'un système (qui en général ne sait pas que ce nombre a des rapports avec la statistique !). Nous ne savons pas à quel point le nombre de signatures de référence a une influence sur les performances du système, car aucune recherche n'a été entreprise en ce sens. Le système aurait-il des performances similaires avec seulement 5 ou 10 signatures de référence ? Une étude ultérieure est nécessaire pour cela.
- La combinaison de deux informations locales n'a pas donné de résultat positif, malgré l'utilisation d'un facteur de pondération. Il est probable que les signatures mal classées avec la meilleure des deux informations le soient également avec l'autre information, mais ce n'a pas été vérifiée. D'autres recherches doivent être poursuivies (si possible avec de nouvelles signatures) pour étudier la représentation des signatures par plusieurs caractéristiques locales, éventuellement combinées avec des caractéristiques locales.
- Le modèle statistique n'a pas été appliqué aux données, notamment afin de déterminer un coefficient de difficulté d'imitation d'un scripteur applicable. Aucune méthode n'a été proposée pour prédire le rapport des écarts-type locaux et vérifier si ce rapport est à peu près constant pour les éléments manuscrits d'un scripteur. Une prédiction de ce rapport permettrait une meilleure classification des signatures. Il serait utile également

de vérifier l'hypothèse de non-polarité sur un plus grand nombre d'imitateurs. Pour ces recherches, il est nécessaire de disposer de nombreux imitateurs car ces études sont de nature statistique.

- Les résultats obtenus avec la banque de signatures utilisée ont permis de faire plusieurs conclusions sur la paramétrisation adéquate du système de VdS, notamment sur la supériorité de l'information locale V_x sur V_y . Les conclusions sont statistiquement significatives d'après le test statistique de comparaison de deux méthodes. Il serait tout de même intéressant de confirmer les résultats obtenus sur une nouvelle banque de signatures.

La direction de recherche que nous privilégions pour la suite consisterait à chercher à implémenter le raisonnement d'un expert en analyse de signatures. Sur quels éléments se base un expert pour prendre une décision sur la nature d'une signature ? Nous n'en avons pour l'instant aucune idée, mais nous savons du moins qu'il se base essentiellement par des comparaisons locales, Nous possédons maintenant un outil permettant l'identification des éléments manuscrits, il est donc en théorie possible de copier le raisonnement d'un analyste.

Une collaboration entre chercheurs en reconnaissance de forme et intelligence artificielle, et experts en analyse de signature pourrait être très intéressante pour les deux partis. Pour les premiers, la compréhension et la représentation algorithmique des mécanismes humains de perception des formes représente un formidable défi. Pour les

seconds, la formalisation d'une discipline pour laquelle l'expérience et l'intuition sont probablement aussi importants que la technique, apporterait probablement une meilleure compréhension des mécanismes déductifs, ainsi qu'un regard nouveau sur la discipline.

Les outils développés dans ce mémoire permettent une puissance d'observation des signatures impressionnante, qui pourra être utile autant aux chercheurs en reconnaissance de formes qu'aux analystes de signatures. On désire connaître les régions d'une signature qui sont difficiles à imiter ? Il suffit de représenter le rapport des écarts-type locaux par une épaisseur proportionnelle des écarts-type sur la signature moyenne pour le voir.

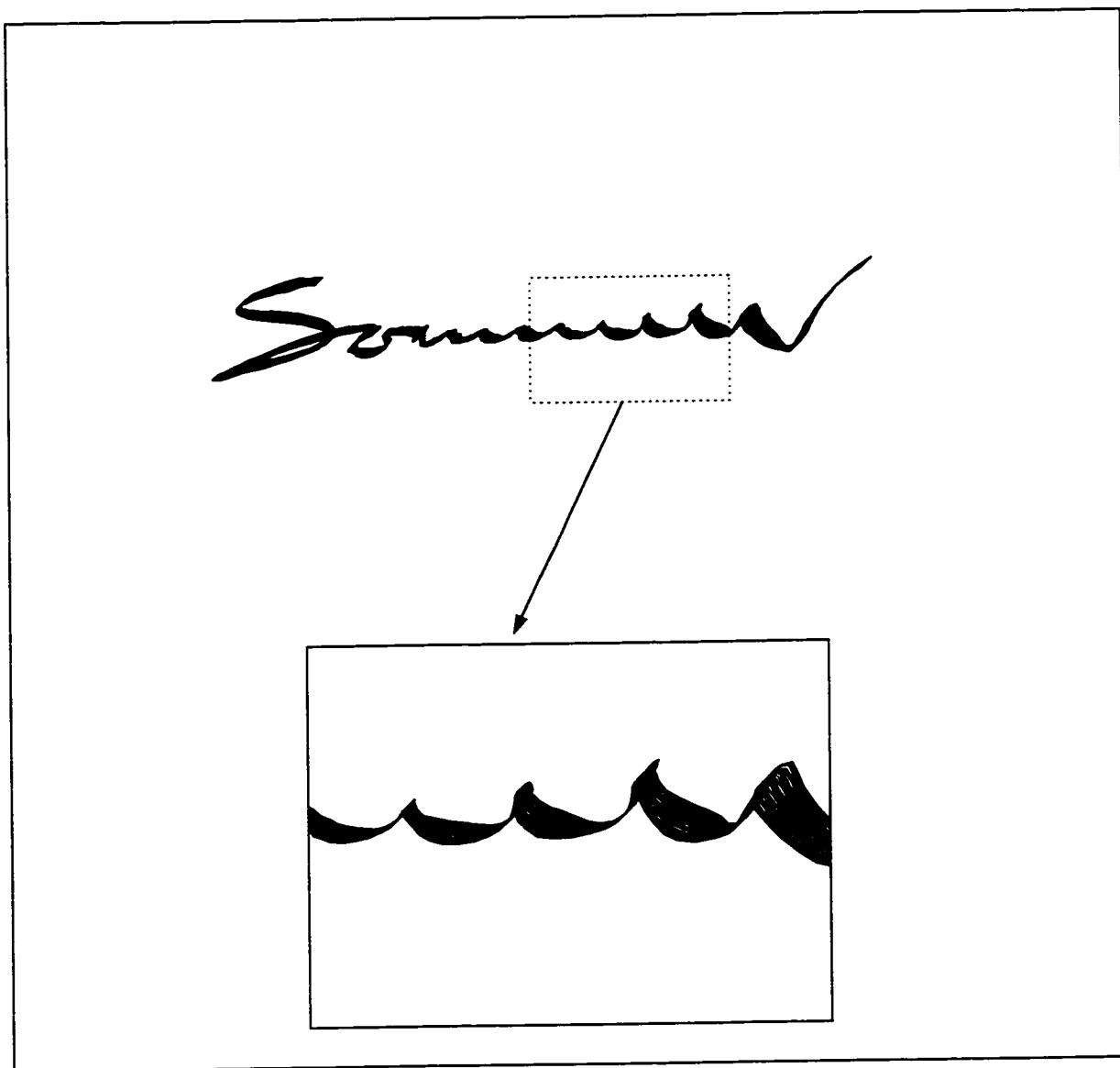


Figure 6-1 : Signature moyenne et rapport des écarts-types locaux

BIBLIOGRAPHIE

- [Babaud 86] BABAUD J., WITKIN A.P., BAUDIN, DUDA R.O., (janvier 1986), Uniqueness of the Gaussian Kernel for Space-Scale Filtering, *IEEE Transactions on Pattern Recognition and Machine Intelligence*, 8(1), 26-33.
- [Barrière 91] BARRIÈRE C., (décembre 1991), *Exploration de l'approche par réseaux neuronaux pour la reconnaissance de symboles manuscrits*, Thèse de doctorat, Ecole Polytechnique de Montréal.
- [Bauer 95] BAUER F, WIRTZ B., (1995), Parameter Reduction and Personalized Parameter Selection for Automatic Signature Verification, *ICDAR 1995*, Montréal, 183-186.
- [Belaid 92] BELAID A., BELAID Y., (1992), *Reconnaissance des Formes : Méthodes et applications*, InterEditions, Paris.
- [Brault 84] BRAULT J.J., PLAMONDON R., (1984), Histogram classifier for characterization of handwritten signature dynamics. In *Proc. Of 7th International Conference on Pattern Recognition*, 619-622.
- [Brault 86a] BRAULT J.J, PLAMONDON R., (mai 1986), Segmenting handwritten signatures at their perceptually important points, *IEEE Transactions on Pattern Recognition and Machine Intelligence*, 15(9), 953-957.

- [Brault 86b] BRAULT J.J., PLAMONDON R., (1986), Coupling visual and dynamic features to study handwritten signature, *Proc. of Graphics Interface 1986*, 375-379.
- [Brault 87a] BRAULT J.J., PLAMONDON R., (juillet 1987),. Handwritten curve partitioning based on geometric and sequential information, *Proceedings of the Third International Symposium on Handwriting and Computer Applications*, Montréal, Canada.
- [Brault 87b] BRAULT J.J., PLAMONDON R., (juillet 1987), Global and local time warping function for handwritten curves comparison, *Proceedings of the Third International Symposium on Handwriting and Computer Applications*, Montréal, Canada.
- [Brault 88] BRAULT J.J., (octobre 1988), *Proposition et vérification d'un coefficient de difficulté d'imitation des signatures manuscrites*, Thèse de doctorat, Ecole Polytechnique de Montréal.
- [Brault 89a] BRAULT J.J., PLAMONDON R., (octobre 1989), How to detect problematic signers for automatic signature verification, *International Carnahan Conference on Security Technology: Electronic Crime Countermeasures*, 127-132.
- [Brault 89b] BRAULT J.J., PLAMONDON R., (1989), Segmentation des signatures manuscrites, *Proc. Of Vision Interface*, 110-116.
- [Brault 91] BRAULT J.J., PLAMONDON R., (1991), Can we measure

the intrinsic complexity of signatures to be imitated?, *Proc. of 5th Handwriting Conference*, 169-171.

- [Brault 93] BRAULT J., PLAMONDON R., (mars-avril 1993), A complexity measure of handwritten curves: modeling of dynamic signature forgeries, *IEEE Transactions on Systems, Man, and Cybernetics*, 23(2), 400-412.
- [Brault 96] BRAULT J.J., (1997), Utilisation d'un réseau de neurones pour localiser les informations pertinentes distribuées sur un signal manuscrit.
- [de Bruyne 90] DE BRUYNE P., (1990), New technologies in credit card authentication, *Proc. Of the 1990 Int. Carnahan Conf. on Security Technology: Crime countermeasures*, 1-5.
- [Dimauro 94] DIMAURO G., IMPEDOVO S., PIRLO G., (1994), Component-oriented algorithms for signature verification, *International Journal of Pattern recognition and Artificial Intelligence*, 8(3), 771-793.
- [Dreyfus 77] DREYFUS S.E., LAW A.M., (1977), *The Art and Theory of Dynamic Programming*, Academic Press, New York.
- [Fischler 86] FISCHLER A., BOLLES R.C., (janvier 1986), Perceptual Organization and Curve Partitionning, journal, *IEEE Transactions on Pattern Recognition and Machine Intelligence*, (8)1, 100-105.
- [Fleiss 81] FLEISS J.L., (1981), *Statistical Methods for Rates and Proportions*, New York, Wiley.
- [Freeman 77] FREEMAN H., DAVIS L., (1977), A corner-Finding Algorithm for Chain-Coded Curves, *IEEE Transactions on*

- Computers*, 3(26), 297-303.
- [Fukunaga 72] FUKUNAGA K., (1972), *Introduction to Statistical Pattern Recognition*, Academic Press, New York.
- [Gonzales 78] GONZALES R.F., THOMASON M.G., (1978), *Syntactic Pattern Recognition, an Introduction*, Addison-Wesley Publishing Co.
- [Han 95] HAN K., SETHI I.K., (1995), Signature Identification via Local Association of Features, *ICDAR 1995*, Montréal, 187-190.
- [Herbst 82] HERBST N.M., LIU C.N., (1982), Automatic Signature Verification, *Computer Analysis and Perception*, volume 1, 83-105, San Diego, CRC Press.
- [Kruskal 83] KRUSKAL J.B., (1983), An Overview of Sequence Comparison, In *Time Warps, String Edits and Macromolecules: the Theory and Practice of Sequence Comparison*, Addison-Wesley.
- [Lalonde 93] LALONDE M., BRAULT J.J., (1993), A neural network approach to handwritten curve partitioning, In *Proc. Of Vision Interface '93*, 136-141.
- [Lalonde 94a] LALONDE M., (août 1994), *Application des réseaux de neurones à la vérification dynamique de signatures*, Ecole Polytechnique.
- [Lalonde 94b] LALONDE M., BRAULT J.J., (1994), Comparison of sequences generated by a self-organizing feature map using dynamic programming, *World Congress on Neural Network*,

page 5, San Diego.

- [Lams 89] LAMS C.F., KAMINS D., (1989), Signature Recognition Through spectral Analysis, journal, *IEEE Transactions on Pattern Recognition and Machine Intelligence*, 22(1), 39-44.

- [Laurière 79] LAURIERE J.L., (1979), *Eléments de Programmation Dynamique*, Bordas, Paris.

- [Leclerc 92] LECLERC F., (1992), Des gaussiennes pour la modélisation des signatures et la segmentation de tracés manuscrits, *Traitement du signal*, 9(4), 347-358.

- [Leclerc 94] LECLERC F, PLAMONDON R., (1994), Automatic Signature Verification: The State of the Art 1989-1993, *International Journal of Pattern recognition and Artificial Intelligence*, 8(3), 643-660.

- [Mandelbrot 84] MANDELBROT B., (1984), *Les objets fractals*, Flammarion, Paris.

- [Marr 82] MARR D., (1982), *Vision*, W.H. Freeman and Co..

- [Menier 95] MENIER G., (septembre 1995), *Système en ligne d'écriture cursive manuscrite*, Université de Rennes I.

- [Moran 68] MORAN P.A.P., (1968), *An introduction to probability theory*, Oxford Scientific Publications.

- [Murshed 95] MURSHED N.A., BORTOLOZZI F., SABOURIN R., (1995), Off-Line signature Verification Without a Priori Knowledge of Class ω_2 : A New Approach, *ICDAR 1995*, Montréal, 191-196.

- [Nelson 94] NELSON W., TURIN W., HASTIE T., (1994), Statistical methods for on -line signature verification, *International Journal of Pattern recognition and Artificial Intelligence*, 8(3), 749-770.
- [Parizeau 90] PARIZEAU M., PLAMONDON R., (juillet 1994), A comparative analysis of regional correlation, dynamic time warping, and skeletal tree matching for signature verification, *IEEE Transactions on Pattern Recognition and Machine Intelligence*, 12(7), 710-717.
- [Plamondon 92a] PLAMONDON R., (septembre 1992), A model-based segmentation framework for computer processing of handwriting, In *Proc. of 11th International Conference on Pattern Recognition*, The Hague, 303-307.
- [Plamondon 92b] PLAMONDON R., YERGEAU P., BRAULT J.J., (1992), A multi-level signature verification system, In *From Pixels to Features III: Frontiers in Handwriting Recognition*, Elsevier Science Publishers B. V. (North-Holland), 363-370.
- [Plamondon 94] PLAMONDON R., (1994), The Design of an On-line Signature Verification System: From Theory to Practice, *International Journal of Pattern recognition and Artificial Intelligence*, 8(3), 795-811.
- [Raphael 74] RAPHAEL D.E., YOUNG J.R., (1974), « automated Personal Identification », *Long Range Planning Service*, Stanford Research Institute, Menlo Park, CA.
- [Ripley 96] RIPLEY B.D., (1996), *Pattern Recognition and Neural*

Networks, Cambridge University Press.

- [Sato 82] SATO Y., KOGURE K., (1982), On-line Signature Verification Based on Shape, Motion, and Writing Pressure, *Proc. International Conference on Pattern Recognition*, Munich, 823-826.
- [Tappert 82] TAPPERT C.C., (novembre 1982), Cursive Script Recognition by Elastic Matching, *IBM J. Res. Develop.*, 6(26), 765-771.
- [Unnikrishnan 91] UNNIKRISHNAN K.P., HOPFIELD J.J., TANK D.W., (1991), Connected-digit speaker-dependant speech recognition using a neural network with time-delayed connections, *IEEE. Trans. On Signal Processing*, 39(3), 698-713.
- [Wagner 74] WAGNER R.A., FISCHER M.J., (janvier 1974), The String-to-String Correction Problem, *Journal of the ACM*, 21(1), 168-173.
- [Wirtz 95] WIRTZ B., (1995), Stroke-Based Time Warping for Signature Verification, *ICDAR 1995*, Montréal, 179-182.
- [Worthington 85] WORTHINGTON T.K., (1985), IBM Dynamic Signature Verification, in *Computer Security, IFIP 1985*, Elsevier Science Publisher, 129-154
- [Zimmerman 85] ZIMMERMAN K.P., VARADY M.J., (1985), Handwriter Identification from One-Bit Quantized Pressure Patterns, *IEEE Transactions on Pattern Recognition and Machine Intelligence*, 18(1), 63-72.

ANNEXE

 D_H D_{N1} D_{N2} V_x

0	0,00%	1,57
1	5,00%	1,42
2	0,00%	1,59
3	0,00%	1,27
4	3,33%	1,24
5	0,00%	1,47
6	0,00%	1,56
7	1,67%	1,23

0	0,00%	1,62
1	5,00%	1,41
2	1,67%	1,26
3	0,00%	1,31
4	1,67%	1,24
5	0,00%	1,51
6	0,00%	1,47
7	0,00%	1,34

0	0,00%	1,72
1	6,67%	1,23
2	3,33%	1,15
3	0,00%	1,30
4	3,33%	1,21
5	1,67%	1,44
6	1,67%	1,34
7	0,00%	1,21

 V_y

0	6,67%	1,42
1	3,33%	1,36
2	0,00%	1,45
3	0,00%	1,25
4	13,33%	1,06
5	5,00%	1,24
6	1,67%	1,23
7	6,67%	1,16

0	5,00%	1,39
1	5,00%	1,26
2	0,00%	1,52
3	0,00%	1,27
4	13,33%	1,07
5	3,33%	1,31
6	3,33%	1,17
7	5,00%	1,15

0	3,33%	1,32
1	5,00%	1,44
2	0,00%	1,48
3	0,00%	1,25
4	15,00%	1,04
5	1,67%	1,42
6	3,33%	1,19
7	6,67%	1,13

 A_x

0	1,67%	1,42
1	16,67%	1,17
2	3,33%	1,30
3	0,00%	1,22
4	20,00%	1,12
5	3,33%	1,37
6	3,33%	1,28
7	6,67%	1,12

0	3,33%	1,30
1	16,67%	1,17
2	3,33%	1,29
3	0,00%	1,25
4	21,67%	1,22
5	1,67%	1,42
6	0,00%	1,33
7	1,67%	1,16

0	5,00%	1,34
1	15,00%	1,10
2	5,00%	1,18
3	0,00%	1,27
4	21,67%	1,05
5	1,67%	1,44
6	0,00%	1,32
7	0,00%	1,21

 A_y

0	11,67%	1,31
1	13,33%	1,13
2	0,00%	1,42
3	1,67%	1,21
4	20,00%	1,04
5	0,00%	1,47
6	1,67%	1,18
7	6,67%	1,16

0	10,00%	1,34
1	13,33%	1,10
2	0,00%	1,43
3	1,67%	1,22
4	20,00%	1,03
5	0,00%	1,46
6	1,67%	1,21
7	5,00%	1,14

0	8,33%	1,25
1	13,33%	1,16
2	0,00%	1,44
3	1,67%	1,24
4	20,00%	1,03
5	0,00%	1,45
6	1,67%	1,23
7	3,33%	1,18

 P_x

0	3,33%	1,58
1	26,67%	0,88
2	1,67%	1,47
3	1,67%	1,55
4	21,67%	1,00
5	13,33%	1,13
6	3,33%	1,53
7	1,67%	1,31

0	5,00%	1,35
1	25,00%	0,83
2	3,33%	1,41
3	1,67%	1,58
4	20,00%	1,03
5	13,33%	1,21
6	3,33%	1,53
7	1,67%	1,30

0	5,00%	1,40
1	26,67%	0,83
2	5,00%	1,37
3	1,67%	1,60
4	18,33%	1,09
5	13,33%	1,21
6	3,33%	1,55
7	1,67%	1,32

Py

0	8,33%	1,29
1	10,00%	1,30
2	10,00%	1,38
3	6,67%	1,33
4	16,67%	0,98
5	6,67%	1,47
6	0,00%	1,58
7	15,00%	1,20

0	8,33%	1,25
1	11,67%	1,16
2	11,67%	1,14
3	8,33%	1,43
4	20,00%	0,94
5	5,00%	1,49
6	0,00%	1,48
7	15,00%	1,03

0	5,00%	1,44
1	11,67%	1,10
2	11,67%	1,23
3	10,00%	1,45
4	20,00%	0,92
5	10,00%	1,32
6	0,00%	1,48
7	18,33%	0,96

Vr

0	8,33%	1,30
1	10,00%	1,33
2	0,00%	1,44
3	0,00%	1,33
4	11,67%	1,05
5	5,00%	1,31
6	0,00%	1,50
7	1,67%	1,22

0	5,00%	1,26
1	6,67%	1,29
2	0,00%	1,43
3	0,00%	1,32
4	11,67%	1,03
5	3,33%	1,38
6	1,67%	1,38
7	1,67%	1,29

0	6,67%	1,28
1	5,00%	1,28
2	1,67%	1,32
3	0,00%	1,29
4	10,00%	1,09
5	1,67%	1,41
6	1,67%	1,35
7	1,67%	1,20

Va

0	0,00%	1,45
1	10,00%	1,15
2	11,67%	1,02
3	0,00%	1,18
4	6,67%	1,15
5	3,33%	1,27
6	0,00%	1,29
7	3,33%	1,24

0	0,00%	1,64
1	6,67%	1,17
2	5,00%	1,09
3	0,00%	1,20
4	5,00%	1,23
5	3,33%	1,29
6	0,00%	1,33
7	3,33%	1,29

0	1,67%	1,40
1	6,67%	1,25
2	6,67%	1,13
3	0,00%	1,24
4	6,67%	1,16
5	3,33%	1,33
6	0,00%	1,39
7	3,33%	1,19

Ar

0	8,33%	1,15
1	20,00%	1,19
2	1,67%	1,39
3	0,00%	1,23
4	20,00%	1,10
5	3,33%	1,44
6	8,33%	1,11
7	6,67%	1,18

0	8,33%	1,16
1	18,33%	1,11
2	0,00%	1,41
3	0,00%	1,24
4	20,00%	1,06
5	3,33%	1,37
6	3,33%	1,28
7	3,33%	1,25

0	8,33%	1,20
1	15,00%	1,10
2	1,67%	1,28
3	0,00%	1,20
4	10,00%	1,12
5	3,33%	1,37
6	1,67%	1,35
7	3,33%	1,17

Aa

0	5,00%	1,36
1	5,00%	1,15
2	11,67%	1,03
3	1,67%	1,17
4	21,67%	1,08
5	0,00%	1,48
6	0,00%	1,23
7	1,67%	1,16

0	3,33%	1,39
1	8,33%	1,10
2	6,67%	1,08
3	0,00%	1,19
4	16,67%	1,09
5	0,00%	1,48
6	0,00%	1,34
7	0,00%	1,22

0	3,33%	1,44
1	6,67%	1,13
2	6,67%	1,29
3	0,00%	1,21
4	13,33%	1,13
5	0,00%	1,55
6	1,67%	1,46
7	1,67%	1,20

Pr

0	1,67%	1,89
1	25,00%	0,89
2	5,00%	1,37
3	1,67%	1,51
4	15,00%	1,05
5	10,00%	1,29
6	1,67%	1,73
7	1,67%	1,32

0	3,33%	1,47
1	25,00%	0,83
2	6,67%	1,52
3	1,67%	1,45
4	15,00%	1,05
5	6,67%	1,38
6	3,33%	1,55
7	1,67%	1,29

0	3,33%	1,61
1	26,67%	0,88
2	6,67%	1,61
3	1,67%	1,57
4	16,67%	1,04
5	5,00%	1,33
6	3,33%	1,59
7	1,67%	1,27

Pa

0	18,33%	1,16
1	46,67%	1,44
2	5,00%	1,46
3	10,00%	1,17
4	25,00%	1,11
5	10,00%	1,21
6	25,00%	1,80
7	18,33%	1,04

0	11,67%	1,28
1	26,67%	0,44
2	3,33%	1,36
3	6,67%	1,33
4	25,00%	0,97
5	3,33%	1,45
6	18,33%	1,02
7	13,33%	0,81

0	13,33%	1,13
1	25,00%	0,37
2	1,67%	1,59
3	11,67%	1,18
4	23,33%	0,95
5	5,00%	1,34
6	13,33%	1,10
7	11,67%	0,74

Sans resynchronisation

[Vx,dH]

0	0,00%	1,51
1	5,00%	1,36
2	0,00%	1,50
3	0,00%	1,27
4	13,33%	1,23
5	1,67%	1,40
6	0,00%	1,52
7	1,67%	1,25

[Vy,dH]

0	8,33%	1,40
1	3,33%	1,29
2	1,67%	1,40
3	0,00%	1,24
4	15,00%	1,04
5	3,33%	1,42
6	3,33%	1,20
7	6,67%	1,22

[Vx,dN1]

0	0,00%	1,60
1	5,00%	1,30
2	1,67%	1,19
3	0,00%	1,27
4	13,33%	1,10
5	1,67%	1,39
6	0,00%	1,45
7	0,00%	1,27

[Vy,dN1]

0	5,00%	1,42
1	1,67%	1,37
2	1,67%	1,40
3	0,00%	1,26
4	15,00%	1,07
5	3,33%	1,40
6	1,67%	1,20
7	6,67%	1,12

Segmentation

[Dx,DH]

0	1.67%	1.5
1	8.33%	1.27
2	5.00%	1.23
3	1.67%	1.39
4	11.67%	1.13
5	5.00%	1.57
6	1.67%	1.56
7	0.00%	1.40

[Dy,DH]

0	6.67%	1.64
1	1.67%	1.36
2	11.67%	1.36
3	0.00%	1.47
4	18.33%	1.10
5	1.67%	1.48
6	1.67%	1.52
7	3.33%	1.28

[Dx,DN1]

0	1.67%	1.34
1	11.67%	0.82
2	11.67%	1.06
3	18.33%	0.51
4	10.00%	1.21
5	8.33%	1.28
6	20.00%	0.42
7	11.67%	3.08

[Dy,dN1]

0	13.33%	1.16
1	8.33%	0.80
2	15.00%	0.85
3	20.00%	0.25
4	16.67%	1.14
5	3.33%	1.58
6	16.67%	0.27
7	11.67%	3.09

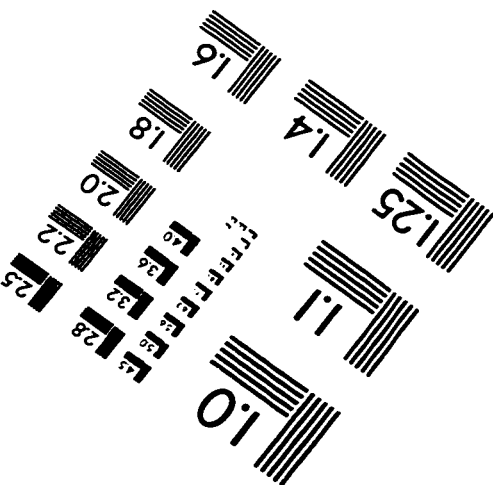
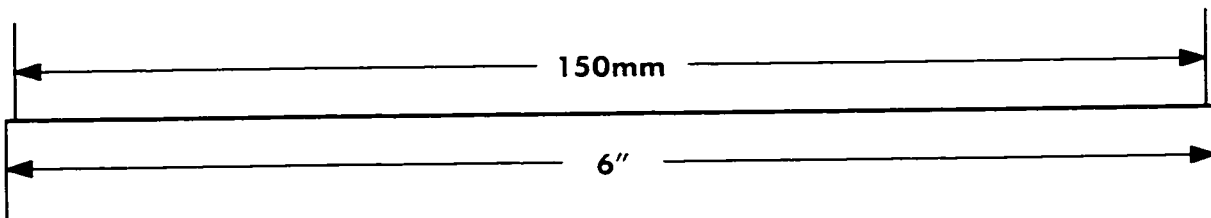
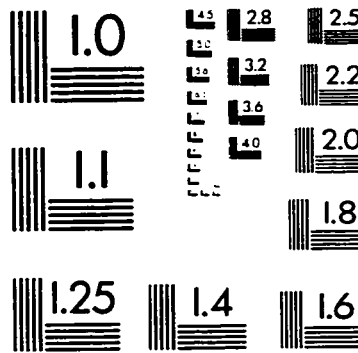
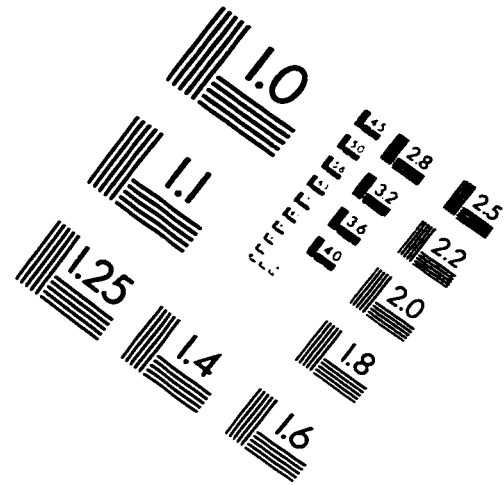
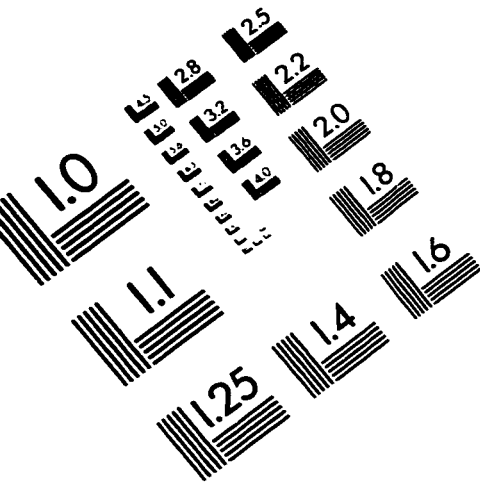
[Dx+Dy,dH]

0	0.00%	1.50
1	6.67%	1.29
2	8.33%	1.33
3	0.00%	1.43
4	6.67%	1.20
5	1.67%	1.41
6	1.67%	1.50
7	0.00%	1.25

[Dx+Dy,dN1]

0	1,67%	1,35
1	10,00%	0,87
2	13,33%	0,87
3	13,33%	0,46
4	6,67%	1,20
5	3,33%	1,31
6	16,67%	0,49
7	11,67%	2,68

IMAGE EVALUATION TEST TARGET (QA-3)



APPLIED IMAGE, Inc.
1653 East Main Street
Rochester, NY 14609 USA
Phone: 716/482-0300
Fax: 716/288-5989

© 1993, Applied Image, Inc., All Rights Reserved

