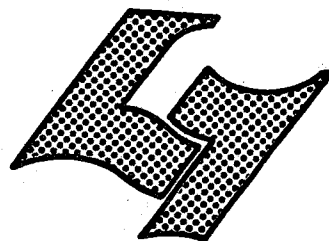


UNIVE

DESS
1980
7
AERNARD LYON-I
11 Novembre 1918
621 VILLEURBANNE

Diplôme d'Etudes Supérieures Spécialisées

Formatique documentaire

* ~~MEMOIRE DE STAGE~~

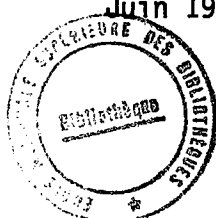
* NOTE DE SYNTHESE

OUTILS STATISTIQUES POUR LE TRAITEMENT AUTOMATIQUE

DES LANGUES.

AUTEUR : VERBAERE Catherine

DATE : Juin 1980

DESS
1980
7
A

I INTRODUCTION

II RECHERCHE BIBLIOGRAPHIQUE

III ANALYSE MORPHOLOGIQUE ET ANALYSE SYNTAXIQUE

A - ANALYSE MORPHOLOGIQUE

B - ANALYSE SYNTAXIQUE

C - AMBIGUITES MORPHOSYNTAXIQUES

IV OUTILS STATISTIQUES

A - PREMIERE APPROCHE

B - GRAMMAIRE PROBABILISTE

C - FILTRE STATISTIQUE

D - CORRECTIONS D'ERREURS SYNTAXIQUES

V CONCLUSION. -

INTRODUCTION

L'utilisation de l'analyse automatique des langues est une voie vers la solution des problèmes de traitement automatique de longs textes en langue naturelle.

Les méthodes linguistiques traditionnelles ne suffisent pas pour lever les ambiguïtés lors de l'analyse morphologique et de l'analyse syntaxique. L'analyse sémantique qui suit l'analyse syntaxique devrait éliminer ces ambiguïtés. Malheureusement l'analyse sémantique est difficilement concevable en traitement automatique des langues.

Depuis quelques années de nouvelles méthodes d'analyses sont apparues. Ce sont des méthodes probabilistes et statistiques. Elles permettent de lever certaines des ambiguïtés citées plus haut. Ces procédés statistiques ont été au départ utilisés pour l'indexation de textes. Ils étaient constitués de simples comptages de mots dans le texte, de calculs de fréquences, de calculs de corrélation. Ces méthodes se sont considérablement développées utilisant la théorie de l'échantillonnage, les filtres statistiques, les fonctions discriminantes, les processus de Markov ou la méthode du maximum de vraisemblance.

Ces méthodes permettent d'obtenir de très bons résultats sur un univers de mots réduit. Ils le seraient beaucoup moins dès qu'on élargit cet univers. On notera que ce problème est la plus grande difficulté du traitement automatique des langues.

Nous allons donc voir les différentes façons d'utiliser les outils statistiques dans le traitement automatique des langues. Les premières pages de cet exposé seront consacrées à la recherche bibliographique faite en vue de cet exposé.

II RECHERCHE BIBLIOGRAPHIQUE.

Il est important d'expliquer de quelles façons a été abordée cette bibliographie. La base de cette recherche est un compte rendu de recherche de l'UER Informatique et Mathématiques en sciences sociales :

SIRAUCO un système d'aide à l'étude morphosyntaxique de corpus - Par BOURGUIGNON HOLLARD et ROUAULT. [1]

On a pu de ce document retirer les mots clés qui semblaient intéressants - On a obtenu la liste suivante -

PROBABILITES
STATISTIQUES
ANALYSE MORPHOLOGIQUE.
ANALYSE SYNTAXIQUE.
ANALYSE MORPHOSYNTAXIQUE.
GRAMMAIRE.
INFORMATION RETRIEVAL.

Faisant un stage dans la société TELESYSTEMES QUESTEL : serveur français de bases de données. J'ai profité de ma formation sur le logiciel MISTRAL. pour interroger le fichier PASCAL : fichier du CNRS -

Il faut préciser pour comprendre le légitime

qui suit que dans le logiciel MISTRAL, les troncatures sont ? (1 ou 0 caractères) et + (un de caractères illimité) et la procédure d'enroulement alphabétique d'un terme est la procédure LE.

On a obtenu le listing suivant (voir page suivante). Il faut noter que parmi les 32 documents 10 seront retenus, les autres étant rejetés à cause de la langue ou du sujet trop éloigné du traitement automatique des langues.

Pour terminer complètement cette recherche. Il a fallu obtenir les documents primaires. Pour cela il existe le centre de documentation du CNRS. A l'aide d'un code des périodiques reçus dans ce centre [LOC - CNRS - 15077 sur l'exemple page 7] on obtient le document.

On peut dire dès à présent que ce sujet n'a pas fait l'objet d'une grande publication surtout en langue anglaise ou française. En effet lors de mon interrogation j'ai dû supprimer environ $\frac{1}{2}$ des références qui étaient en langue russe allemande ou japonaise.

COM

??FRT0440804801401 V0204 P02
?? ST00 YOUR NAME IS CL02

DAY:0128,HOUR:0013,MIN:0058

??CENTSSTS IS CONNECTED
PLEASE LOGIN: 0402,04201:89
PASSWORD: ■■■■

DAY:0128,HOUR:0013,MIN:0058

LOGIN SUCCESSFUL. YOUR ID IS X287
NOUVELLES SUR: BIBLIOS, PAPYRUS, IALINE.
!MISTRAL

QUESTEL mistral V5-00-01

FRANCAIS, FRAPPER :1
ENGLISH, TYPE :2

?1

TERMINAL (T/VXX)

?T

PROCEDURE: ..INFO, ..SOS, ..BASE ?

?..BA PASCAL

BASE CONNEXÉE: PASCAL
PROCEDURE, OU ETAPE DE RECHERCHE 1

?STATISTIQUE?

TERME MULTISENS STATISTIQUE?: 4

T1 STATISTIQUE /DE
T2 STATISTIQUE /UT
T3 STATISTIQUES /DE
T4 STATISTIQUES /UT
SELECTIONNER OU NON ?

?S TT

1 RESULTAT 25973
PROCEDURE, OU ETAPE DE RECHERCHE 2

?PROBABILITE?

TERME MULTISENS PROBABILITE?: 3

T1 PROBABILITE /DE
T2 PROBABILITE /UT
T3 PROBABILITES /UT
SELECTIONNER OU NON ?

?S TT

2 RESULTAT 6658
PROCEDURE, OU ETAPE DE RECHERCHE 3

?STOCHASTIQUE?

TERME MULTISENS STOCHASTIQUE?: 2

T1 STOCHASTIQUE /UT
T2 STOCHASTIQUES /UT
SELECTIONNER OU NON ?

?S TT

3 RESULTAT 6030

?ANALYSE ET SYNTAXIQUE

TERME MULTISENS ANALYSE: 2

T1 ANALYSE /DE

T2 ANALYSE /UT

SELECTIONNER OU NON ?

?S TT

RESULTAT 130236

4 RESULTAT 740

PROCEDURE, OU ETAPE DE RECHERCHE 5

?ANALYSE ET MORPHOLOGIQUE

TERME MULTISENS ANALYSE: 2

T1 ANALYSE /DE

T2 ANALYSE /UT

SELECTIONNER OU NON ?

?S TT

RESULTAT 130236

5 RESULTAT 240

PROCEDURE, OU ETAPE DE RECHERCHE 6

?ANALYSE ET MORPHOSYNTAXIQUE

TERME MULTISENS ANALYSE: 2

T1 ANALYSE /DE

T2 ANALYSE /UT

SELECTIONNER OU NON ?

?S TT

RESULTAT 130236

6 RESULTAT 4

PROCEDURE, OU ETAPE DE RECHERCHE 7

?GRAMMAIRE

TERME MULTISENS GRAMMAIRE: 2

T1 GRAMMAIRE /DE

T2 GRAMMAIRE /UT

SELECTIONNER OU NON ?

?S TT

7 RESULTAT 565

PROCEDURE, OU ETAPE DE RECHERCHE 8

?..HI

ETAPE FREQ

1 25973 STATISTIQUE?

2 6658 PROBABILITE?

3 6030 STOCHASTIQUE?

4 740 ANALYSE ET SYNTAXIQUE

5 240 ANALYSE ET MORPHOLOGIQUE

6 4 ANALYSE ET MORPHOSYNTAXIQUE

7 565 GRAMMAIRE

PROCEDURE, OU ETAPE DE RECHERCHE 9

?1 OU 2 OU 3

9 RESULTAT 35863

PROCEDURE, OU ETAPE DE RECHERCHE 10

?4 OU 5 OU 6 OU 7

10 RESULTAT 1211

PROCEDURE, OU ETAPE DE RECHERCHE 11

?9 ET 10

11 RESULTAT 52

PROCEDURE, OU ETAPE DE RECHERCHE 12

PROCEDURE, OU ETAPE DE RECHERCHE 13

?ANALYSE MORPHOLOGIQUE

13 RESULTAT 50

PROCEDURE, OU ETAPE DE RECHERCHE 14

?4 OU 6 OU 7 OU 13

14 RESULTAT 1025

PROCEDURE, OU ETAPE DE RECHERCHE 15

?9 ET 14

15 RESULTAT 32

2684 C.PASCAL

: 78-3-0295447

: DELATTE L.; DUCHESNE-DEGEY M.; GOVAERTS S.; DENOOZ J.

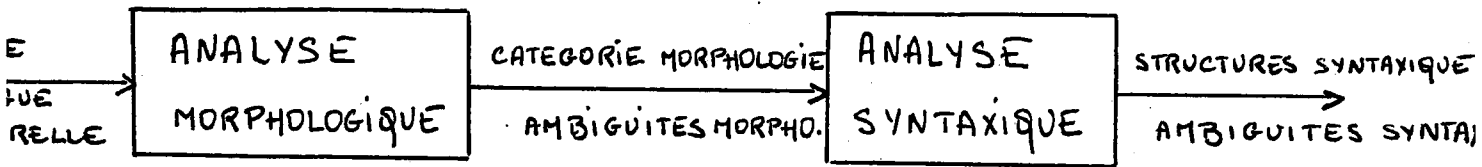
: NIV. LIEGE, LAB. ANAL. STAT. LANGUES ANC., LIEGE, BELG.

: LE TRAITEMENT AUTOMATIQUE DE LA LANGUE FRANCAISE AU LABORATOIRE
D'ANALYSE STATISTIQUE DES LANGUES ANCIENNES.

: TP;LM;DU

: ORGANIS. INTERNATION. ET. LANGUES ANC. ORDINAT., REV.; BELG.; DA. 1977;
NO 4; PP. 1-55; BIBL. 6 REF.; LOC. CNRS-15077

I ANALYSE MORPHOLOGIQUE ET ANALYSE SYNTAXIQUE.



A. ANALYSE MORPHOLOGIQUE.

L'analyse morphologique est l'analyse des différentes unités lexicales composant le texte en langue naturelle [1] Une décomposition sera faite en éléments suivants : la base (ou racine), les affixes (préfixe ou suffixe) et la desinence.

PREFIXE	BASE	SUFFIXE	DESINENCE
---------	------	---------	-----------

On traite un texte en langue naturelle - la langue est un ensemble infini - ou il existe un nombre de phénomènes différents - Ainsi on aura des difficultés dans

→ le choix de la base pour un terme
|IR|REMEDI|ABLE| base REMEDI
dans ce cas REMEDE sera analysé à part

→ le choix de dislocation ou non d'expression
POMME DE TERRE.

→ le choix des catégories morphologiques

catégorie NOM. ou catégorie NOM PROPRE

catégorie NOM COMMUN

On peut donc affiner ou non une catégorie à l'aide de variables.

Ainsi l'analyseur morphologique avec l'aide d'un dictionnaire de références ou sont répertoriées : bases, dérives, affixes, fera correspondre le terme à analyser avec un modèle suivant les règles morphologiques. Il y aura des ambiguïtés morphologiques.

B. ANALYSE SYNTAXIQUE

L'analyse syntaxique d'un texte en langue naturelle fonctionne avec une grammaire donnée a priori. L'ensemble des règles appartenant à la grammaire ne doit pas être trop grand, cela augmenterait les difficultés. Il faut donc analyser le plus grand nombre de phrase avec un nombre de règles correcte.

L'analyseur syntaxique valide une structure de phrase en fonction de la grammaire choisie au départ. Il pourra ainsi résoudre des ambiguïtés morphologiques. Il y aura cependant des ambiguïtés syntaxiques.

C - LES AMBIGUITES MORPHOSYNTAXIQUES.

Le traitement automatique des langues comporte une analyse morphologique et une analyse syntaxique - étant entendu que l'analyse sémantique ne sera pas abordée ici.

Lors de l'analyse morphosyntaxique d'un texte plusieurs obstacles se présentent. Ces obstacles sont dus à la grande richesse de la langue naturelle - les ambiguïtés sont donc morphologiques [5]

Exemple : SI adverbe affirmation
adverbe intensité
conjonction de subordination

POUVOIR	}	substantif verbe.
ETRE		
AVOIR		

UN : numéral.
pronom indéfini.
article.

Ces exemples montrent que lors de l'analyse morphologique d'une phrase le nombre de combinaisons possibles devient vite très grand.

Cependant l'analyse syntaxique éliminera beaucoup de ces ambiguïtés - Mais elle en générera

d'autres : des ambiguïtés syntaxiques -

Exemple: IL Y A UNE MAISON.

Cette phrase peut être analysée.

[IL Y A] [UNE MAISON]

Ou

[IL] [Y] [A] [UNE MAISON]

sous entendu à cet endroit.

Pour ce type d'ambiguïtés seule une étude du contexte de la phrase permet de savoir quelle est l'analyse correcte. Cet aspect est abordé normalement dans l'analyse sémantique.

IV OUTILS STATISTIQUES

Nous nous apercevons qu'une nouvelle approche du traitement automatique des langues est nécessaire, les outils classiques ne suffisent pas pour analyser automatiquement un texte.

On a été amené à introduire des méthodes statistiques - Un approfondissement de ces méthodes a permis à l'aide d'échantillons de déduire des propriétés linguistiques, des structures morphosyntaxiques d'un texte et d'éliminer la plupart des ambiguïtés, après les analyses morphologiques et syntaxiques. Ces méthodes sont quantitatives et permettent de prendre des décisions lors de chaque ambiguïté réduisant ainsi au minimum le nombre de ces ambiguïtés.

Nous allons étudier ces différentes utilisations des statistiques et probabilités dans le traitement automatique des langues.

Nous verrons d'abord les simples comptages, les calculs de fréquences, les calculs de corrélation, puis les grammaires probabilistes, les filtres statistiques, enfin une approche de corrections d'erreurs syntaxiques.

1. PREMIERE APPROCHE.

Comme nous venons de le dire les premiers outils statistiques utilisés ont été les fréquences les coefficients de corrélation. L'inconvénient de ces méthodes est qu'elles ne peuvent être appliquées que sur un univers assez restreint. Dès qu'un texte est long ces calculs sont alors beaucoup trop longs et donc ils ne sont plus intéressants.

Nous verrons que ces méthodes sont utilisées pour construire des théories plus élaborées telles que les filtres statistiques. Elles sont cependant beaucoup exploitées en indexation automatique et en recherche d'évaluation de pertinence [3]

Nous donnerons quelques exemples de SALTON sur la recherche d'évaluation de pertinence.

On pose $P(x_i/w_1)$ = probabilité qu'un terme x_i apparaisse dans un enregistrement donné et que cet enregistrement soit pertinent

$P(x_i/w_2)$ = probabilité qu'un terme x_i apparaisse dans un enregistrement donné et que cet enregistrement ne soit pas pertinent.

En utilisant la formule de Bayes la fonction de recherche

$$P(w_i|x) = \frac{P(x|w_i) P(w_i)}{\sum_{i=1}^2 P(x|w_i) P(w_i)}$$

ou $P(w_i)$ est la probabilité a priori de pertinence

x est un vecteur de un ou plusieurs termes

En utilisant une approximation la fonction discriminante

$$g(x) = \frac{P(w_1|x)}{P(w_2|x)}$$

Le calcul pratique de $P(x|w_i)$ se fait en calculant la fréquence de x_i ou en utilisant un modèle de distribution de poisson -

Considérant le cas d'indépendance le calcul des fréquences se fait sur un échantillon. ($x_i=1$ si x_i est un terme de x)

$$P(x|w_i) = \prod_{i=1}^n P(x_i=1|w_1)^{x_i} (1 - P(x_i=1|w_1))^{1-x_i}$$

avec $P(x_i=1|w_1) = \frac{x_i}{R}$ $P(x_i=1|w_2) = \frac{n_i - x_i}{N - R}$

	w_1	w_2	
$x_i=1$	x_i	$n_i - x_i$	n_i
$x_i=0$	$R - x_i$	$N - n_i - R + x_i$	$N - n_i$
	R	$N - R$	N

table d'occurrence d'un terme x_i
et d'une question Q .

1. GRAMMAIRES PROBABILISTES.

La théorie probabiliste a ensuite été utilisée à un niveau supérieur = au niveau des grammaires choisies lors de l'analyse syntaxique.

Rappel sur les grammaires [4]

Une grammaire formelle G est définie par le quadruplet

$$G = (V_T, V_N, R, A)$$

où V_T est l'ensemble des symboles terminaux

V_N est l'ensemble des symboles non terminaux

R est l'ensemble des règles

A est le symbole de départ.

Exemple $G = (V_T, V_N, R, S)$

$$V_N = \{S, B, C\}$$

$$V_T = \{a, b, c\}$$

$$R \left\{ \begin{array}{l} S \rightarrow aSBC \\ S \rightarrow a b C \\ CB \rightarrow BC \\ bB \rightarrow bb \\ bC \rightarrow bc \\ cC \rightarrow cc \end{array} \right.$$

On définira une grammaire probabiliste à l'aide d'une grammaire formelle. Il existera en plus, des probabilités associées à chaque règle de la grammaire. Une telle grammaire sera définie par un quintuplet [5]

Exemple :

$$G = (V_T, V_N, R, S, P)$$

P est l'ensemble des probabilités,
les autres symboles sont les mêmes que pour une
grammaire formelle.

$$V_N = \{ \sigma, A_2 \}$$

$$V_T = \{ a, b \}$$

$$R = \begin{cases} \sigma \rightarrow a A_2 A_2 & \textcircled{1} & P_{11} \\ \sigma \rightarrow b & \textcircled{2} & P_{12} \\ A_2 \rightarrow A_2 \sigma & \textcircled{3} & P_{21} \\ A_2 \rightarrow a & \textcircled{4} & P_{22} \end{cases}$$

$$P_{11} = P(a A_2 A_2 | \sigma)$$

$$P_{12} = P(b | \sigma)$$

$$P_{21} = P(A_2 \sigma | A_2)$$

$$P_{22} = P(a | A_2)$$

Si on a la chaîne suivante $a b a = v$
 v est obtenu

$$\sigma \xrightarrow{\textcircled{1}} a A_2 A_2 \xrightarrow{\textcircled{3}} a A_2 \sigma A_2 \xrightarrow{\textcircled{4}} a a \sigma A_2 \xrightarrow{\textcircled{2}} a a b A_2 \xrightarrow{\textcircled{4}} v$$

$p(v)$ est obtenu

$$\begin{aligned} p(v) &= P(a A_2 A_2 | \sigma) P(A_2 \sigma | A_2) P(a | A_2) P(b | \sigma) P(a | A_2) \\ &= P_{11} P_{21} P_{22} P_{1,2} P_{22} \end{aligned}$$

Les probabilités associées à chaque règle d'une grammaire
sont calculées en comptant le nombre moyen de
fois que chaque règle est utilisée dans le
langage généré par cette grammaire. Ce n'est pas
toujours facile l'univers est trop grand. Aussi on travaille
souvent sur un échantillon représentatif de ce langage.

Notons l'article de ПАРЯНСКИ qui traite un aspect un peu particulier des grammaires probabilistes. En effet il traite les grammaires déterministes avec un langage fini. Après quelques propriétés sur ces grammaires particulières, il propose l'utilisation des processus de Markov.

Ces processus permettent de calculer la probabilité d'une chaîne de caractères. Pour cela on doit d'abord calculer la matrice de transition du processus.

L'utilisation la plus intéressante qu'il fait de ces grammaires est la construction d'une méthode heuristique pour trouver une grammaire qui s'adaptera le mieux à un échantillon source.

Les outils statistiques nécessaires sont le test du χ^2 . Ce test permettra de prendre la décision à propos de la grammaire proposée.

En effet on pose

$$\chi^2 = \sum_{i=1}^n (f_i - \Pi p(\pi_i))^2 / \Pi p(\pi_i)$$

avec S l'échantillon tel que $S = \{(\pi_1, f_1) \dots (\pi_n, f_n)\}$

tel que $\sum_{i=1}^n f_i = N$. $\pi_i \in L(G)$ f_i sa fréquence

et tel que la probabilité d'occurrence d'un échantillon S donné est.

$$P(S) = \prod_{i=1}^n p(x_i)^{f_i}$$

ou $p(x_i)$ est la probabilité de la chaîne x_i dans le langage $L(G)$.

Pour le test χ^2 on est obligé de définir un seuil α .

Ainsi la règle de décision sera

- Si $\chi^2 \geq \chi^2_\alpha$ on rejette l'hypothèse que G est une grammaire correcte
- Si $\chi^2 < \chi^2_\alpha$ on accepte l'hypothèse que G est une grammaire correcte.

Il est important de calculer les probabilités des règles de la grammaire G .

Ce calcul que ne nous ne développerons pas ici est en fait le calcul de l'estimation par la méthode du maximum de vraisemblance.

On notera seulement le résultat

$$p_i = \frac{n_i}{n(A_i)}$$

avec n_i : nombre de fois que la règle r_i avec la prémiss

A_i apparaît dans les chaînes de S .

Ainsi l'auteur nous propose l'algorithme suivant. Nous le donnerons sous forme très schématique pour définir plus clairement l'idée principale. Puis nous le détaillerons de façon plus mathématique pour montrer à quels niveaux se trouveront les tests du χ_2

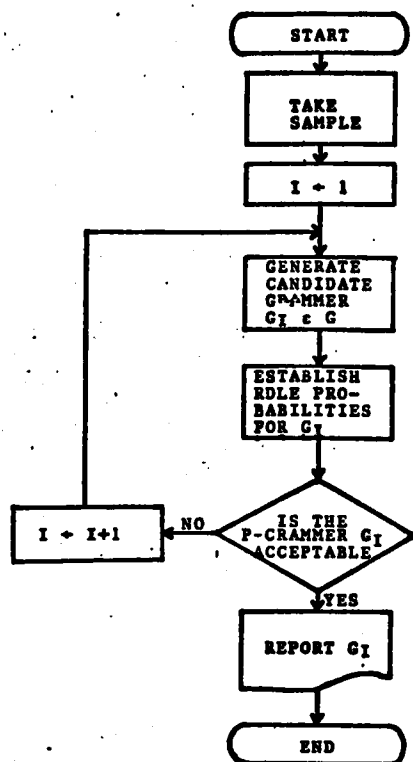


Fig. 1. General inference algorithm.

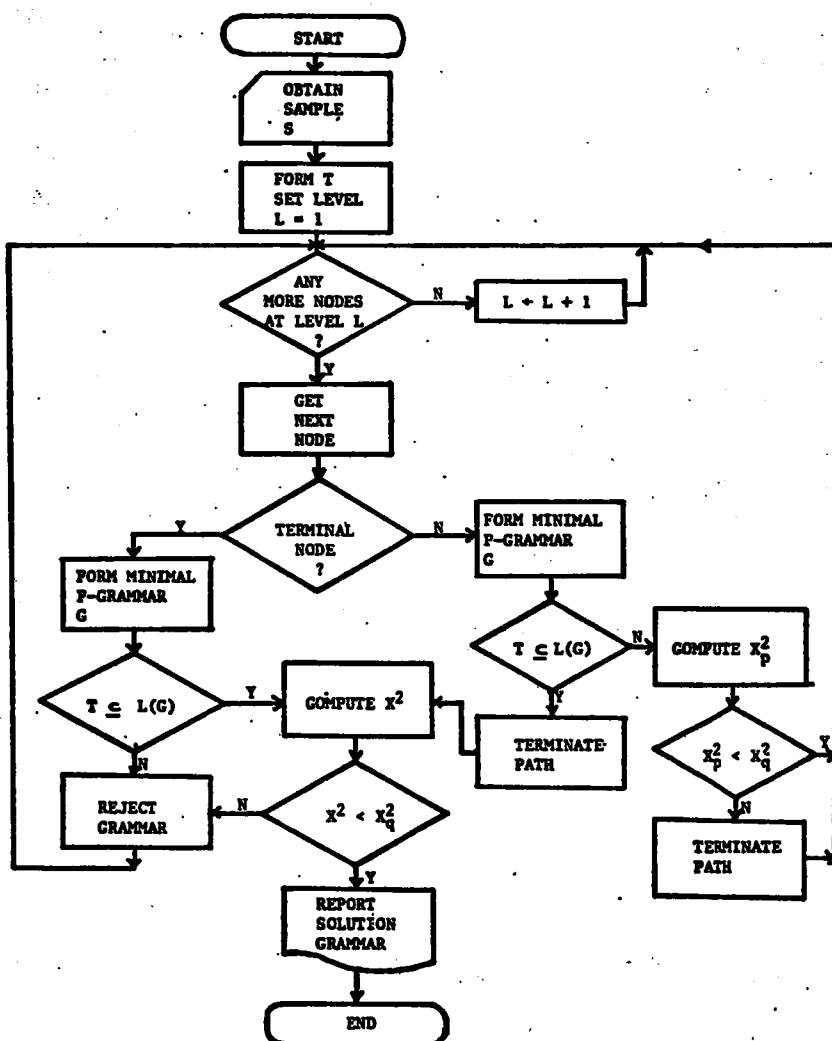


Fig. 3. Detailed inference algorithm.

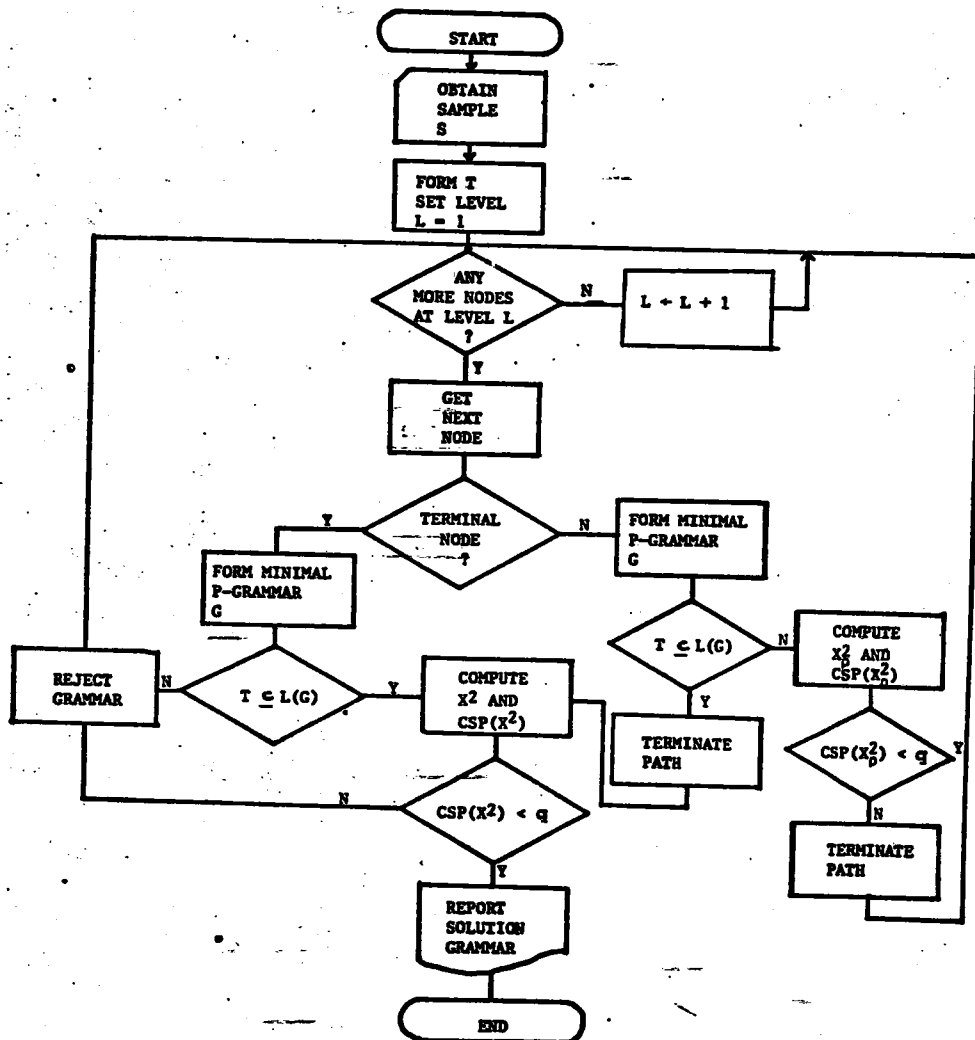


Fig. 4. Heuristic inference algorithm.

2. FILTRES STATISTIQUES

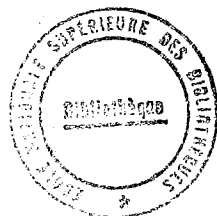
les grammaires probabilistes apportent des améliorations à l'analyse syntaxique. Les probabilités associées aux règles de grammaire permet une meilleure adaptation à un texte donné.

On a introduit en plus des grammaires probabilistes des filtres statistiques. Ce filtre permet, à la sortie de l'analyse morphologique, de supprimer la plupart des ambiguïtés [2].

Le principe du filtre est construit sur les cooccurrences des catégories morphologiques.

Ainsi on prend un échantillon de phrases tirées de façon aléatoire. Cet échantillon est analysé manuellement pour qu'il n'y ait plus d'ambiguïté. On calcule alors la matrice de effectifs de tous les couples de catégories morphologiques. On obtient la matrice des fréquences qui permettra d'étudier la dépendance de deux catégories morphologiques. Ceci se fait à l'aide d'un test en χ^2 et de la table des contingences.

On définira aussi un test pour la table des effectifs pour évaluer la représentativité de l'échantillon.



Tous ces calculs statistiques établis sur l'échantillon forment la base du filtre -

Soit une phrase

$M_1(C_1) M_2(C_2) \dots M_n(C_n)$

M_i est un mot

C_i une catégorie morphologique

Après l'analyse morphologique on aura par exemple

$A(a_1 \dots a_n) B(b_1 \dots b_n) C(c_1 \dots c_p) D(d_1 \dots d_q) E(e_1 \dots e_r)$

On examine AB .

1. Si $n=1$ $m=1 \rightarrow$ pas d'ambiguïté
2. Si $n \neq 1$ $m \neq 1 \rightarrow$ on utilise les résultats statistiques.

Après décision

3. Si $n=1$ $m=1 \rightarrow$ plus d'ambiguïté
4. Si $n \neq 1$ $m \neq 1 \rightarrow$ on ne peut rien dire sur A
5. Si $\left. \begin{array}{l} n=1 \text{ et } m \neq 1 \\ \text{ou} \\ n \neq 1 \text{ et } m=1 \end{array} \right\} \rightarrow$ on observe BC comme en 2

Après décision

- Si $m=1 \rightarrow$ on reexamine $A(a_1 \dots a_n) B(b_1)$
ou passe à CD comme en 1
- Si $m \neq 1 \rightarrow$ on ne décide rien sur A et B
on passe à CD comme en 1

Cet algorithme a l'avantage de ne pas être de taille très grande et minimise ainsi le temps.

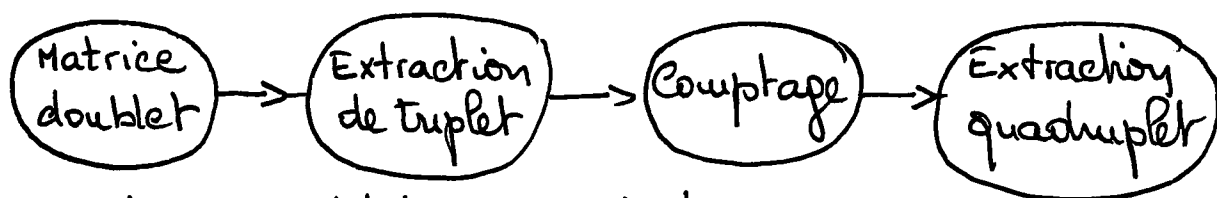
les résultats expérimentaux montrent que sur un échantillon de 1491 mots sur un texte de 7000 mots représentant environ 40% de mots comportant au moins une ambiguïté soit 3000 mots ambigus.

On obtient après le filtre
250 mots ambigus soit moins de 10%
60 mots comportant une erreur soit moins de 2%.

L'analyse syntaxique se fait par un système de recherche de suites de catégories morphologiques autres que les couples.

On propose ainsi de traiter les triplets (quadruplets, quintuplets) consécutifs et significatifs.

On calcule les fréquences relatives correspondantes sur l'échantillon. Et on choisit à l'aide du test du χ^2 les triplets significatifs et ainsi de suite.



exemple doublet (DN, N) (N, AN)

triplet (DN, N, AN)

La construction du noyau grammatical se fait de façon dynamique - la grammaire évolue suivant des classes du corpus -

Mouneur HYA HYA utilise aussi la théorie du filtre statistique pour l'analyse morphologique [9] [10]

En effet on essaie d'obtenir automatiquement à l'aide de deux filtres statistiques les chaînes initiales et terminales pertinentes en vue de l'analyse.

On veut obtenir un compromis entre la taille minimum des dictionnaires et le maximum de bruit morphologique.

Exemple la chaîne terminale TIONS
TIONS chaîne terminale de (verbe, substantif)
NOUS CHANTIONS
(Pronom) (Verbe substantif)

↓
bruit morphologique

La méthode consiste à effectuer un traitement statistique sur le texte d'entrée pour obtenir les différentes chaînes et informations grammaticales et statistiques associées.

Ces mots non vides généreront des chaînes terminales et initiales et ces chaînes seront candidates pour être sélectionnées chaînes pertinentes.

Plusieurs critères sont nécessaires pour établir ces deux filtres statistiques.

D'abord un critère de non redondance - Une chaîne est non redondante par rapport à une autre chaîne si elle n'apporte pas d'information supplémentaire.

On définit ensuite les chaînes primaires qui seront des chaînes suffisamment longues et représentatives - On aura première espèce et deuxième espèce.

Enfin on décrit un critère de pertinence

- pertinence morphologique
- pertinence grammaticale.

Le critère de pertinence morphologique minimise les chaînes représentatives

Le critère de pertinence grammaticale fixe le bruit morphologique tolérable.

Le premier filtre permet alors de sélectionner sur les chaînes primaires les chaînes qui vérifient

- le critère de pertinence grammaticale
- le critère de pertinence morphologique
- le critère de non redondance

Le deuxième filtre après un tri sur les chaînes du premier filtre traite en fait les chaînes primaires dites de deuxième espèce.

On peut après ces deux filtrages passer à

l'analyse morphologique - on a généré le système avec la consultation de différents lexiques - La consultation de ces lexiques se fait par ordre de priorité -

Le taux de résolution est de l'ordre de 56%
Les résultats expérimentaux montrent que la différence des taux de résolution entre l'apprentissage et les textes nouveaux diminue lorsque la longueur des textes augmente (ordre de grandeur 6 à 3%).

Cette différence est due essentiellement au mot inconnu employé exceptionnellement.

D. CORRECTIONS D'ERREURS SYNTAXIQUES

Les ambiguïtés ne sont pas les seuls problèmes de l'analyse syntaxique - les erreurs en sont d'autres tout aussi importantes - les méthodes de corrections d'erreurs syntaxiques ont été largement développées avec la compilation principalement [11]

Mousour Fu introduit pour les erreurs de l'analyse syntaxique un modèle statistique en vue de la correction de ces erreurs -

Il distingue 3 types d'erreurs -

- substitution T_S
- suppression T_D
- insertion T_I

Soit L un langage et Σ un alphabet fini, une chaîne de L peut être déformée en une chaîne différente n'appartenant pas à ce langage -

$$\begin{array}{ll} xby \in T_S(xay) & b \neq a \\ xay \in T_I(xy) & a \text{ dans } \Sigma \\ xy \in T_D(xay) & a \text{ dans } \Sigma \end{array}$$

On peut alors associer à chaque cas une probabilité, définissant ainsi le modèle stochastique.

la probabilité de substituer a par b est $q_s(b|a)$
 la probabilité d'insérer b avant a est $q_I(b|a)$
 la probabilité de supprimer a dans une chaîne est $q_D(a)$
 la probabilité d'insérer a à la fin d'une chaîne est $q_I(a)$

On calculera alors la probabilité qu'un symbole a soit déformé en x

$$\begin{aligned}
 q(x|a) &= q_D(a) \quad \text{si } x = \lambda \\
 &= \max(q(a), q_I(a|a)q_D(a)) \quad \text{si } x = a \\
 &= \max(q_s(b|a), q_I(b|a)q_D(a)) \quad \text{si } x = b \neq a \\
 &= q_I(b_1|a) \dots q_I(b_{l-1}|a) \max\{q_s(b_l|a), q_I(b_l|a)q_D(a)\} \\
 &\quad \text{si } x = b_1 \dots b_l \quad b_l \neq a \quad l > 1 \\
 &= q_I(b_1|a) \dots q_I(b_{l-1}|a) \max\{q(a), q_I(b_l|a)q_D(a)\} \\
 &\quad \text{si } x = b_1 \dots b_l \quad b_l = a \quad l > 1
 \end{aligned}$$

la méthode utilisée pour essayer de corriger ces erreurs est un algorithme utilisant la théorie du maximum de vraisemblance.

Elle est définie ainsi y n'appartenant pas au langage $L(G)$ et x appartenant à $L(G)$

$$q(y|x) p(x) = \max_z \{q(y|z) p(z) \mid z \in L(G)\}.$$

où $p(z)$ est la probabilité de générer z avec G grammaire probabiliste.

Une telle méthode est utilisée sur des chaînes de caractères - ce n'est pas en relation directe avec

Ce qui nous interesse. Il est cependant important de noter la méthode qui par certain aspect pourrait être exploitée dans de nouvelles recherches sur l'analyse syntaxique.

V CONCLUSION

Le traitement automatique des langues est utilisé dans différents projets: la traduction automatique, l'analyse de contenu, la recherche de l'information.

Des problèmes théoriques du traitement automatique ont amené les chercheurs à utiliser les théories statistiques pour palier ses difficultés. L'analyse morphologique et l'analyse syntaxique se trouvent améliorées par des filtres statistiques, des grammaires probabilistes et la construction d'algorithmes. Il est à remarquer que le temps de traitement des algorithmes est de plus en plus court.

L'évolution la plus intéressante est la génération de noyaux de grammaire qui peuvent ne pas rester statiques.

Cependant il est un inconvénient majeur, qu'on a retrouvé tout le long de l'exposé. Le traitement automatique des langues est très difficile à réaliser lorsqu'on travaille sur des textes très longs.

- BIBLIOGRAPHIE -

- 1 BOURGUIGNON C., HOLLARD S., ROVAULT J. ; SIRAUCO UN SYSTEME D'AIDE A L'ETUDE MORPHO-SYNTAXIQUE DE CORPUS ; DOC UER INFORMATIQUE ET MATHEMATIQUE EN SCIENCES SOCIALES ; DOC N°5 SERIE C ; UNIVERSITE SCIENCES SOCIALES GRENOBLE .
- 2 CARRETERO J. ; MESURES DE PROPRIETES STATISTIQUES DE TEXTE EN LANGUE NATURELLE POUR LA DEFINITION DE MODELES LINGUISTIQUES UTILISABLES EN ANALYSE MORPHO-SYNTAXIQUE ; THESE DOCTORALES SCIENCES GRENOBLE ; OCTOBRE 77 .
- 3 BALTON G. ; MATHEMATICS AND INFORMATION RETRIEVAL ; J. DOCUMENT. ; 1979 ; V 35 ; 1-29 ; 71 BIBL .
- 4 TAYLOR L. BOOTH , RICHARD A. THOMPSON ; APPLYING PROBABILITY MEASURES TO ABSTRACT LANGUAGES ; IEEE TRANS. COMP. ; 1973 ; V 22 ; N 5 ; 442 - 450 .
- 5 DELATTE L., DUCHESNE-DEGEY M. , GOVAERTS S., DENOOZ J. . LE TRAITEMENT AUTOMATIQUE DE LA LANGUE FRANCAISE AU LABORATOIRE D'ANALYSE STATISTIQUE DES LANGUES ANCIENNES ; ORGANIS. INTERN. ET. LANGUES ANC. ORDINATEUR ; 1977 ; N 4 ; 1-55 ; 6 BIB
- 6 FU K. S. ; SOME APPLICATIONS OF STOCHASTIC LANGUAGES ; IN. APPL. STAT. SYMP. PROC. , DAYTON OHIO ; 1976 ; 183-200 ; 1P BIBL .
- 7 MARYANSKI F. T. , BOOTH T. L. ; INFERENCE OF FINITE-STATE PROBABILISTIC GRAMMARS ; IEEE. TRANS. COMP. ; 1977 ; V 26 ; N 3 ; 521-536 ; 12 BIBL .
- 8 VAN DER MUDE A. , WALKER A. ; ON THE INFERENCE OF STOCHASTIC REGULAR GRAMMARS ; INFORM. AND CONTROL ; 1978 ; V 38 ; N 3 ; 310-329 ; 19 BIBL .

9. HLAL YAHYA ; ANALYSEUR MORPHOSYNTAXIQUE ET SYNTAXIQUE
DE LANGUES NATURELLES OBTENU PAR APPRENTISSAGE ; NOTE CEA
N 2053/1978, 2 FASC. ; 64 + 47 ; 106 BIBL.

10. HLAL YAHYA ; METHODES D'APPRENTISSAGE POUR L'ANALYSE
MORPHOSYNTAXIQUE ; THESE SCI. PHYS. INTELLIGENCE ARTIFICIELLE
LINGUIST. AUTOM. PARIS, 1979, (8) - 285 ; 111 BIBL.

11. SHIN-YEE LU, KING-SUNFU. STOCHASTIC ERROR-CORRECTING
SYNTAX ANALYSIS FOR RECOGNITION OF NOISY PATTERNS ; IEEE.
TRANS. COMPUTERS.) 1977, V 26, N 12, 1268-1276, 17 BIBL.

