

La construction de lexiques de formes fléchies et l'analyse morphologique du coréen

Université Paris 7
Laboratoire d'Automatique Documentaire et Linguistique
Institut Gaspard Monge

Thèse de Doctorat en Informatique

ChangYeol LEE

Directeur de thèse : Maurice GROSS

Membres du Jury : Dominique PERRIN
Eric LAPORTE
Gaston GROSS
Maurice GROSS
Maxime CROCHEMORE

Juin 1997

La construction de lexiques de formes fléchies et l'analyse morphologique du coréen

Université Paris 7
Laboratoire d'Automatique Documentaire et Linguistique
Institut Gaspard Monge

Thèse de Doctorat en Informatique

ChangYeol LEE

Directeur de thèse : Maurice GROSS

Membres du Jury : Dominique PERRIN
Eric LAPORTE
Gaston GROSS
Maurice GROSS
Maxime CROCHEMORE

Juin 1997

La construction de lexiques de formes fléchies et l'analyse morphologique du coréen

Résumé

Ce travail a pour objectif de développer un système d'analyse morphologique du coréen. La particularité de cette étude est qu'il repose sur un dictionnaire du coréen construit systématiquement pour cette application informatique. Les dictionnaires traditionnels, qu'ils soient destinés à un usage humain ou à une consultation informatique, n'ont pas de principes formels de représentation. Le système du dictionnaire utilisé dans notre étude est par contre défini par des principes généraux. Par exemple, il ne contient pas plusieurs fois les mêmes entrées, c'est un des critères formels de recensement linguistique que nous avons choisis.

Le mot "simple" est la plus petite unité dans nos dictionnaires. Dans les phrases coréennes, il est fréquent qu'une séquence délimitée par des séparateurs (comme des blancs, des points, ...) soit composée de plusieurs mots "simples". Ces mots simples peuvent être analysés et construits comme un jeu de LEGO et son démontage est facile.

Il n'y a pas de langues naturelles qui n'ait pas d'ambiguïté. Les méthodes de résolution d'ambiguïté utilisées actuellement sont basées sur un calcul statistique ou sur des règles. C'est parce que la connaissance qu'on a actuellement des langues naturelles ne garantit pas la complétude que l'on a dans un monde clos. Autrement dit, la connaissance linguistique ne peut pas être accumulée automatiquement, et donc elle doit être faite par l'être humain. C'est un des problèmes du monde ouvert.

Dans notre étude donc, la connaissance linguistique est accumulée seulement par les experts humains, non par une méthode statistique ou heuristique. Cela veut dire que la méthode linguistique est plus assurée, même si elle demande beaucoup de temps. L'information de désambiguïsation construite par les experts humains est utilisée dans notre système d'Analyseur morphologique. Pour la construction de l'information, il est inévitable d'analyser les structures internes ambiguës qu'on observe dans les phrases réelles.

Actuellement, on a construit des dictionnaires de Noms simples, de Postpositions nominales, de Préfixes, de Suffixes, de Verbes, d'Adjectifs et de Postpositions prédicatives. Sur la base de ces dictionnaires, on a dérivé deux dictionnaires des formes fléchies: le dictionnaire des Formes Nominalisées et celui des Formes fléchies prédicatives, aussi que des règles pour les Noms Dérivés. Les dictionnaires et l'information de désambiguïsation sont utilisés dans le processeur d'Analyseur morphologique.

Mots Clés:

Dictionnaire de mots simples, Dictionnaire de formes fléchies, Syllabes finales du radical, Postpositions prédicatives, Formes Nominalisées, Formes fléchies prédicatives, Lexique-grammaire, Consultation du dictionnaire, Analyseur morphologique.

한국어 활용어 사전 구축과 형태소 분석

요약

이 논문은 한국어 문장 분석에 필요한 사전 구축과 형태소 분석에 관한 연구 논문이다. 본 논문의 가장 큰 특징은 체계적으로 구성된 사전에 있다. 전통적 사전이 인간을 위한 것이든 기계를 위한 것이든 그 내용이 임의성이 많은데 비하여 여기의 사전은 단순어에 기반을 가지고 체계적이며, 중복성이 없고, 항목 선택의 명확한 기준이 있으며, 또한 구성적으로 구성되었다.

단순어란 의미의 최소 단위로서 통사적 독립성을 유지하고 있다. 한국어에 있어서 쉼표나 빈칸등에 의해서 만들어지는 글자들은 여러개의 이들 단순어로 구성되어있다. 이 단순어들은 레고를 가지고 임의의 인조물을 만드는 것과 같이 문장의 합성과 분해에서 어휘의 확장성을 가지고 있으며 어휘의 특징을 추론할 수 있게 한다. 그들은 특별한 변화없이 문장에서 활용어 형태와 일치된다. 그러므로 활용어 사전의 구축은 형태소 분석기의 대부분을 차지하게 된다.

문장을 분해(해석)하는 데 있어서 모호성이 없는 시스템은 물론 없다. 모호성에 대하여 많은 시스템은 통계나 규칙을 가지고 문제에 대하여 접근하고 있다. 그것은 닫힌 체계에서의 지식의 완전성과 건전성 (Completeness and Soundness)을 자연어에서는 보장하지 않기 때문에 생기는 문제이다. 다시말하면 언어적 지식은 자동적으로 축적될 수 없으며 그러므로 인간에 의해서 행해져야 한다.

우리의 연구에 있어서, 언어적 지식은 통계나 휴리스틱한 방법이 아닌 오직 언어적 전문가에 의해서 축적을 한다. 다시말하면 이런 언어적 방법이 비록 많은 시간이 걸리지만 가장 정확성을 보장해준다. 언어적 방법에 의해서 구축된 모호성 해결 정보는 우리의 형태소 분석기에서 사용될 것이고, 그 정보의 구축을 위하여 어떤 모호성 구조가 있는가를 분석하는 것이 물론 전체 조건임은 언급할 필요가 없다.

현재 구축한 단순어 사전으로 단순 명사, 명사 구성 접두사, 명사 구성 접미사, 명사 후치사, 동사, 동사 활용어, 형용사, 형용사 활용어 사전등이 제공되고 있다. 이 사전들을 바탕으로, 문장 분석에 필요한 두개의 활용형 사전인 명사화 유형 사전, 술어-활용어 사전, 그리고 명사를 위한 문법적 규칙을 이 논문에서 구축하였다. 이 사전 정보와 모호성 해결 정보가 본 논문의 형태소 해석기인 품사 표시기(EPL)에서 활용되고 있다.

중요 단어

단순어 사전, 활용어 사전, 어간 끝 음절, 어미, 명사화 유형, 술어와 활용어, 사전 문법, 사전 검색, 형태소 해석기

The construction of lexicons of inflectional forms and morphological analysis in Korean

Abstract

This paper contains a study about the Korean sentence analysis. The characteristics of this study is that it is based on a well-formed dictionary system. Traditional dictionaries, either for human use or in a machine readable form, do not have any clear principle in their construction, but the dictionary system used in this study does not have any duplicate entries between sub-dictionaries; it has clear principles in the selection of entries, and has been systematically constructed with linguistic knowledge.

A simple word is the smallest unit in our dictionaries. A string made up by separators can be composed of several simple words in Korean sentences. These simple words permit to be analyzed and constructed like LEGO system which can freely be constructed and dis-assembled to any artifacts. The constructed words are matched with the inflectional forms in sentence.

There is no system which has no ambiguity. Several systems are based on statistics or heuristics methods to the ambiguity problems. It is caused by the facts that linguistic knowledge doesn't guarantee the completeness and soundness of the closed world assumption. In other words, linguistic knowledge can not be automatically accumulated, and therefore it must be constructed by human. It is an open world problem.

In this paper, the linguistic knowledge is accumulated only by human experts, not by statistics and heuristics methods. It means that the linguistic method is the best way we assure, although it consumes much time. Disambiguation information constructed by humans is used in our morphological analyser. For the construction of this information, it is prerequisite to analyze the ambiguity structures in real sentences.

Currently, we have several simple dictionaries such as Simple Noun, Nominal-Postposition, Prefix, Suffix, Verb, Adjective, and Inflectional Suffix. Upon these dictionaries, we constructed two inflectional dictionaries, such as Nominalized Form dictionary and Predicate containing Inflectional Suffixes dictionary, and one grammatical rule for Noun phrases. The dictionaries and the disambiguation information will be used in our lexical language processor.

Keywords

Simple Dictionary, Dictionary of Inflectional Forms, Final Syllables of Root, Inflectional Suffix, Nominalized Form, Predicate with Inflectional Suffix, Lexicon-Grammar, Dictionary Lookup, Morphological Analyser.

TABLE DES MATIERES

Principes de notations

1.	Introduction	1
1.1.	Modèle conceptuel	4
1.2.	Organisation du travail	7
2.	Structure du dictionnaire établi pour ce travail	9
2.1.	Structure de type 1	10
2.1.1.	Nom Simple	13
2.1.2.	Préfixe	16
2.1.3.	Suffixe	19
2.1.4.	Postposition du Nom	23
2.2.	Structure de type 2	25
2.2.1.	Verbes Nominalisés	26
2.2.2.	Adjectifs Nominalisés	28
2.2.3.	Formes Spéciales	30
2.2.4.	Procédures de traitement des ambiguïtés	33
2.3.	Structure de type 3	35
2.3.1.	Noms Prédicatifs	37
2.3.2.	Verbes	40
2.3.3.	Adjectifs	41
2.3.4.	Syllabe Finale du Radical	42
2.3.5.	Postpositions Prédicatives	46
3.	Organisation du dictionnaire	49
3.1.	Plan de la spécification de l'Automate	51
3.1.1.	Définition des fonctions	51
3.1.2.	Exemples de procédures	53
3.1.3.	Algorithme pour l'automate fini	56
3.2.	Consultation Basée sur la Position(CBP)	59
3.2.1.	Procédures de la construction CBP	61
3.2.2.	Evaluation de la performance	65
4.	Etiquetage des parties du discours basé sur le lexique	66
4.1.	Analyseur Lexical	67
4.1.1.	Exemples de test	68
4.2.	Analyseur Morphologique	71
4.2.1.	Exemples de test	75
5.	Conclusion	78
	BIBLIOGRAPHIE	80
	ANNEXE A. B-1. B-2. C-1. C-2. D. E. F.	82

Principes de notation

Notations

Coréen	Français	Exemples
명사	Nom(N)	국세청장, 나무
단순명사	Nom Simple(NS)	음악, 나무
명사 접두어	Préfixe(PF)	국 de (국세청장)
명사 접미어	Suffixe(SF)	청, 장 de (국세청장)
파생명사	Noms Dérivés(ND)	신세대
동사	Verbe(V)	가다
형용사	Adjectif(A)	아름답다
서술명사	Nom Prédicatif (NP)	공부 de (공부하다)
명사 어미	Postpositions nominales (PPN)	은, 는, 이, 가, 께서도
명사화 유형	Forme Nominalisée(FN)	잠, 자기, 자는것,
명사화 동사	Verbes Nominalisés(VN)	잠
명사화 형용사	Adjectif Nominalisé(AN)	아름다움
술어와 활용어	Formes fléchies prédicatives(FFP)	자-시거나
어절	Séquences de mots	학교에
음절	Syllabes	학, 교, 에
단어	Mots	학교, 에
기본형	Formes de Base	기쁘다
어간	Radicaux	기쁘
활용어미	Postpositions prédicatives (PPP)	느냐, 구나, 으므로
어간끝 음절	Syllabes Finales du Radical (SFR)	쁘 de (기쁘다)
초성	Premières Consonnes(PC)	ㅎ de (학)
중성	Voyelles(VO)	ㅏ de (학)
종성	Consonnes Finales(CF)	ㄱ de (학)

Symboles

$$X^+ \quad (= \quad X^n \quad (n \geq 1))$$

$$X^* \quad (= \quad X^n \quad (n \geq 0))$$

$$X^? \quad (= \quad X^n \quad (n = 0) \text{ ou } (n = 1))$$

$$\prod_{i=1}^n X_i \quad (= \quad X_1 * X_2 * X_3 * \dots X_{n-1} * X_n)$$

1. Introduction

Nous développons ici un système d'analyse automatique de textes écrits en coréen. A l'heure actuelle, de nombreux documents sont disponibles sous forme électronique que l'on peut consulter directement par ordinateur. Ainsi, pour obtenir des informations contenues dans des textes libres, il sera indispensable d'analyser la langue naturelle dans un système électronique.

Toutes les activités en traitement automatique des langues naturelles, notamment l'analyse automatique du texte, la documentation automatique et la traduction ou la génération automatique de textes nécessitent des dictionnaires électroniques[33]. Dans les applications à base d'analyse de textes, si une unité segmentée n'est pas trouvée dans le dictionnaire, l'analyse plus approfondie du texte qui devrait suivre la première phase sera bloquée. Il est donc indispensable de disposer de bases de données ainsi complètes que possible.

Dans le cas du français, la segmentation d'un texte s'effectue en unités graphiquement séparées par les blancs (et/ou aussi par les tirets, apostrophe ou point). Une chaîne de caractères comprise entre deux séparateurs sera alors considérée comme mot simple, plus précisément comme forme fléchie de mot simple canonique. Ainsi, le dictionnaire électronique des mots simples du LADL (DELAS) [4] contient le codage méthodique des entrées qui permet d'engendrer automatiquement les formes fléchies du dictionnaire électronique des mots fléchis du LADL (DELAF) [5], c'est-à-dire les formes rencontrées dans les textes.

Pour le coréen, la situation est différente. Les unités séparées par les blancs et/ou par le point (les tirets et apostrophe n'existant pas en coréen) ne sont pas toujours des formes fléchies d'unités lexicalement simples, elles peuvent être constituées de séquences composées, éventuellement des syntagmes, des types :

- Nom - Postposition nominale (PPN)
- (Verbe+Adjectif) - Postposition prédictif (PPP) (suffixes flexionnels)

Les postpositions, comparables aux prépositions en français, se trouvent obligatoirement soudées à un substantif pour indiquer le rôle grammatical du substantif associé ; les postpositions prédictives sont soudées aux racines des verbes et des adjectifs (notons que les adjectifs coréens se conjuguent, comme les verbes, à l'aide de postpositions prédictives, sans prendre de copule).

Le fait que certains morphèmes se trouvent attachés à des éléments lexicaux complique l'analyse morphologique. Surtout quand il existe plusieurs possibilités de segmenter des mots qui forment une seule séquence, les problèmes d'ambiguïtés deviennent alors plus complexes. Supposons qu'on ait une séquence de type :

abcd

Nous pouvons la segmenter des 8 manières suivantes :

- | | |
|--------|---------------|
| (i) | a / b / c / d |
| (ii) | ab / c / d |
| (iii) | ab / cd |
| (iv) | a / b / cd |
| (v) | a / bc / d |
| (vi) | abc / d |
| (vii) | a / bcd |
| (viii) | abcd |

Pour obtenir la(les) segmentation(s) correcte(s), un dictionnaire aussi complet que possible doit être fourni : s'il n'existe pas de mots de type "a" ni de type "d" dans le dictionnaire, les possibilités (i), (ii), (iv), (v), (vi), et (vii) seront rejetées, et donc on n'aura que deux segmentations possibles (iii) et (viii). Si on observe ces deux cas, c'est-à-dire si on trouve les mots "ab", "cd" et "abcd" dans le dictionnaire, le choix entre deux analyses (iii) et (viii) ne peut pas être effectué au niveau lexical ou morphologique, mais il faudra le traiter au niveau syntaxique.

Par exemple, la séquence suivante présente une ambiguïté qu'on ne peut pas lever au niveau morphologique :

정치가 *jengchiga*

car, ce peut être une des deux séquences suivantes :

a / b (a = 정치 (*jengchi*) [politique] / b = 가 (*ga*) [PPP.nmtf],
 donc **ab** = 정치-가 [politique-Sujet])
 ab (**ab** = 정치가 (*jengchiga*) [politicien])

Les phrases suivantes illustrent respectivement ces deux interprétations :

- (1) 한국 정치가 변하고 있다.
 hangug jengchi-ga byenha-go iss-da
 Corée **politique**-PPN changer-PPP être en train de-PPP
 (La **politique** coréenne est en train de changer)

- (2) 한국 정치가 협회가 결성되었다.
 hangug jengchiga hyebhoi-ga gyelsengdoi-ess-da
 Corée **politicien** association-PPN être organisé-PPP
 (Une association des **politiciens** coréens a été organisée)

Il est évident que ces deux interprétations ne peuvent être obtenues que quand le dictionnaire fournit tous les mots acceptables. Or, les dictionnaires traditionnels sont loin d'être complets de ce point de vue. D'une part, la liste des postpositions qui

s'attachent à un nom (e.g. postpositions casuelles comme *nmtf*, *Acc*, *Dat* ou *Loc*) n'est pas complète, et surtout les combinaisons de postpositions ne sont pas étudiées systématiquement. D'autre part, des noms dérivés comme celui dans (2) (i.e. 정치가 (*jengchi-ga*) [politique-Sfx:personne = politicien]) ne sont pas recensés d'une manière exhaustive dans dictionnaires classiques. Dès lors, il serait difficile d'espérer un étiquetage correcte pour les séquences de type (1) ou (2).

Dans ce travail, nous allons construire des dictionnaires aussi complets que possible, dans le but d'analyser toutes les séquences complexes, y compris les formes fléchies de termes prédicatifs (i.e. verbes et adjectifs), formes directement observables dans des textes écrits. Dans le paragraphe suivant, nous allons présenter la structure de notre système d'analyse morphologique comportant des dictionnaires de formes fléchies.

1.1. Modèle conceptuel

Le système d'analyse morphologique que nous allons implémenter dans ce système est construit de la manière suivante <figure 1-1> :

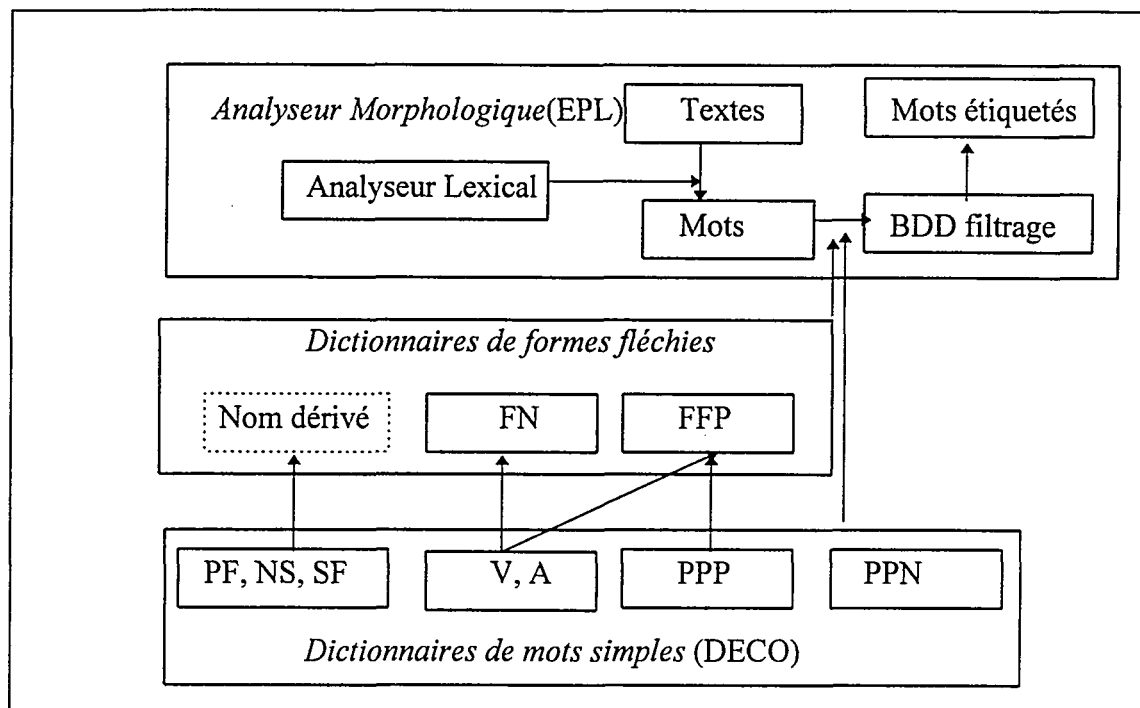


Figure 1-1. Modèle conceptuel du système de l'analyse morphologique.

Notre système est constitué de trois parties : le système lexical DECO, le dictionnaire des formes fléchies et l'analyseur morphologique. Le système DECO contient les sous-lexiques suivants qui ont été décrits et qui figurent dans [24], [25], [27], [30] :

Dictionnaire des	Noms simples	NS
Dictionnaire des	Adjectifs simples	A
Dictionnaire des	Verbes simples	V
Dictionnaire des	Préfixes et Suffixes	PF/SF
Dictionnaire des	Postpositions prédictives et Postpositions nominales	PPP/PPN

Les deux derniers dictionnaires (PF/SF, PPP/PPN) permettent d'engendrer des séquences complexes en associant avec les trois premiers : il y a d'une part, des formes dérivées basées sur des noms simples (une association de NS avec PF/SF) et d'autre part, des formes fléchies basées sur des prédictives (une association de A ou V avec PPP).

Basée sur le système lexical DECO, on va construire une grammaire et engendrer les dictionnaires suivants :

Grammaire équivalente au dictionnaire des Noms dérivés	ND
Dictionnaire des Formes Nominalisées des termes prédicatifs	FN
Dictionnaire des Formes fléchies prédicatives	FFP

Les deux derniers dictionnaires comporteront toutes les formes conjuguées des verbes et des adjectifs, y compris toutes les nominalisations basées sur les suffixes dérivationnels -음 (-eum) et -기 (-gi). La grammaire locale que nous allons construire sera équivalente au dictionnaire des noms dérivés. Autrement dit, cette grammaire permettra de reconnaître toutes les formes dérivées, et donc de distinguer ces formes, par exemple, les séquences composées de type Nom-PPN. Considérons les exemples suivants figure 1-2 :

Cas 1	Cas 2	Cas 3	Cas 4
모양가	예술가	친구가	정치가
<i>moyangga</i> <i>moyang-ga</i>	<i>yesulga</i> <i>yesul-ga</i>	<i>chinguga</i> <i>chingu-ga</i>	<i>jengchiga</i> <i>jengchi-ga</i>
*	art-PPN	ami-PPN	<ambiguë>
	artiste	ami- <i>Sujet</i>	(a) politicien (b) politique- <i>Sujet</i>

Figure 1-2. Cas divers de type “X-ga” .

Le Cas 1 représente des séquences inacceptables, bien que la partie qui précède le morphème *ga* soit un nom simple (= *appearance*). C’est parce que les substantifs dont la dernière lettre est une consonne (comme dans le *Cas 1*) ne peuvent pas être suivis d’une postposition *ga* que le mot est interdit. Mais, le morphème *ga* n’est pas non plus un suffixe compatible.

Par contre, dans le Cas 2, *ga* est un suffixe désignant une personne, et ainsi la séquence 예술-가 (*yesul-ga*) est un nom dérivé à partir d’un nom simple 예술 (*yesul*) [art], et il figure dans notre dictionnaire. Notre grammaire permettra donc de reconnaître ce mot.

De même, le morphème *ga* dans le Cas 3 ne sera reconnu que comme postposition de nominatif, puisque le nom 친구 (*chingu*) [ami] ne peut pas être suivi du suffixe *ga*, et qu’on ne trouvera pas cette séquence dans la liste des formes dérivées. Notons que nous ne construirons pas de dictionnaires des formes dérivées, mais nous comparerons la liste des suffixes et celle des postpositions, et nous n’établirons les listes de noms dérivés que quand il existe des homographes entre ces deux catégories.

Par exemple, le morphème *ga* est un des homographes typiques, nous établirons la liste des noms dérivés à base de suffixe *ga* : le nombre de ces noms est de 597 dans la version actuelle de notre dictionnaire. Cette liste est intégrée dans le module BDD

<figure 1-1> et sera utilisée pour la désambiguation. Voici un extrait des noms de type *N-ga* <figure 1-3> :

공예가 (<i>gongyega</i>)	계몽가 (<i>gyemongga</i>)
극작가 (<i>geugjagga</i>)	달변가 (<i>dalbyenga</i>)
무용가 (<i>muyongga</i>)	문예가 (<i>munyega</i>)
세력가 (<i>selyegga</i>)	수필가 (<i>supilga</i>)
연설가 (<i>yenselga</i>)	이론가 (<i>ilonga</i>)

Figure 1-3. Extrait des noms de type *N-ga*.

Par conséquent, toutes les séquences de type *X-ga* qui ne se trouvent pas dans le module BDD seront considérées automatiquement comme des séquences constituées d'un nom et d'une postposition de nominatif *ga*. La situation est en fait plus complexe, car il y a des noms dérivés et composés qui s'attachent à *ga*. Nous en parlerons plus tard.

Le Cas 4 représente des cas ambigus dûs à des homographes (i.e. le suffixe *-ga* et la postposition *-ga*), étant donné que le nom 정치 (*jengchi*) [politique] est compatible avec le suffixe *-ga*. Cette sorte d'ambiguïté ne peut pas être levée au niveau morphologique, l'information syntaxique doit être prise en compte.

Avec les dictionnaires que nous avons mentionnés au-dessus, nous allons implémenter un analyseur morphologique du coréen (Étiquetage des Parties du discours basé sur le Lexique : EPL). Quand un texte est donné, la première phase sera de le transformer en des chaînes de séquences où les séparateurs sont reconnus. Les structures de toutes les séquences ainsi obtenues seront analysées sur la base des dictionnaires de formes fléchies : tous les découpages possibles seront ainsi proposés. Le Base de Données de Désambiguï sation (BDD) filtrage sera appliquée pour lever des ambiguïtés d'étiquetage des mots.

1.2. Organisation du travail

Ce travail est constitué de trois chapitres. Dans le chapitre 2, nous allons présenter les dictionnaires que nous avons utilisés et engendrés. Comme nous l'avons mentionné ci-dessus, nous nous sommes basé sur des dictionnaires de mots simples et des lexiques d'uffixes et de postpositions. Les dictionnaires des formes fléchies sont alors engendrés à partir de ces dictionnaires et de la grammaire que nous avons construite dans ce travail.

Ensuite, nous classons toutes les séquences qu'on observe dans les textes écrits du coréen. À part les *adverbes* que nous ne traiterons pas dans ce travail, on peut diviser les séquences en deux types : séquences *nominales* et séquences *prédicatives* (*verbales* et *adjectivales*). Plus précisément, on les classera en trois catégories, comme suit :

Type 1 (T1)	Toutes les séquences comportant au moins un nom simple. On peut les formaliser comme $PF^* NS^+ SF^* PPN^?$. Ces séquences sont éventuellement constituées d'un ou plusieurs préfixe/s/ et/ou d'un ou plusieurs suffixes. Les PPN sont des postpositions qui s'attachent à des substantifs. Voici un exemple : 국세청장이 <i>gugseichengjangi</i> [PF-NS-SF-SF-PPN].
Type 2 (T2)	Toutes les séquences comportant au moins une forme nominalisée. On peut les formaliser comme $(PF^? (FN NS)^+ SF^? \text{Formes Spéciales}) PPN^?$. Ces séquences ressemblent à celles du Type 1, sauf qu'elles comportent une forme nominalisée à la place d'un nom simple. Par exemple : 늦잠자기가 <i>neujjamjagineun</i> [PF-FN-FN-PPN].
Type 3 (T3)	Toutes les séquences comportant un verbe ou un adjectif. On peut les formaliser comme $(PF^? (FN NS)^+ SF^?)^* (V_{\text{root}} A_{\text{root}}) PPP$. Ces séquences correspondent à des formes fléchies de verbe en français, mais en coréen, ce type de séquences présente une complexité très importante : d'une part, les adjectifs prennent autant de formes fléchies que les verbes ; et d'autre part, le nombre des combinaisons de suffixes de conjugaison est extrêmement élevé. . . Prenons un exemple : 노래하다가 <i>jabilhayeneundei</i> [NS-V-PPP].

Table 1-1. Les types de mot coréen.

Dans le chapitre 3, nous discuterons de la structure de dictionnaire et de la méthode qu'on utilisera pour retrouver des mots dans les dictionnaires. Nous allons comparer plusieurs méthodes déjà proposées ainsi que la méthode que nous utilisons dans ce travail.

Le chapitre 4 est consacré à la construction de notre système d'analyse morphologique du coréen. Ce système permettra de reconnaître toutes les analyses possibles pour une séquence donnée : ce système garantira tous les découpages corrects. Nous allons le discuter en détail avec des exemples concrets.

2. Structure du dictionnaire établi pour ce travail

Dans les dictionnaires à usage humain, une entrée lexicale contient plusieurs types d'information [17], [21]. Les dictionnaires utilisables par la machine pour le TNL doivent comporter des informations plus explicites que les dictionnaires faits pour l'usage humain. Une entrée lexicale peut contenir au moins quatre types d'information: la forme phonologique, la catégorie syntaxique, l'indication des irrégularités morphologiques, et le sens [8], [40].

Cependant les dictionnaires coréens existants (sous forme imprimée ou électronique) ont été conçus pour l'usage humain, ils ne sont donc guère satisfaisant pour notre projet. La classification des entrées lexicales en parties du discours n'est pas faite de façon explicite et cohérente. Par exemple, il n'y a pas de critères formels qui permettent de distinguer les verbes des adjectifs. L'information sur la relation dérivationnelle entre les entrées n'a pas été étudiée de manière systématique et exhaustive [24].

Notre système DECO prend en considération de telles informations [24]. L'Analyseur Morphologique, EPL, basé sur DECO, analysera les phrases entrées. L'unité minimale de l'analyse de EPL est le mot et celui-ci est classifié en trois types de structures **T1**, **T2** et **T3**.

Toutes ces structures sont définies dans la table 1-1 et le mot est analysé selon l'une d'elles par l'analyseur morphologique. Chaque classe contient au moins un dictionnaire de formes fléchies la grammaire équivalente au dictionnaire du ND appartient à la classe T1, le dictionnaire du FN à la classe T2, et le dictionnaire du FFP à la classe T3. Nous discuterons plus en détail ces classes dans les sections suivantes.

2.1. Structure de type 1

Le nom se trouve dans T1. Le mot classé T1 est défini ainsi:

$$T1 = PF^* NS^+ SF^* PPN^?$$

La structure T1 contient au moins un NS. Les autres composantes sont facultatives. Lors de l'analyse d'un mot, nous rencontrons nécessairement des PF et SF qui ont des structures complexes dans T1.

En réalité, l'occurrence répétitive des catégories PF et SF, représentée comme PF^* et SF^* , est très rare, mais NS est très productif. Cela nous donne l'idée de changer T1 en T1' qui n'a pas de structures répétitives dans les cas de PF' et SF'. T1' sera donc défini de la façon suivante :

$$T1' = PF'^? NS^+ SF'^? PPN^? = ND PPN^? \quad (ND = PF'^? NS^+ SF'^?)$$

ND correspond aux noms dérivés de PF, NS et SF. PF' et SF' dans T1' doivent inclure les cas rares de PF et SF dans T1. Par exemple, “상” (haut) et “품” (qualité) sont des PF de T1, mais “상품” (haute qualité) n'est pas un PF de T1, parce que PF est dans le dictionnaire simple. Dans ce cas, nous ajoutons le mot “상품” (haute qualité) dans PF' de T1' avec la marque de répétition.

Pour construire les dictionnaires PF' et SF' de T1', nous devons examiner linguistiquement toutes les possibilités d'entrées répétitives qui peuvent se produire dans PF et SF de T1. Les structures répétitives peuvent être formulées de la façon suivante :

- R1** $(\exists x) (\exists y) ((X \subset \text{Fact}(T1)) \wedge x \in X \wedge y \in X \text{ et } (x \bullet y \in X))$
 (• est l'opérateur de la concaténation)
 alors $\text{Fact}(T1) = PF^* | NS^+ | SF^* | PPN^? | PF^* NS^+ | NS^+ SF^* | SF^* PPN^?$
 $| PF^* NS^+ SF^* | NS^+ SF^* PPN^? | T1$
- R2** $(\exists x) (\exists y) ((X^* \subset \text{Fact}(T1)) \wedge x \in X \wedge y \in X \text{ et } (x \bullet y \notin X))$

Elles peuvent produire des ambiguïtés interne et externe.

R1 est la structure concernant l'ambiguïté interne dans le dictionnaire. Par exemple, chaque “연” (total) (x) et “두” (tête) (y) sont dans PF, et la concaténation “연두” (vert pâle) ($x \bullet y$) est aussi dans PF. Dans ce cas, il génère les marques ambiguës. Mais toutes les entrées des dictionnaires contenant PF, SF et NS sont des *unités simples*. Cela signifie qu'aucune entrée des dictionnaires ne peut être séparée en deux entrées. “연두” (vert pâle) n'a donc ni le sens de “연” (total), ni celui de “두” (tête). “연두” (vert pâle) est l'unité autonome la plus petite. Notre système DECO ne génère pas le problème **R1**, car il est basé sur des *unités simples*.

R2 est la structure concernant l'ambiguïté externe dans le dictionnaire. PF et SF de T1 ont rarement des cas de répétition. Nous pouvons alors énumérer linguistiquement toutes les entrées répétitives dans PF et SF. Les nouvelles données obtenues seront gardées dans PF et SF de T1'. Ainsi on supprime les problèmes de répétition dans PF et SF de T1'.

D'autres types d'ambiguïtés existent encore dans T1'. Il ne s'agit plus d'ambiguïtés de répétition, mais d'ambiguïtés d'interférence entre les dictionnaires. Pour NS, il existe toujours des ambiguïtés de répétition, mais nous ne sommes pas intéressé au problème des répétitions dans ce travail, à cause de leur caractère productif.

Les ambiguïtés d'interférence sont définies ainsi [19]:

P1 $(x \in X \text{ et } x \in Y) \text{ et } (X \neq Y) \text{ } X \subset \text{Fact}(T1') \text{ et } Y \subset \text{Fact}(T1')$

P1 est le problème de l'homonymie entre les dictionnaires x et y. Par exemple, “ㅈ” est dans SF et aussi dans PPN.

P2 $((XY \subset \text{Fact}(T1') \wedge x \bullet z \in X \wedge y \in Y \text{ et } (x \in X \wedge z \bullet y \in Y))$

P2 est le problème du croisement mutuel. Par exemple, le mot “복사실” (Salle de photocopie) peut être marqué au niveau de la morphologie comme les trois candidats suivants :

- 1) “복ㅈ” (la pêche) (PF) (le mot sino-coréen) et “실” (la salle) (NS),
- 2) “복ㅈ” (la photocopie) (NS) et “실” (la salle) (SF),
- 3) “복” (duplicatif) (PF) et “ㅈ실” (le fait) (NS).

Les cas 1) et 3) produisent l'ambiguïté **P2**. Pour ces cas, nous devons fournir l'information linguistiquement correcte. Le cas 2) sera sélectionné et cette information sera gardée dans BDD.

P3 $((XY \subset \text{Fact}(T1') \text{ et } XZ \subset \text{Fact}(T1') \wedge x \bullet z \in X \wedge y \in Y \text{ et } (x \in X \wedge z \bullet y \in Z) \text{ et } (X \neq Y \neq Z))$

P3 est le problème de l'interférence. Ce problème peut se produire seulement entre NS, SF, et PPN.

EPL peut détecter les ambiguïtés dans l'analyse du mot. Ces ambiguïtés sont dues à la répétition, l'interférence, la syntaxe et la sémantique, etc. Dans notre étude, nous nous limiterons à la répétition et l'interférence. La sélection de l'information correcte par rapport aux problèmes d'interférence **P1**, **P2** et **P3** se fera de manière semi-automatique.

Toutes les séquences acceptables dans T1' appartiennent à la sphère 1 de la figure 2-1. La sphère 1 est un ensemble. Les données de cette sphère sont produites automatiquement.

Toutes les données comprises dans la sphère 2 sont ambiguës. Les données proviennent automatiquement de la répétition et de l'interférence que nous venons de mentionner.

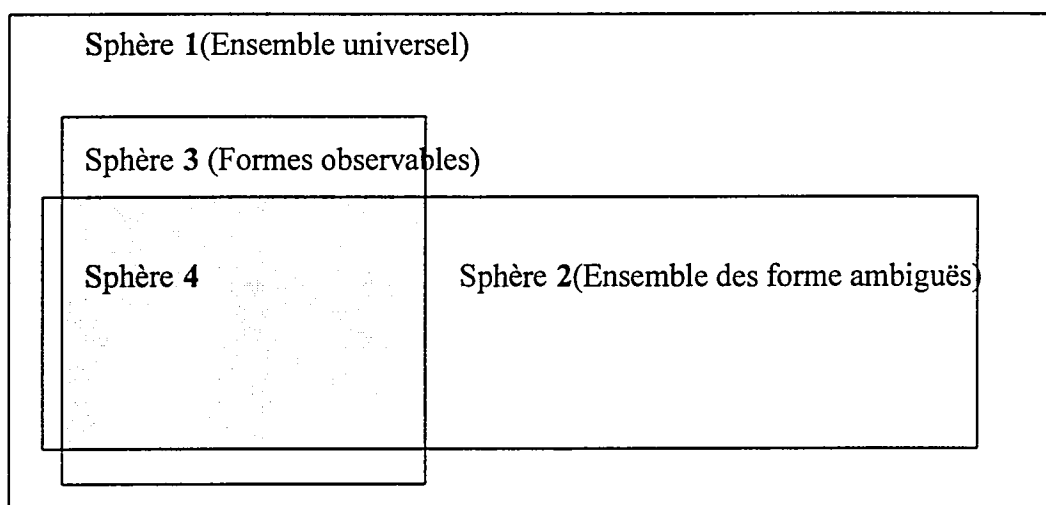


Figure 2-1. Procédure de désambiguïsation dans notre système.

Les données de la sphère 2 ne sont pas toutes observables dans la langue réelle. Par exemple, “우두통” est analysé en deux formes comme “우” (la droite ou la vache) (PF) avec “두통” (mal à la tête) (NS) et “우두” (le chef) (PF) avec “통” (le tonneau) (NS). Bien que la forme présente théoriquement une ambiguïté P2 entre PF et NS, cette ambiguïté n'est pas effective. Nous devons donc choisir linguistiquement les données réelles dans la sphère 2 et elles seront gardées dans le module BDD. Par conséquent, nous re jetterons l'exemple donné ci-dessus.

Les données sélectionnées linguistiquement occupent la sphère 4 dans la figure 2-1. La procédure de sélection dans les sphères 2 à 4 est exécutée par la sphère 3. La sphère 3 représente l'être humain. Autrement dit, elle implique 'linguistiquement'. L'expression 'linguistiquement' peut signifier 'par des linguistes' ou ' par des êtres humains à l'aide des dictionnaires publiés dans le commerce'[37].

2.1.1. Nom Simple

Le nom est la plus importante des parties du discours. Il peut être forme d'une combinaison de parties du discours telles que PF-SN, SF-SN, ou PPN-SN. Il est difficile de reconnaître chaque catégorie sans ambiguïté dans le mot qui comporte des éléments. Les exemples suivants illustrent le type NS <Figure 2-2> [24].

NS	CODE	Information
가	NS	
가간	NS	
가감	NS	/PRED1/PRED3
가객	NS	/HUM
가게	NS	/*PREDHA
가격	NS	
가결	NS	/PRED1/PRED3

Figure 2-2. Exemples de entrées de NS.

Notre dictionnaire a été amélioré à plusieurs reprises, et nous avons obtenu finalement 13 331 entrées de NS.

2.1.1.1. Structure de répétition (NS⁺)

Les structures de noms coréens peuvent être décrites de la manière suivante (PF désigne un préfixe, défini ici comme élément non-autonome s'ajoutant en tête d'un mot; SF est un suffixe, élément non-autonome s'ajoutant à la fin d'un mot; NS désigne un nom simple qui est un élément autonome) (Pour le détail, consulter *Nam 94* [24]). Considérons :

- (1)
 - a. NS
 - b. PF-NS
 - c. PF-NS-SF
 - d. PF-PF-NS
 - e. NS-SF-SF
- (2)
 - a. NS-NS
 - b. PF-NS-NS
 - c. NS-NS-SF
 - d. NS-NS-NS

Seule la forme (1a) sera traitée comme nom simple; les formes (1b) ~ (1e) sont dérivées; (2) sont des noms composés. La distinction entre (1a) et (1b) ~ (1e) est faite en fonction de l'existence de préfixes et/ou de suffixes en dehors d'un élément autonome (i.e. NS).

La distinction entre (1) et (2) est effectuée selon le nombre de noms(NS), c'est-à-dire dans (2) le nombre de noms est égal ou supérieur à 2. Soulignons donc que la

typologie décrite ici ne correspond pas aux modèles présentés par certaines grammaires scolaires qui, d'une part, distinguent différents niveaux morphologiques, et d'autre part considèrent des unités de type PF-X ou X-SF (X étant un élément non-autonome) comme dérivées, alors que pour nous ce sont des noms simples que nous ne décrivons pas autrement que NS (ainsi, nous ne intéressons pas à l'analyse en morphèmes des mots qui ne possèdent aucune partie autonome à l'intérieur). Par conséquent, même avec le même préfixe, si l'élément qui suit est une unité autonome, cette combinaison sera reconnue comme dérivée (ou affixée), sinon comme simple, donc nom simple.

Voici des exemples qui illustrent respectivement les types de combinaisons (1) et (2) :

- (3)
- | | | | |
|----|--------|---------------------------------------|---|
| a. | 가을 | (automne) | |
| b. | 재-결합 | (ré-unification | = réunification) |
| c. | 포도-주 | (alcool de raison | = vin) |
| d. | 동-성-애 | (même-sexe-amour | = amour homosexuel) |
| e. | 가-출-옥 | (provisoire-sortir-prison | = sortie provisoire de prison) |
| f. | 경리-부-장 | (comptabilité-département-responsable | = responsable du département de comptabilité) |
- (4)
- | | | | |
|----|----------|--------------------------------|-----------------------------------|
| a. | 구슬-땀 | (perle-sueur | = goutte de sueur de forme perle) |
| b. | 반-독재-투쟁 | (anti-dictature-révolte | = révolte anti-dictature) |
| c. | 간호-보조-원 | (infirmier-assistance-personne | = assistant d'infirmier) |
| d. | 가운뎃-손-가락 | (centre-main-doigt | = majeur) |

On pourrait supposer a priori d'autres types de combinaisons possibles (e.g. PF-PF-NS-SF), mais le nombre des formes autre que (1) et (2) est extrêmement réduit. Par contre, des séquences chaînées de noms simples de type (5) s'observent fréquemment:

- (5)
- | | | | | | |
|--------|-------------|-------|-------------|---------|-------------|
| 국어 | 맞춤 | 법 | 통일 | 개정 | 안 |
| coréen | orthographe | règle | unification | réforme | proposition |
- (La proposition de réforme pour l'unification des règles d'orthographe du coréen)

La séquence (5) est composée de six noms simples (NS NS NS NS NS NS). Rappelons que la d'elimination des mots typographique pas de blancs ne peut par servir de critère opératoire, car on observe également:

- (6)
- | | | | | |
|----|--------------------------|---------------------------------|-------------|---------------------|
| a. | 국어 | 맞춤법 | 통일 | 개정안 |
| | coréen | orthographe-règle | unification | réforme-proposition |
| b. | 국어맞춤법 | 통일개정안 | | |
| | coréen-orthographe-règle | unification-réforme-proposition | | |

Pourtant, bien que cette d'elimination typographique ne soit pas normalis鉶, il existe des contraintes:

- (7) a. *국어맞춤법 통일개정안
 *coréen-orthographe règle unification-réforme proposition

Cela est probablement dû au fait que les formations des mots composés à l'intérieur de cette séquence ne se font pas au même niveau morphologique.

Ce type de nom composé est la structure que nous avons définie comme **R2** de NS. Il s'agit d'une structure où plusieurs noms apparaissent successivement. Ce type (symbole + dans NS⁺) peut engendrer plusieurs problèmes. Bien que peu productive cette structure ne peut pas être ignorée. La répétition de noms simples peut avoir lieu deux ou trois fois selon la règle générale [15]. Nous n'énumérerons pas ici toutes les entrées de ce type de Nom Composé.

NS sera utilisé obligatoirement dans T1', mais PF' et SF' sont différents. C'est pourquoi PF' et SF' sont définis comme PF'[?] et SF'[?], alors que NS est défini comme NS⁺ dans la structure T1'. PF' se trouve toujours en tête de NS, alors que SF' est placé à la fin de NS. S'il y a un conflit entre les séquences répétitives de NS et PF' ou SF', on pourra résoudre le problème par le principe que la structure répétitive est autorisée uniquement pour NS.

2.1.2. Préfixe

Le PF, en tant qu’affixe, s’attache en tête du nom. La forme complexe basée sur le nom restera toujours nom. Les entrées du dictionnaire de PF sont des unités simples. Autrement dit, aucune entrée du dictionnaire de PF ne peut être divisée en plus petites unités sans changer le sens de l’entrée. Actuellement, il y a 505 entrées distinctes dans le dictionnaire de PF. La figure 2-3 montre les entrées simples de PF. Par exemple, l’entrée ‘가등기’ (l’enregistrement provisoire) dans la figure 2-3 est composée du PF ‘가’ (provisoire) et du NS ‘등기’ (l’enregistrement).

PF	CODE	Information	Exemples
가	PF.		가극, 가곡
가	PF.		가등
가	PF.		가등기, 가석방
가	PF.		가법, 가례
가	PF.		가홍
가	PF.	PREDH	가해, 가열
가랑	PF.		가랑비, 가랑잎
가시랑	PF.		가시랑비

Figure 2-3. Exemples des entrées de PF.

2.1.2.1. Structure répétitive PFⁿ (si n >= 2))

Le PF est utilisé de manière récursive dans la structure T1. Cela rend difficile l’analyse des mots contenant des PF et peut produire des ambiguïtés. Nous allons transformer la structure répétitive du PF en celle non-répétitive qui a été définie dans T1’. Cela veut dire que nous allons chercher toutes les entrées répétitives dans le dictionnaire des PF de T1. Les entrées trouvées seront insérées dans le dictionnaire de PF. La procédure consiste en la transformation de PF^{*} dans T1 en PF[’] dans T1’. Cette procédure permettra de supprimer la structure R2 dans PF.

Nous avons extrait linguistiquement toutes les entrées combinées de PF dans T1. Les entrées sélectionnées se trouvent dans la sphère 4 de la figure 2-1.

Supposons que le nombre de reprises est deux(si, n=2). Le nombre candidats de R2 correspond au Numéro_de (PF) multiplié par le Numéro_de (PF). C’est-à-dire que 505*505 (255,025) est le nombre d’entrées de l’expression PF². Parmi les 255,025 entrées, nous avons trouvé de façon manuelle 26 nouvelles entrées qui sont représentées dans la liste 2-1 ci-dessous, et elles sont ajoutées au dictionnaire PF’. Les autre cas PFⁿ (si, n >= 2)) n’existent pas dans le monde réel des mots.

고품 극단 극대 극소 극한 급회 남고동 냉난 다갈 맨헛 맹혹 상품 생까막 순까막 순헛 역순 온갖 온난 저품 적녹 중품 진고동 청녹 최단 하품 한냉(26)
--

Liste 2-1. Nouvelles entrées du PF’.

2.1.2.2. Ambiguïté P1

Il y a 262 entrées ambiguës de **P1** produites entre PF et NS comme le montre l'Annexe A. Les PF sont toujours employés immédiatement avant un NS, les NS peuvent être employés sans PF. Cela permet de déterminer si une forme est PF ou NS en cas d'ambiguïté. Une autre solution serait d'énumérer toutes les formes possibles entre PF et NS. Cela demandera beaucoup de temps.

2.1.2.3. Ambiguïté P2

Afin d'établir la liste des données ambiguës de **P2** entre PF et NS autrement dit, afin d'extraire toutes les données se trouvant dans la sphère 2 sous la structure d'ambiguïté **P2** de T1', nous avons construit le programme de la figure 2-5 qui est la forme algorithmique de la figure 2-4. Les nombres maximum de syllabes du PF et du NS sont réciproquement 3 et 6.

$$\text{mots} = \prod X_m, \text{ donc } \text{PF} = \prod_{k=1}^j X_k, \text{ NS} = \prod_{k=j+1}^n X_k, \text{ SF} = \prod_{k=1}^l X_k, \text{ PPN} = \prod_{k=l+1}^m X_k$$

Figure 2-4. Représentation des mots incluant les entrées de PF et de NS.

```
for(j = 0; j < Maximum_Length_Of_PF; j++)
  for(i = j+1; i <= Maximum_Length_Of_PF; i++)
    if( (Look_Up(PF,  $\prod_{k=1}^j X_k$ ) && Look_Up(NS,  $\prod_{k=j+1}^n X_k$ ) &&
        (Look_Up(PF,  $\prod_{k=1}^i X_k$ ) && Look_Up(NS,  $\prod_{k=i+1}^n X_k$ ))
        printf("%s %s »,  $\prod_{k=1}^j X_k$   $\prod_{k=j+1}^n X_k$ ,  $\prod_{k=1}^i X_k$   $\prod_{k=i+1}^n X_k$ ); /* Ambiguïté P2 */
```

Figure 2-5. Programme d'extraction des données **P2**.

Il n'y a pas d'ambiguïté **P2** entre $j=3$ et $j=2$, et entre $j=3$ et $j=1$. Dans d'autres cas, le mot peut être marqué selon l'un des trois cas suivants:

- 1) PF (longueur 0), NS (la longueur est n) ou
- 2) PF (longueur 1), NS (longueur $n-1$), ou
- 3) PF (longueur 2), NS (longueur $n-2$).

Les formes qui présentent une ambiguïté **P2** sont 2,739. Elles sont extraites automatiquement à partir du programme dans la figure 2-5. Ces formes se trouvent dans la sphère 2 de la figure 2-1 sous l'ambiguïté **P2**. La liste figure dans l'Annexe B-1. Nous avons extrait manuellement de cette liste les formes susceptibles d'être dans la

sphère 4. l'Annexe B-2 montre la liste de ces formes. Elle sera gardée dans BDD de la figure 1-1 avec les marques correctes. La marque 'S' dans l'Annexe B-2 signifie que le nom se trouve dans les dictionnaires publiés [37].

Le SF s'attache à la fin du nom. La forme combinée avec le nom est classée comme nom. Le problème de la répétition se pose comme pour le PF. Actuellement, 468 entrées distinctes sont dans le dictionnaire SF. La figure 2-6 montre les exemples de SF dans les entrées. Par exemple, 대학가 (le quartier universitaire) est composé de 대학 (l'université) et de 가 (le quartier) et il est marqué comme NS-SF.

Figure 2-6. Les exemples des entrées du SF

P1 = {댁 (chez), 땡 (camarade), 석 (place), 실 (bureau),
 역 (membre), 측 (parti), 파 (parti)}

P2 = {네 (chez), 하 (cher)}

P3 = {전 (à), 후 (après)}

S1 = {님 (cher)}

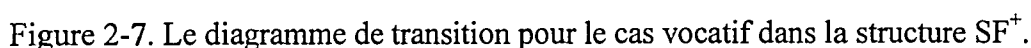
S2 = {들 (-s)}

S3 = {간 (entre), 외 (sauf)}

T1 = {선생 (professeur), 여사 (madame) }

T2 = {양 (mademoiselle, 씨 (monsieur), 군 (garçon)}

E = Vide



19

des formes du *vocatif*. Les éléments terminaux du diagramme sont des éléments vocatifs.

Certains SF représentent la *position sociale*. Pour exprimer les positions sociales variées des gens, plusieurs types des suffixes sont susceptibles d'être attachés à la fin du nom. Ils correspondraient à monsieur, docteur, président, membre, etc. en français.

La figure 2-8 est un diagramme de transition des suffixes de *position* pour le nom. Dans ce diagramme, S2(pluriel) est utilisé librement. Il se trouve en tête ou en fin des entrées P1 ou P2.

A l'exception de ces deux cas, SF n'a pas de structure répétitive dans le mot. Nous pourrions établir le processus pour les deux cas avec des règles, car ils comprennent des formes très régulières. Les règles figurent dans le Module BDD filtrage de la figure 1-1. Dans le Lexique-Grammaire, on appelle cette opération régulière une grammaire locale[13] [18]. Sauf pour les deux cas que nous venons de mentionner, nous pouvons redéfinir facilement la structure SF* dans T1 comme SF? dans T1'.

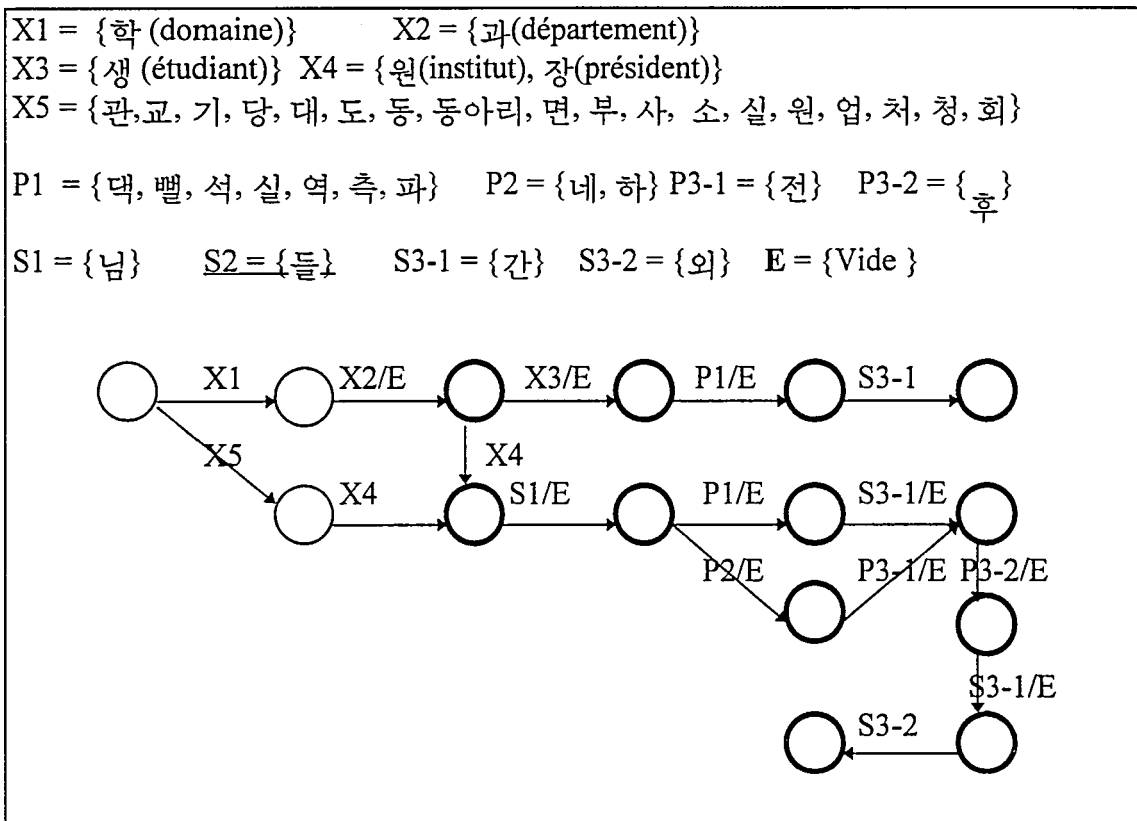


Figure 2-8. Diagramme de transition pour *position sociale*.

2.1.3.2. Ambiguïté P1

Il y a 253 entrées P1 entre NS et SF comme le montre la liste 2-2. Le conflit entre NS et SF est résolu de la même manière que pour les ambiguïtés P1 entre PF et NS. Les SF sont toujours employés immédiatement après un NS, les NS peuvent être employés sans SF. Cela permet de déterminer si une forme est SF ou NS en cas d'ambiguïté. Nous pouvons donc distinguer le SF de NS sans difficulté.

가 가락 가리 각 간 감 갑 강 개 객 거리 거미 건 건너 검 격 결 경 계 곡 골 공 과 괄
관 광 구리 구원 국 군 굶 꺾 귀 금 급 기 길 깨 꺼 풀 꼬 리 폴 끝 난 남 납 날 내
내 기 널 노 농 눈 님 다리 단 담 답 당 대 대가 리 대기 닥 도 독 동 동아 리 동이 동자
들 딱 지 띠 마 마루 막 만 말 망 망나 니 매 맥 머리 면 명 모 목 무 무리 묵 문 물 미
바 가지 바 구 니 바 라 지 바 람 바 지 박 반 발 밥 방 방면 배 배기 백 번 벌 범 벼 락 벽
변 별 보 라 복 본 봉 부 부리 부 지 북 송 이 분 불 비 빵 사 산 살 상 새 생 원 서 리 서 방
선 선생 설 성 세 소 속 손 술 수 순 술 시 식 신 실 씨 아 비 악 안 암 애 액 약 양 업
역 연 염 영 오 리 요 욕 용 운 원 위 윤 의 이 이전 인 임 입 자 자 락 자 위 잔 장 재
쟁 저 적 전 절 점 정 젓 조 중 좌 주 주부 주사 주의 죽 지 중 지 개 지 기 지방 진 질
집 짝 차 참 창 채 책 처 천 첩 청 체 초 총 추 축 춤 치 기 침 칼 탄 태 기 톱 토 톱 통
통수 파 판 패 기 펄 포 품 풍 피 학 한 함 합 해 향 호 혼 화 회 효 훈 (253)

Liste 2-2. Les entrées ambiguës P1 entre NS et SF

2.1.3.3. Ambiguïté P2 avec NS.

Dans l'expression de la figure 2-4, une ambiguïté P2 entre NS et SF peut se produire. La procédure d'extraction de P1 entre NS et SF est la même que pour le cas de P2 entre PF et NS. Les longueurs maximum du NS et du SF sont réciproquement 6 et 3 syllabes. Les variables l , e , et m sont utilisées dans l'expression à la place de j , et n pour le cas de PF.

Nous avons obtenu automatiquement 2,272 données P2 qui appartiennent à la sphère 2 sous la structure T1'. La liste de ces données est dans l'Annexe C-1. A partir de cette liste, nous avons extrait les mots réels de la sphère 4. Les marques correctes des formes sont enregistrées dans BDD. La liste est dans l'Annexe C-2.

2.1.3.4. Ambiguïté P1 avec PPN

Il existe 13 SF qui peuvent produire l'ambiguïté P1 avec PPN comme dans la liste 2-3 Si le mot contient un des 13 SF, EPL l'analysera comme forme ambiguë. Nous avons appliqué la méthode précédente à cette forme. D'abord, nous avons obtenu automatiquement les formes combinées de NS avec les 13 SF. La liste de ces formes combinées se trouvent dans la sphère 2 et elles peuvent être marquées comme NS-SF ou NS-PPN. Le nombre minimum de formes atteint 173,303 (= 13,331 * 13) (Nombre de (NS) * Nombre de (Ambiguïté P1 entre SF et PPN)).

Nous avons examiné cette liste et extrait manuellement les formes correctes, qui sont dans l'Annexe D. Les données seront sera enregistrées dans BDD. L'extrait est présenté dans la table 2-1.

가과 관게도로론만서선이 의이(13)

Liste 2-3. L’ambiguïté P1 entre SF et PPN.

NS-SF forme	Nombre	NS-SF forme	Nombre	NS-SF forme	Nombre
NS + 가(SF)	597	NS + 과(SF)	352	NS + 관(SF)	366
NS + 께(SF)	13	NS + 도(SF)	456	NS + 료(SF)	94
NS + 론(SF)	885	NS + 만(SF)	8	NS + 셔(SF)	410
NS + 선(SF)	442	NS + 아(SF)	32		
NS + 의(SF)	16	NS + 이(SF)	18		

Table 2-1. L’information BDD pour les formes NS-SF.

Par exemple, les 597 formes pour “가” sont marquées comme NS-SF dans la table 2-1. Cela signifie que si on trouve un mot qui peut être marqué comme NS-SF et NS-PPN (i.e. il est ambigu), et s’il n’est pas dans BDD, il sera marqué comme NS-PPN. Le nombre de formes dans la table 2-1 est celui des formes de NS qui peuvent être combinées avec les formes SF.

2.1.4. Postposition du Nom

En français, l'unité minimale pour l'analyse lexicale est, en général, la phrase contenant les blancs. Le dictionnaire doit donc contenir de l'information au delà du niveau morphologique.

Le dictionnaire de PPN est la base de données pour les noms, il fournit les indications sur les variantes morphologiques. Il existe des marques de fonction grammaticale comme les cas : *nominatif*, *accusatif*, *datif*, *locatif*, etc. qui se combinent avec les noms sans blanc. Une séquence composée de deux éléments comme 'à Paris' correspond à un élément 'Paris-Locatif' en coréen. PPN se situe après le nom. La figure 2-9 donne des exemples de la liste du dictionnaire PPN.

PPN	CODE	Information	
가	PPN	\PN1	.nmtf .postp
가라도	PPN	\PN1	.nmtf .postp
같이	PPN	\PN1 \PN2	.postp
까지	PPN	\PN1 \PN2	.nmtf .acc .postp
까지나	PPN	\PN1 \PN2	.nmtf .acc .postp

Figure 2-9. Exemples de la liste PPN.

2.1.4.1. Ambiguïté P3.

L'ambiguïté **P3** comporte des formes très complexes. Les deux formules ci-dessous peuvent générer l'ambiguïté **P3** entre NS, SF, et PPN. Nous avons déjà décrit **P3** en détail. On a les deux possibilités ci-dessous pour **P3**:

- 1) $x \bullet z \in \text{NS}$ et $y \in \text{SF}$ et $x \in \text{NS}$ et $z \bullet y \in \text{PPN}$.
- 2) $x \bullet z \in \text{NS}$ et $y \in \text{PPN}$ et $x \in \text{NS}$ et $z \bullet y \in \text{SF}$.

Les 189 formes ont été extraites automatiquement au moyen des deux formules (Liste 2-4). Les 189 candidats de la liste peuvent produire l'ambiguïté **P3**. De cette liste, nous avons obtenu manuellement les trois formes de la liste 2-5.

가동아 가동이 각주의 간주의 감동아 감동이 감주의 강박이 객주의 격동아 격동이 결박이 경쟁이
경주의 골동아 골동이 공동아 공동이 공박이 공주의 군주의 금주의 기동아 기동이 난동아 난동이
냉동아 냉동이 노동아 노동이 논박이 논쟁이 단주의 대동아 대동이 도박이 도주의 독주의
뒤주의 맥주의 면박이 명주의 목동아 목동이 무동아 무동이 물주의 미동아 미동이 밀주의 반동아
반동이 반박이 반주의 발동아 발동이 발주의 백주의 범주의 변동아 변동이 변주의 부동아 부동이
부주의 분쟁이 사동아 사동이 사주의 상동아 상동이 상주의 선동아 선동이 선박이 선주의 소동아
소동이 소박이 소주의 속박이 수동아 수동이 수박이 시동아 시동이 시주의 신동아 신동이
악동아 악동이 안주의 애주의 약주의 양주의 연동아 연동이 연주의 염주의 영주의 요동아 요동이
우동아 우동이 우박이 우주의 운동아 운동이 위주의 음주의 이동아 이동이 이주의 인주의 임박이
입주의 차동아 차동이 차주의 재주의 저주의 전동아 전동이 전쟁이 전주의 점주의 정박이 정쟁이
정주의 종주의 주동아 주동이 진동아 진동이 진주의 질주의 찬동아 찬동이 책동아 책동이 청동아
청동이 청주의 초동아 초동이 축주의 탈주의 태동아 태동이 터주의 토주의 투쟁이 파동아 파동이
패주의 편주의 폐갱이 포박이 포주의 폭동아 폭동이 폭주의 피동아 피동이 학동아 학동이 함박이
합동아 합동이 합주의 해동아 해동이 향동아 향동이 호박이 호주의 혼동아 혼동이 활동아 활동이
활주의 황동아 황동이 회동아 회동이 동동주의 조물주의 구세주의 두레박이 (183)

Liste 2-4. Les candidats à l'ambiguïté **P3** entre NS, SF et PPN.

대동아	:	대(PF) 동아(NS), 대동아(Nom Propre)
반동아	:	반동(NS) 아(SF)
부주의	:	부(PF) 주의(NS)

Liste 2-5. La liste des ambiguïtés **P3** avec les marques.

2.2. Structure de type 2

La Forme Nominalisée est construite à partir du verbe ou de l'adjectif avec un suffixe de nominalisation(SFN). SFN est un cas spécial de PPP. FN est défini comme la classe T2 avec PPN. FN comporte le Verbe Nominalisé(VN), l'Adjectif Nominalisé (AN), et les Formes Spéciales.

Les FN fonctionnent comme le nom dans la phrase. Mais elles sont un peu différentes de la notion conventionnelle de nom. C'est-à-dire, FN peut représenter une séquence phrastique. Dans cette séquence, FN joue le rôle de l'infinitif en français ou celui du gérondif en anglais.

Par exemple, '뛰기' (courir) fonctionne comme l'infinitif en position de sujet en français ou comme le gérondif en anglais dans la phrase "뛰기가 어렵다"(Courir est difficile.) (Running is difficult). Dans une autre phrase, comme "뛰는것이 어렵다"(Courir est difficile.) (To run is difficult), "뛰는것 (Courir) (To run) est une autre forme de FN correspondant à l'infinitif en anglais.

Au niveau de la morphologie, il n'est pas important de distinguer la notion de FN de celle de NS. Dans l'interprétation syntaxique, la différence d'information est importante, parce que FN exige une structure de phrase différente de celle de NS.

T2 a une structure supplémentaire par rapport à T1'. FN fonctionne comme un nom dans la phrase. L'analyse de la structure T2 est la même que celle de T1' :

$$T2 = (PF^2(FN | NS) + SF^2 | \text{Formes Spéciales}) PPN^2$$

FN est composé d'un radical et d'un SN. Par exemple, le FN "뛰기" (courir) est analysé comme "뛰" (cour) (*radical*) et "기" (suffixe complémentaire). Dans la structure T2, PPN peut être soudé à la fin de FN. "뛰기가" (courir est) est analysé comme la combinaison de FN "뛰기" (courir) et de PPN "가" (*nominatif*).

La liste 2-6 donne des suffixes typiques pour FN[36]. "이", "개" et "애" qui s'utilisent dans des expressions figées. Les mots contenant ces suffixes sont considérés comme des noms et leur liste est déjà incluse dans NS. Dans le cas de '기', cette syllabe peut être soudée indépendante de la syllabe finale de radical à la fin du radical. Mais le choix des autres éléments de la liste comme '음' (d'un *suffixe complémentaire*) et 'ㅁ' (d'un *suffixe complémentaire*) dépend de la Syllabe Finale du Radical (SFR). Les formes fléchies des prédicats dépendent de SFR. Les formes fléchies des SFN dépendent aussi de SFR.

이 기 개 애 음 ㅁ(6)

Liste 2-6. SFN pour FN.

2.2.1. Verbes Nominalisés

Pour construire le lexique VN, on ajoute tout simplement SFN “기”(suffixe complémentaire) à la fin du radical des verbes. La forme de base de l'exemple 1) est “받다”(recevoir). Le radical “받” est la forme dont la dernière syllabe “다” est supprimée. La forme nominalisée du verbe se fait en ajoutant le SFN “기” après le radical “받”. Dans ce cas, SFR et le radical sont identiques. Nous obtenons la forme nominalisée 1). La même procédure fournit l'exemple 2).

- 1) 받기 (la reçue)
- 2) 하기 (faire)

TYPE	CF dans SFR	Substitution	Nombre	Total	Exemple
NVT1		음	353	353	많다(많음)
NVT2-1		ㅁ	7,857	7,860	하다(함)
NVT2-2		non exister	3		건너다
NVT3	ㄷ	ㄷ + 음	44	79	받다(받음)
NVT4		ㄷ + 음	35		깨달다(깨달음)
NVT5-1	ㄱ	ㅁ	259	267	길다(김)
NVT5-2		ㄱ + 음			길다(길음)
NVT6		ㅅ	8		살다(삶)
NVT7	ㅂ	ㅂ + 음	68	84	꼬집다(꼬집음)
NVT8		음	14		눅다(누음)
NVT9		ㅁ	2+1(NVT8)		뵈다(뵈음)
NVT10-1	ㄴ	ㅁ	28	47	낫다(남)
NVT10-2		음			낫다(나음)
NVT11		ㄴ + 음	19		벗다(벗음)
NVT12-1	ㅎ	ㅁ	73	73	놓다(눔)
NVT12-2		ㅎ + 음			놓다(놓음)
				8,763	

Table 2-2. Les Classes des VN.

Il existe d'autres formes de SFN, comme “음” et “ㅁ”. En principe, la première forme se combine avec SFR qui se termine par une consonne, alors que la dernière se combine avec SFR à voyelle.

Examinons les exemples 3), 4) et 5). Ils ont la même consonne finale (CF) de SFR “ㄷ”. Tous les verbes qui ont comme CF du SFR “ㄷ” sont classés NVT3 ou NVT4 dans la table 2-2. L'exemple 3) est classé comme NVT3, mais les deux autres verbes sont classés comme NVT4. Il n'y a pas de principe pour choisir les formes correctes. Parmi 79 verbes, 44 verbes entrent dans la classe NVT3 et le reste des verbes entrent dans NVT4.

- | | | | | |
|----|------------------|-----|---|------------------------------|
| 3) | 받다 (recevoir) | (V) | : | 받음 (la reçue) (FN) |
| 4) | 듣다 (écouter) | (V) | : | 들음 (l'écoute) (FN) |
| 5) | 깨달다 (comprendre) | (V) | : | 깨달음 (la compréhension) (FN). |

L'irrégularité des formes nominalisées dépend de SFR. Autrement dit, si deux verbes différents ont la même SFR, ils auront la même forme nominalisée.

Nous avons établi linguistiquement la liste des VN dérivés du verbes. Ce sont les formes du présent, et leur classification est décrite dans la table 2-2.

En ce qui concerne la construction des VN, l'élément le plus important est le SFR. Le choix des PPP dépend de la SFR. NVT1 et NVT2(-1, -2) sont des classes typiques des suffixes nominalisateurs. Les verbes qui se combinent avec une forme de NVT1 comportent la CF dans SFR, à la différence des verbes de NVT2-1. NVT2-2 est une classe spéciale. Selon le système de code du coréen, appelé **KSC5601-1992**, il n'existe pas de forme nominalisée pour les trois verbes de NVT2-2 [3].

Les formes des VN de NVT5-1, NVT10-1 et NVT12-1 sont obtenues respectivement de NVT5-2, NVT10-2 et NVT12-2. Bien qu'elles s'utilisent très rarement dans les phrases écrites, ces formes sont acceptables.

NVT2-2:	견니다 뇌니다 뛰다(3)
NVT4 :	결문다 결다 곧이듣다 귀넘어듣다 귀담아듣다 귀여겨듣다 그러문다 길다 깨달다 겨문다 내달다 눕다 덧문다 되문다 되받아문다 듣다 묻다 새겨듣다 새기어듣다 알아듣다 얻어듣다 여겨듣다 옛듣다 일컫다 주어듣다 주워듣다 좃어듣다 지내듣다 처신타 치달다 캐문다 캐어문다 헛듣다 설듣다 싣다(35)
NVT6 :	되살다 막살다 못살다 살다 숨어살다 알다 엇혀살다 쥐여살다(8)
NVT8 :	가로눕다 곱다 곱다 깎다 곱다 눕다 돌아눕다 돕다 뒤눕다 드러눕다 몸져눕다 줍다 짜깁다 여쭙다(14)
NVT9 :	뵈웁다 여쭙웁다 여쭙다(2 + 1)
NVT11:	내숫다 드숫다 들숫다 들이웃다 발가벗다 벌거벗다 벗다 비웃다 빋다 빨가벗다 빼앗다 뺏다 빨거벗다 숫다 씻다 앗다 웃다 치숫다 헐벗다(19)

Liste 2-7. Liste des verbes de la table 2-2.

La liste 2-7 donne les verbes des classes de la table 2-2. Les verbes à CF comme “ㄷ” sont classés soit dans NVT3, soit dans NVT4. Les verbes qui entrent dans NVT4 sont énumérés dans la liste 2-7. Remarquons que les verbes qui ne se trouvent pas dans NVT4 tout en comportant CF “ㄷ” sont classés dans NVT3.

2.2.2. Adjectifs Nominalisés

Les procédures de construction des AN sont les mêmes que pour les VN. “기” (suffixe complémentaire) peut être attaché à SFR de façon automatique, alors que “음” et “ㅁ” se combinent différemment avec les adjectifs en fonction des formes de SFR. Les suffixes engendrent beaucoup de formes fléchies et ils dépendent des formes de SFR des adjectifs.

Les deux exemples ci-dessous sont classés comme NAT10 et NAT11. Les 90 adjectifs qui ont le CF “ㅎ” de SFR entrent normalement dans NAT10 comme dans le cas 1). Mais 3 de ces 90 adjectifs entrent dans NAT11. C’est le cas 2). Il n’existe pas de critère pour classer comme NAT10 ou NAT11.

- 1) 노랗다 (est jaune) (A) : 노람 (le jaune) (FN)
- 2) 좋다 (est bon) (A) : 좋음 (la bonté) (FN).

Le dictionnaire de AN est construit linguistiquement. Les résultats sont décrits dans la table 2-3.

TYPE	CF dans SFR	Substitution	Nombre	Total	Exemple
NAT1		음	276	276	맑다(맑음)
NAT2-1		ㅁ	4,303	4,304	가깝하다(가깝함)
NAT2-2		non exister	1		애교다
NAT3-1	ㄷ	ㅁ	40	40	가늘다(가늠)
NAT3-2		ㄷ + 음	40		가늘다(가늘음)
NAT4		름	1		서툴다(서투름)
NAT5	ㅂ	음	391 + 143	538	노엽다(노여음)
NAT6		ㅂ + 음	4		비좁다(비좁음)
NAT7		ㅁ	391 + 70		노엽다(노염)
NAT8	ㅅ	음	1	2	낫다(나음)
NAT9		ㅅ + 음	1		헐벗다(헐벗음)
NAT10	ㅎ	ㅁ	87	90	노랗다(노람)
NAT11		ㅎ + 음	3		좋다(좋음)
				5,250	

Table 2-3. Les Classes des AN.

La liste des adjectifs classés dans la table 2-3 se trouve dans la liste 2-8. Les AN de NAT4 dans la table 2-3 sont les formes qui ont subis la liaison à partir des adjectifs de NAT3-2. Les 70 entrées dans NAT7 ont été sélectionnées parmi les 143 entrées de NAT5. Il s’agit d’exemples typiques qui ne peuvent être sélectionnés que par le linguiste. NAT2-2 est un cas particulier. La forme de AN correspondante n’existe pas dans le système de code du coréen standard[32].

Certains VN et AN s'utilisent très rarement mais on en trouve des formes dans les dialectes et dans la langue parlée.

NAT6 :	비좁다 수좁다 좁다 좁디좁다
NAT7*:	가럽다 가소럽다 가파럽다 감미럽다 건조럽다 경사럽다 경이럽다 고럽다 공교럽다 괴까다럽다 괴럽다 굵다 권태럽다 귀엽다 까다럽다 꼬깝다 꼴사납다 피까다럽다 날카럽다 노엽다 놀랍다 다사럽다 다채럽다 단조럽다 도탑다 두렵다 두텁다 따사럽다 따습다 똥마렵다 뜨습다 마렵다 매섭다 매스껍다 메스껍다 명예럽다 무섭다 반지럽다 번거럽다 보드럽다 보배럽다 사사럽다 상서럽다 새럽다 순조럽다 신비럽다 아니꼽다 아쉽다 애처럽다 어둡다 영예럽다 외럽다 <u>오</u> 애럽다 위태럽다 의럽다 의외럽다 이럽다 자비럽다 자애럽다 자혜럽다 쟁그럽다 정겹다 정교럽다 정의럽다 조화럽다 지혜럽다 평화럽다 한가럽다 향기럽다 호화럽다(70)
NAT12:	맷집중다 의중다 중다(3)

Liste 2-8. Liste des adjectifs de la table 2-3.

2.2.3. Formes Spéciales

Les verbes que nous avons discutés dans les paragraphes 2.2.1 et 2.2.2 sont au présent. Nous discuterons ici les autres *temps* ou *modalités*.

Les trois formes de NPT1 dans la liste 2-6 représentent trois SFN du *mode*. La première “냐기” exprime le mode *interrogatif* et le *discours indirect*. La forme nominalisée avec suffixe flexionnel implique le sens ‘demander si ...’. La deuxième “다기” est la forme transformée du prédicat de *discours indirect*. La dernière “자기” est le SFN *prépositif*. “가자기” (d’aller ensemble) (VN), par exemple, il est composé de “가” (le radical) et de “자기” (SFN).

NPT2, NPT3, NPT4, NPT5, NPT6 et NPT7 sont aussi complexes que NPT1. Les séries “젯” (autrement dit, les SFN dont la première syllabe est “젯”) dans NPT3 dénotent la *supposition* comme modalité et aussi le *futur*. Les séries ‘젯’ dans NPT4, ‘시’ dans NPT5, ‘우’ dans NPT6 et ‘으’ dans NPT7 expriment la *politesse* comme *modalité*. Toutes les formes de la liste 2-9 s’attachent après le SFR et les mots qui les contiennent seront marqués comme FN.

NPT1 :	냐기 다기 자기(3)
NPT2 :	는데(1)
NPT3 :	젯기 젯냐기 젯다기 젯음(4)
NPT4 :	젯젯기 젯젯냐기 젯젯다기 젯젯음 젯기 젯냐기 젯다기 젯젯젯냐기 젯젯젯다기 젯젯젯냐기 젯젯젯다기 젯젯젯음 젯젯음 젯음(14)
NPT5 :	시기 시냐기 시다기 시젯기 시젯냐기 시젯다기 시젯음 시解未 시젯젯냐기 시젯젯다기 시젯젯음 시젯기 시젯냐기 시젯다기 시젯젯젯기 시젯젯젯냐기 시젯젯젯다기 시젯젯젯음 시젯젯기 시젯젯냐기 시젯젯다기 시젯젯젯음 시젯젯음 시자기(24)
NPT6 :	우라기 우젯젯기 우젯젯냐기 우젯젯다기 우젯젯음 우젯기 우젯냐기 우젯다기 우젯젯젯다기 우젯젯젯다기 우젯젯젯기 우젯젯젯냐기 우젯젯젯다기 우젯젯젯음 우젯젯음 우시젯기 우시젯냐기 우시젯다기 우시젯음 우시젯젯기 우시젯젯냐기 우시젯젯다기 우시젯젯음 우시젯기 우시젯냐기 우시젯다기 우시젯젯젯기 우시젯젯젯냐기 우시젯젯젯다기 우시젯젯젯음 우시젯젯기 우시젯젯냐기 우시젯젯다기 우시젯젯음 우시젯음 우시기 우시냐기 우시다기 우시자기 우심(40)
NPT7 :	으라기 으젯젯기 으젯젯냐기 으젯젯다기 으젯젯음 으젯기 으젯냐기 으젯다기 으젯젯젯냐기 으젯젯젯다기 으젯젯젯기 으젯젯젯냐기 으젯젯젯다기 으젯젯젯음 으젯젯음 으시젯기 으시젯냐기 으시젯다기 으시젯음 으시젯젯기 으시젯젯냐기 으시젯젯다기 으시젯젯음 으시젯기 으시젯냐기 으시젯다기 으시젯젯젯기 으시젯젯젯냐기 으시젯젯젯다기 으시젯젯젯음 으시젯젯기 으시젯젯냐기 으시젯젯다기 으시젯젯음 으시젯음 으시기 으시냐기 으시다기 으심 으시자기 (40) (Total 218)

Liste 2-9. Les formes particulières de SFN

Pour dénoter le *passé*, six syllabes sont disponibles. Il s'agit de 렸 (NPT8), 았 (NPT9), 었 (NPT10), 였 (NPT11), 왔 (NPT12), et 왔 (NPT13) dans la liste 2-10. Ces suffixes sont attachés après SFR.

Les verbes et les adjectifs peuvent être classés en fonction de la forme de leur SFR. Pour construire les Formes Spéciales de FN, nous devons savoir *quel SFR se combine avec quel SFN*. Nous devons donc avoir l'information sur la connexion entre SFR et le prédicat. Dans ce but, nous en avons fait une liste des SFR dans la figure 2-10. Dans cette figure, SFN est incluse dans PPP.

NPT8 :	렸겠기 렸겠냐기 렸겠다기 렸겠음 렸기 렸냐기 렸다기 렸었겠기 렸었겠냐기 렸었겠다기 렸었겠음 렸었기 렸었냐기 렸었다기 렸었음 렸음(16)
NPT9 :	았겠기 았겠냐기 았겠다기 았겠음 았기 았냐기 았다기 았었겠기 았었겠냐기 았었겠다기 았었겠음 았었기 았었냐기 았었다기 았었음 았음(16)
NPT10:	었겠기 었겠냐기 었겠다기 었겠음 었기 었냐기 었다기 었었겠기 었었겠냐기 었었겠다기 었었겠음 었었기 었었냐기 었었다기 었었음 었음(16)
NPT11:	였겠냐기 였겠다기 였기 였냐기 였다기 였었겠냐기 였었겠다기 였었기 였었냐기 였었다기 였었음 였음(12)
NPT12:	왔겠기 왔겠냐기 왔겠다기 왔겠음 왔기 왔냐기 왔다기 왔었겠기 왔었겠냐기 왔었겠다기 왔었겠음 왔었기 왔었냐기 왔었다기 왔었음 왔음(16)
NPT13:	웠겠기 왔겠냐기 왔겠다기 왔겠음 왔기 왔냐기 왔다기 왔었겠기 왔었겠냐기 왔었겠다기 왔었겠음 왔었기 왔었냐기 왔었다기 왔었음 왔음(16)

Liste 2-10. Les formes du *passé* pour SFN.

Classe 4	SFR	Type PPP
듣다	듣	SU1-5 6
		SU2-1
		SU3-1 2 4 5 7 8 9 10 12 13 14 15 17 21 23 24
		SU4-1
		PV1-1 3
		PV4-1
	들	SU2-2 3 4 5
		PV4-1

Figure 2-10. Exemple d'un verbe ‘듣다’ du dictionnaire SFR.

La figure 2-10 est un exemple de (e.g. ‘듣다’(écouter)). Le dictionnaire comporte l'information sur la connexion entre le verbe et le PPP. Le verbe a deux formes de SFR “듣” et “들”. Dans le cas de ‘듣’, il a l'information incluant SU3-1, 2, 4, ... dans la zone **Type PPP**. Cela signifie que les données de SU3-1, 2, 4, ... peuvent être soudées après SFR ‘듣’.

Un autre exemple, celui de ‘졌기’, a l’information SU3-4 dans la zone **Type PPP** de la figure 2-11. Par conséquent, SFR “듣” peut être combinée avec le suffixe “졌기”. Le résultat “들겠기” est marqué comme FN.

Il existe une autre structure qui correspond à l’*infinitif* en français. C’est celle de la forme *modifieur-”것”* (suffixe complémentaire) en coréen. Elle comporte plusieurs formes fléchies selon le temps. Nous appelons ce type de modifieur **PPP modificatif**. PPP modifieur est un sous-ensemble de PPP.

SFN entré	PDD ¹	de V/A	PPN Type	Type PPP
졌기	SFN	V et A	PPN	SU3-4-NM
졌냐기	SFN	V et A	PPN -PN1	SU3-4-ID
졌다기	SFN	V et A	PPN -PN1	SU3-4-ID
졌음	SFN	V et A	PPN	SU3-4-NM
기	SFN	V et A	PPN	SU3-7-NM
냐기	SFN	V et A	PPN -PN1	SU3-9-ID
는데	SFN	V	PPN	SU3-13-CJ-OV
다기	SFN	V et A	PPN -PN1	SU3-14-ID
렸겠기	SFN	V et A	PPN	PP7-2-NM
렸겠냐기	SFN	V et A	PPN -PN1	PP7-2-ID

Figure 2-11. L’exemple des SFN.

¹ PDD : parties du discours

2.2.4. Procédures de traitement des ambiguïtés

2.2.4.1. Ambiguïté P1

Nous avons déjà discuté la différence entre NS et FN. L'ambiguïté **P1** n'est pas importante au niveau morphologique. La différence d'information a lieu aux niveaux syntaxique ou pragmatique. Dans la liste 2-11 les formes d'ambiguïté **P1** sont énumérées, en vue de l'analyse syntaxique ultérieure.

가늌 가뭇 간지럼 같기 감 감기 개기 겹 겹침 결함 고기 고름 고향 곱 괴로움 구름
귀염 그리움 금 기름 기울기 기함 김 깨달음 꿈 끄스름 남 내기 노름 노여움 녹음
놀음 담 대감 대기 대함 댐 뎀 도움 들기 등침 두려움 뒤집기 딸기 땀 땀 땀 매기 땀
매기 명함 모기 몸 무서움 미끄럼 미움 밀림 바람 반가움 배기 뱀 범 봄 봉함
부끄러움 부끄럼 부럼 부음 부품 사기 살기 살림 삶 삼 샘 서기 서러움 석기 섬
세기 썸 속기 슬기 시기 식기 식음 신기 신음 실기 얇 어두움 어둠 어름 어름
어지러움 어지럼 얼뜨기 얼뜨기 영름 연임 염 오금 오기 움 움 이기 이름 이음 임
자기 잡기 재기 잦 전함 절기 점 조림 줌 주기 줄기 지겨움 지기 지름 지침 직함
집기 째 차기 참 추기 춤 치기 침 탐 태기 톱 파기 패기 폐함 품 함 회기(143)

Liste 2-11. Les ambiguïtés **P1** entre NS et FN.

감 겹 곱 금 남 담 범 심 엄 염 점 줌 참 품 함(15)

Liste 2-12. Les ambiguïtés **P2** entre PF et FN.

감 겹 금 남 내기 담 대기 데기 때기 뜨기 땀 바람 배기 범 부럼 빔 뺨기 심 싸대기
썸 염 임 잠 점 지기 짐 짜기 참 춤 치기 침 태기 패기 품 함(35)

Liste 2-13. Les ambiguïtés **P1** entre SF et FN.

Les listes 2-12 et 2-13 figurent dans la liste des ambiguïtés **P1**. SF et PF sont utilisés facultativement, mais FN est obligatoire dans la structure T2. Nous pouvons utiliser cette information pour lever l'ambiguïté.

Nous avons choisi en priorité le FN dans le cas d'ambiguïté **P1** avec SF et PF.

2.2.4.2. Ambiguïté P2

FN peut provoquer l'ambiguïté **P2** entre PF, SF, et PPN. La liste 2-14 est la liste des candidats localisés dans la sphère 2 ayant l'ambiguïté **P2** entre PF et FN. Nous n'avons pas découvert les données susceptibles d'appartenir à la sphère 4 sous T2. Autrement dit, il n'y a pas d'ambiguïté **P2** dans FN. Les variables x, y, et z dans la liste sont expliquées dans la formule **P2**.

강남(xy)	국적이기 국적임 기 기기 김 다르기 다름 루하기 루함 부끄러움 부끄럼 부끄럽기 부럽잖기 부럽잖음 상거리기 상거림 상남상하기 상남상함 성적이기 성적임 실거리기 실거림 실남실하기 실남실함 아넘치기 아넘침 아돌기 아돌아가기 아돌아감 아돌음 아돔 음 짓하기 (z) (33)
보금(xy)	실금실하기 실금실함 처놓기 하기 함 (z) (5)
어금(xy)	실금실하기 실금실함 처놓기 하기 함 (z) (5)
징검(xy)	기 누래지기 누래짐 누럼 누렁기 누르기 누름 디검기 디검음 뜯기 뜯기기 뜯김 뜯음 붉기 붉음 붉히기 붉힘 뿌염 뿌영기 소하기 소함 송검송하기 송검송함 송하기 송함 실거리기 실거림 실검실하기 실검실함 음 적검적하기 적검적함 치기 침 퍼래지기 퍼래짐 퍼럼 퍼렁기 푸르기 푸르접접하기 푸르접접함 푸르죽죽하기 푸르죽죽함 푸름 (z) (44)

Liste 2-14. Les candidats à l'ambiguïté P2 dans FN.

2.3. Structure de type 3

Le type 3 concerne le prédicat à PPP dans l'interprétation traditionnelle de la phrase, la prédicat joue le rôle principal dans la structure de la phrase. Les entrées prédictives relatifs aux éléments lexicaux de PPP peuvent être attachées à [10] [28]. En français, il y a trois types des séquences prédictives. PPP est attaché seulement au verbe (la deuxième colonne la table 2-4 est le cas du verbe ordinaire, la troisième colonne horizontale est pour la copule, et la quatrième est lexicalement vide mais il est grammaticalement le verbe) [28].

En coréen, nous avons observé trois types de séquences prédictives. PPP n'est pas seulement attaché aux verbes mais aussi aux adjectifs et aux unités grammaticales liées (GL) aux noms.

PDD	Type français	Coréen	Ex. en français	Ex. Coréen
verbes	N V-PPP	N V-PPP	Jean dort	존은 잔다
adjectif	N Etre-PPPAdj	N A-PPP	Jean est grand	존은 크다
nom prédicatif	N V-PPP Npréd	N N-GL- PPP	Jean fait une plaisanterie (= Jean plaisante)	존은 농담한다

Table 2-4. L'exemple des types des mots prédictifs.

A partir de ces trois types, nous avons défini la structure T3. Les mots qui contiennent un prédicat et un SF sont classés comme T3.

$$\begin{aligned}
 T3 &= (PF^2(NS | FN)^+ SF^?)^* (V_{root} | A_{root}) \\
 &\quad \text{(Postposition verbale | Postposition adjectivale)} \\
 &= (ND | PF^2 FN^+ SF^?)^* (V_{root} | A_{root}) PPP \\
 &= (ND | PF^2 FN^+ SF^?)^* FFP
 \end{aligned}$$

T3 est constitué de trois parties. La première concerne le nom prédictif. Le nom prédictif est le nom qui autorise l'attachement du verbe. Il y a d'autres formes à la place des noms de la première partie, comme avec les quatre types suivants :

- 1) “물컹하다” (être mou) est formé de l'adverbe “물컹” (mou) et du prédicat (dans ce cas, l'adjectif) “하다”(faire). Pour “하다”, nous avons observé que les adverbes qui peuvent être combinés avec “하다” afin de former un adjectif n'acceptent pas le suffixe verbal “하다”. Par contre, ces adverbes peuvent être combinés avec le suffixe verbal “거리다”. Ainsi une verbe comme “물컹거리다” peut être produit. Ce mécanisme est observé très fréquemment. Cependant, même si la procédure est très productive, elle n'est pas applicable à tous les adverbes qui entrent dans l'une de ces relations dérivationnelles.

- 2) Certains prédicats peuvent être attachés à la tête de l'autre prédicat. Par exemple, “놀고먹다” (manger en s’amusant) est composé des deux verbes “놀고” (s’amuser) et “먹다” (manger).
- 3) Les deux autres formes complexes du prédicat telles que préfixe-prédicat et prédicat-suffixe peuvent être considérées. Pour la composition de préfixe-prédicat, par exemple “엇몰리다” est composé du préfixe “엇” (obliquement) et du verbe “몰리다” (lier).
- 4) La forme prédicat-suffixe est dérivée du verbe simple par l’ajout des affixes. Le suffixe peut être un morphème grammatical qui transforme *le verbe transitif* au *passif*, ou *le verbe intransitif* à la forme *causative*. “박히다” (être enfoncé) est un exemple du premier cas. La forme ‘enfoncer-suffixe.pass’ transforme en ‘être enfoncé’. “굶기다” représente l’exemple du second cas. La forme ‘être affamé-suffixe.caus’ se transforme en ‘affamer’.

Nous avons enregistré les quatre types complexes, ci-dessus, dans notre lexique avec la marque VC (Verbe Complexe).

La deuxième partie concerne la radical des prédicats. Cette partie ne connaît pas de variations. Pour ajouter la partie finale de T3 à la deuxième partie, il nous faudra définir un nouveau concept appelé SFR. C’est la dernière syllabe du radical. Pour attacher la dernière partie à la fin du radical, il est nécessaire de connaître l’information sur le type de SFR. La forme PPP de la dernière partie dépend de SFR. SFR joue le rôle du connecteur de l’interface entre le radical (deuxième partie) et le PPP (dernière partie) dans T3.

SFR permet de classer 8,763 verbes en plusieurs sous-classes. Si deux radicaux possèdent le même caractère final, elles ont généralement le même type de SFR.

La dernière partie est sur PPP. PPP ne peut exister seul. Il a toujours besoin du verbe ou de l’adjectif. Les verbes et les adjectifs utilisent environ 6,000 formes de PPP dans les mots réels [26]. Nous donnerons l’information exacte sur la connexion entre la partie du prédicat et celle de PPP.

2.3.1. Noms Prédicatifs

Le nom prédicatif est un nom qui peut s'attacher du prédicat. Par exemple, “공부하다” (étudier) est composé du nom (N) “공부” (étude) et du verbe (V) “하다” (faire). Ce mot peut être transformé en deux formes disjointes N-PPN et V, comme “공부를 하다” (N-PPN V). Ces deux unités disjointes sont des éléments syntaxiques, car on peut insérer d'autres éléments grammaticaux entre elles. C'est pourquoi nous ne considérons pas cette forme comme un verbe simple. Par contre, nous n'avons pas observé deux unités disjointes du point de vue syntaxique, quand il s'agit d'une forme basée sur l'adjectif [28] [29].

Voici un exemple de phrase à un verbe support ‘주다’ (*juda*) [donner] :

- (1) 민우는 인아에게 도움을 주었다
 Minu-PPN Ina-PPN aide-PPN donner-PPP
 (Minu a donné de l'aide à Ina)

Cette phrase correspond à la construction à verbe distributionnel ‘돕다’ (*dobda*) [aider], lexicalement relié au substantif ‘도움’ (*doum*) [aide] :

- (2) = 민우는 인아를 도왔다
 Minu-PPN Ina-PPN aider-PPP
 (Minu a aidé Ina)

On peut décrire cette relation comme suit [14] [31] :

- (3a) *N0 N1 Npréd Vsup*
 (3b) = *N0 N1 V*

La relation (3) s'observe quand il existe un verbe morphologiquement relié au substantif prédicatif. Cependant, avec le verbe ‘하다’ (*hada*) [faire], on a systématiquement ce type de paires de phrases. Ainsi, la phrase suivante [18] :

- (4) 민우가 산책을 한다
 Minu-PPN promenade-PPN faire-PPP
 (Minu fait une promenade)

est en relation transformationnelle avec :

- (5) = 민우가 산책한다
 Minu-PPN se promener-PPP
 (Minu se promène)

Or, la relation entre la phrase à support *hada* et celle à verbe distributionnel correspondant est particulière du point de vue morphologique : l'une contient une séquence Nom-PPN *hada*, et l'autre Nom-*hada*. Autrement dit, le changement de forme

entre ces deux types de prédicats est minimal, il se limite à la présence d'une postposition de l'accusatif. Par ailleurs, cette transformation est permise, a priori, dans toutes les phrases à support *hada*, sans exception

Dans des dictionnaires actuels, les formes Nom-*hada* dont les substantifs (Nom) sont considérés comme prédicatifs sont intégrées sous une entrée de verbe dérivé. Par conséquent, on compte d'une part ces substantifs dans le lexique des substantifs, et d'autre part les verbes basés sur ces substantifs (Nom-*hada*) dans le lexique des verbes. Cette situation est certainement gênante quand les noms de ce type sont nombreux et surtout quand leur liste n'a pas été établie par un procédé systématique. On risquerait de dédoubler le lexique partiellement avec les noms associables à *hada*.

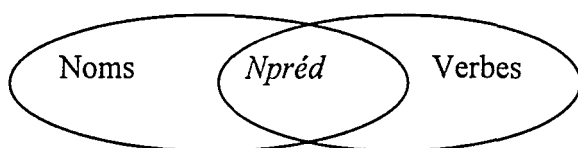


Figure 2-12. Position des Noms Prédicatifs

Il y a des noms qui ne fonctionnent pas comme nom prédicatif. Mais s'ils sont accompagnés d'un suffixe “질” (acte) ou “화” (devenir), les formes N-SF-Prédicat deviennent acceptables, comme ci-dessous :

- (1) 가위-질-하다 (ciseaux- acte -faire) (N-SF- Prédicat)
- (2) 세계-화-하다 (monde-devenir -faire) (N-SF- Prédicat).

Nous possédons le dictionnaire des noms prédicatifs qui peuvent être accompagnés du verbe ‘하다’ (faire) ou ‘되다’ (devenir), mais, il n'existe ni dictionnaire de tous les noms prédicatifs ni liste complète des prédicats qui peuvent être attachés aux noms prédicatifs.

가다 거리다 끼치다 나다 내다 당하다 되다 떨어다 미치다 받다 부리다 시키다 입다
주다 지다 짓다 치다 피우다 하다

Liste 2-15. Liste des verbes utilisés fréquemment avec les noms prédicatifs.

가지다 같다 아니다 없다 이다 있다(6)

Liste 2-16. Liste des adjectifs utilisés fréquemment avec les noms prédicatifs.

Comme le dictionnaire complet des noms prédicatifs n'existe pas, nous utilisons la règle T3. Les listes 2-15 et 2-16 contiennent des prédicats très productifs, qui, entre autres, se combinent avec le nom prédicatif Nom-*hada* sous T3. Mais on n'a pas observé de restrictions sur les prédicats.

Dans le cas de l'adjectif, les noms prédicatifs sont ajoutés avec des restrictions à la tête de “하다” (faire). Même si la liste est longue, nous l'avons construite par un recensement linguistique des données [26]. En effet, l'information est décrite dans le dictionnaire des adjectifs. Celui-ci possède 3,231 entrées dont la forme est composée d'un Nom Prédicatif et de “하다”.

2.3.2. Verbes

Notre système comporte une liste de 8,763 verbes simples ou complexes. Comme nous l'avons déjà remarqué, les verbes complexes peuvent avoir quatre types de marque, comme par exemple VC dans le dictionnaire des verbes.

Il existe 80 verbes simples qui sont formés d'éléments non-autonomes et du suffixe “하다”, comme “택하다”. Tous ces éléments non-autonomes sont sino-coréens à l'origine.

VERBE	PDD	Information
가까와지다	VS	/INT /JMA
가꾸다	VS	/TRA
가닐거리다	VS	/INT /ADV KM
가동거리다	VS	/TRA /INT /RKM
가꾸로박히다	VC	/INT /PS
가라앉히다	VC	/TRA /CS
나자빠지다	VC	/INT /JM

Figure 2-13. Extrait de la liste des verbes simples.

Dans la figure 2-13 sont donnés quelque exemples de la liste des verbes. Pour établir l'information de la connexion avec PPP, nous classerons les verbes en plusieurs groupes au paragraphe 2.3.4.

2.3.3. Adjectifs

La catégorie Adjectif représente les formes prédicatives qui se conjuguent au moyen de certains suffixes grammaticaux, de façon analogue au verbe. Dans notre système, il y a 5,250 adjectifs et ils sont classés en fonction de plusieurs formes finales. Par exemple, la classe CM du dictionnaire indique l'adjectif avec le suffixe “이다” (être). Les adjectifs entrant dans cette classe connaissent 800 formes flexionnelles environ. Les formes conjuguées permettent nécessairement de reconnaître les entrées lexicales. La figure suivante est un extrait de la liste des adjectifs [23].

Adjectif	PDD	Information
가공적이다	ADJ	/CM
가깝다	ADJ	/RM
가깝디가깝다	ADJ	/RM
가깝하다	ADJ	/HM
가난하다	ADJ	/HM

Figure 2-14. Exemples d'adjectifs.

Pour établir l'information de la connexion avec PPP, nous allons classer les adjectifs en plusieurs groupes dans ce qui suit.

2.3.4. Syllabe Finale du Radical

Quel prédicat peut être combiné avec quelle série flexionnelle ? Supposons que les combinaisons soient sans restriction. Le nombre N de combinaisons possibles (Prédicat + PPP) sera le suivant:

$$\begin{aligned} N &= \text{Nombre des verbes} * \text{Nombre de Postposition verbale} \\ &\quad + \text{Nombre des adjectifs} * \text{Nombre de Postposition adjectivale} \\ &= 8,763 * 20,853 + 5,250 * 18,893 = 281,923,089. \end{aligned}$$

Dans le cas normal, un prédicat peut être combiné avec environ 6,000 PPP [26]. Le nombre estimé des formes combinées est le suivant:

$$\begin{aligned} \text{Le nombre estimé} &= \text{Nombre des verbes} * 6,000 + \text{Nombre des adjectifs} * 6,000 \\ &= 8,763 * 6,000 + 5,250 * 6,000 = 84,078,000. \end{aligned}$$

Pour extraire de manière exacte environ 84,078,000 formes de la combinatoire entre les prédicats et PPP, nous devons examiner 281,923,089 données. Le nombre est trop énorme pour une recherche manuelle.

Il est indispensable d'appliquer une méthode semi-automatique pour sélectionner les formes correctes FFP. Après avoir classé les prédicats et les PPP, nous fournirons l'information de connexion entre le groupe des prédicats et celui des PPP. La classification des SF sera discutée dans le paragraphe suivant.

Pour classer les prédicats, le concept de SFR est indispensable. Postposition verbale (PPV) et Postposition adjectivale (PPA) s'attachent après la radical du verbe ou de l'adjectif. Si deux prédicats comportent la même SFR avec le même sens, ils auront le même PPP. Par exemple, avec les deux verbes “나가다” (sortir) et “가져가다” (emporter). Les radicaux de ces verbes sont “나가” et “가져가”. Ils ont la même SFR “가”. Ils ont également les mêmes formes fléchies, comme dans la figure 2-15.

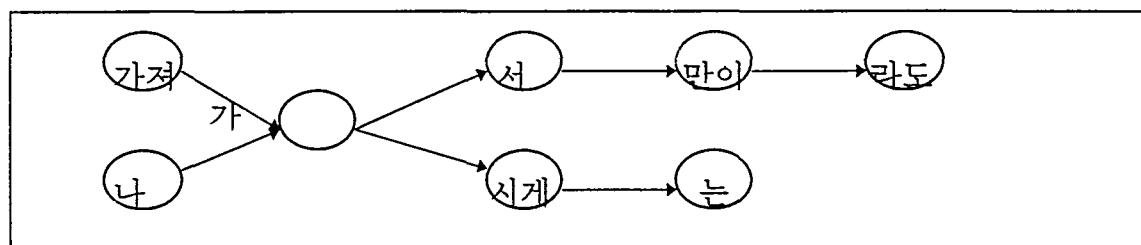


Figure 2-15. L'exemple pour les deux verbes ayant la même SFR.

SFR précède PPP, parce que SFR fait partie du prédicat, et le prédicat précède PPP dans la structure T3. Il se réfère toujours à la forme de PPP. Dans certains cas, une compression entre SFR et la première syllabe du PPP peut avoir lieu, et dans d'autre cas, la dernière consonne du SFR peut être supprimée ou changée.

Les figures 2-16-1 et 2-16-2 donnent des exemples de formes fléchies. La différence entre ces deux figures porte sur un problème de code du coréen. Actuellement on utilise plusieurs systèmes de code [3] [32]. Cela provoque des représentations différentes et des analyses différentes pour les formes de SFR.

Nous avons décidé d'utiliser le système de code à 2 octets (appelé **KSC5601 WANSUNG Code**) comme code standard du coréen pour représenter les syllabes [33]. Le concept est hérité du système **Unicode**. Il est difficile de représenter l'unité phonémique en coréen avec ce dernier car le système est basé sur la syllabe.

La figure 2-16-1 comporte les valeurs de phonème “ㄴ”, “ㄷ”, et “ㅈ”. Le premier noeud “가” et les phonèmes dans la figure 2-16-1 ne peuvent pas se combiner naturellement, parce que le code standard n'a pas de règle pour la combinaison entre une syllabe et un phonème. Selon les règles d'écriture du coréen, la syllabe “가” et le phonème “ㄴ” doivent produire la syllabe “간”. C'est impossible dans le code standard du coréen sans programme particulier de composition. Dans d'autres systèmes de code du coréen, ces formes se combinent sans difficulté. C'est-à-dire que la syllabe “간” contient l'information sur la syllabe “가” et sur le phonème “ㄴ” dans le système du code.

La présentation de code **WANSUNG** de la figure 2-16-1 est décrite dans la figure 2-16-2. Même si la forme est différente de celle de la figure 2-16-1, le contenu est identique. On peut reconstituer les formes de la figure 2-16-2 au moyen de celles de la figure 2-16-1. Par exemple, le premier noeud “가” et l'élément phonémique “ㄴ” du noeud “ㄴ지” de la figure 2-16-1 sont combinés linguistiquement, et le résultat “간” est inscrit dans le nouveau noeud de la figure 2-16-2. Nous avons appliqué le même principe à tous les éléments phonémiques dans la figure 2-16-1 et nous avons obtenu la figure 2-16-2. Dans la figure 2-16-1, il n'y a qu'un seul noeud de SFR “가”, mais dans l'autre, il y a les noeuds de SFR “간”, “갈” et “갸”.

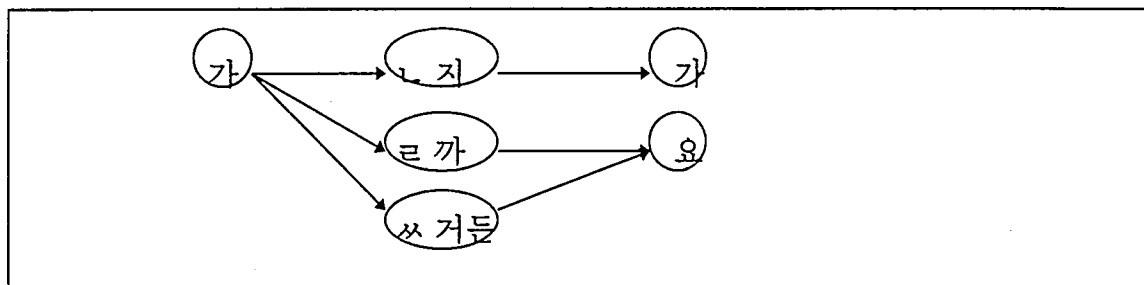


Figure 2-16-1. Une SFR et PPPs avec l'information phonémique.

Par conséquent, dans la figure 2-16-1, il y a un nœud SFR “가”, mais dans la figure 2-16-2, il y a trois nœuds SFR. Si on ne tient pas compte de la notation et du système de code, il n'y a pas de différence entre les figures 2-16-1 et 2-16-2.

```

graph LR
    A1((간)) --> B1(지)
    B1 --> C1((자))
    A2((갈)) --> B2(까)
    B2 --> C2((요))
    A3((값)) --> B3(저든)
    B3 --> C3((요))
  
```

Nous avons énumérés 333 types de SFR pour les verbes et 151 types de SFR pour les adjectifs. Les listes 2-17 et 2-18 donnent les SFR respectivement compatibles avec les verbes et avec les adjectifs. Nous remarquons que le SFR “크” que l'on trouve dans les verbes et dans les adjectifs a des variations particulières. Dans ce cas, le nombre des SFR est 2. Le détail est décrit dans les classes de SFR. Supposons que deux formes différentes de SFR aient le même PPP. Les deux SFR entrent dans la même classe de SFR.

Liste 2-17. Liste des SFR apparaissant dans les verbes.

Liste 2-18. Liste des SFR apparaissant dans les adjectifs.

Avec cette méthode, nous avons classé des SFR dans la liste en petits groupes avec l'information de PPP. Pour le cas du verbe, les SFR et l'information sur leur PPP sont décrites dans l'Annexe E, l'adjectif, dans l'Annexe F. Un extrait de la liste avec l'information sur PPP ont été donné dans le paragraphe 2.2.

2.3.5. Postpositions Prédicatives

L'information sur les formes fléchies des noms et les PPP des prédicats doit être indiquée dans le dictionnaire, car le dictionnaire contenant seulement les entrées canoniques ne permettrait pas de reconnaître les formes conjuguées telles que 'goes', 'went' et 'gone' pour le verbe 'go' en anglais.

Dans le cas du français, des informations plus complexes sont exigées. Un nom masculin peut avoir les formes de féminin, masculin pluriel, et féminin pluriel, comme *chanteur* (singer:m:sing), *chanteuse* (singer:f:sing), *chanteurs* (singer:m:pl), et *chanteuses* (singer:f:pl). Un adjectif peut aussi avoir plusieurs formes fléchies. Un verbe apparaît rarement sous sa forme canonique. Il est plus souvent conjugué, par exemple, *vais*, *vas*, *va*, *allons*, *allez*, *vont*, *allais* et *allait* pour le verbe 'aller'. Le dictionnaire doit donc fournir l'information sur toutes les formes conjuguées.

Examinons maintenant le cas du coréen. Il n'existe ni de genre pour le nom ni de forme conjuguée pour le verbe. Par contre, il existe des marques de fonction grammaticale, comme les marques de *temps*, de *modalité*, d'*aspect*, de *politesse* et de *mode de la phrase*. L'ordre et les contraintes de combinaison sont extrêmement complexes. Les adjectifs en coréen doivent être suivis de PPP. Les suffixes indiquent également les fonctions grammaticales des adjectifs, alors qu'en français et en anglais, c'est une copule, comme 'être' ou 'be' ou des verbes équivalents, qui prennent les marqueurs indiquant les fonctions grammaticales des mots adjectivaux.

Il était évident qu'un analyseur morphologique du coréen ne peut reconnaître les formes fléchies des noms, des verbes et des adjectifs sans l'information associée. Par exemple, "갔다" (FFP) (est allé) est la forme du passé de "가다" (aller). L'entrée est composée du radical et des suffixes grammaticaux. Donc, il faut avoir le dictionnaire qui donne l'information sur les formes fléchies.

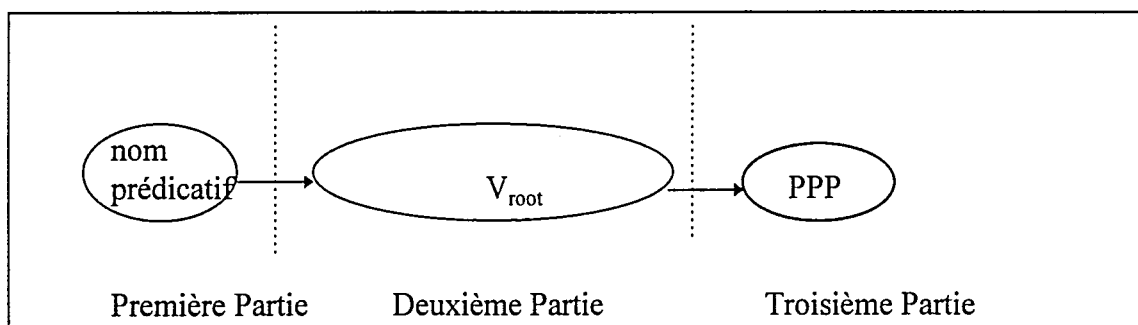


Figure 2-17. Forme de la structure T3.

FFP inclut toutes les formes combinées du radical avec PPP dans la figure 2-17. SFR occupe la fin de la deuxième partie. Pour construire FFP, nous construirons l'information de la connexion entre SFR et PPP, puisque la troisième partie dépend du SFR.

Nous avons classé les entrées PPP à l'aide de l'information de connexion avec le SFR dans la table 2-5.

Par exemple, les données dans la troisième colonne horizontale dans la table 2-5 sont classées comme type SU3-1. Toutes les entrées entre l'entrée “거나” (dans la colonne ‘de’) et l'entrée “건만은” (dans la colonne ‘à’) dans le dictionnaire de PPP selon l'ordre lexical sont classées comme types SU3-1. Pour comprendre les types des classes, examinons l'extrait de la liste des PPP de la figure 2-16.

de	à	Type	de	à	Type
ㄴ가	ㄴ지조차도	SU1-1	ㄹ것같으면	ㄹ텐데도	SU1-2
ㅁ세		SU1-3	ㅂ니까	ㅂ시다	SU1-4
거나	건만은	SU3-1	게	게조차도	SU3-2
젠들		SU3-3	겠거나	겠지요	SU3-4
고	고저	SU3-5	구나	구면	SU3-6
기	길레	SU3-7	나		SU3-8 PA1-1 SU1-5
나니	나니라	SU3-8	나마	나마나	PA1-1 SU1-5
나이다		SU3-8	냐		SU1-6
냐거나	냐하	SU3-9	난	난듯이	SU3-10
날	날텐데도	SU3-11	네	느라고	SU3-12
너라		PV1-2 /* V */	는	는지조차 도	PA1-3 SU3-13 /* V */
니	니이까	SU1-6	다	다하	SU3-14
단	단듯이	SU3-15	달	달텐데도	SU3-16
대요		SU3-17	더구나	든지	SU3-17
들		PA1-4	듯이	디	SU3-18
라		PA1-5	라거나	란듯이	SU1-7
랄	랴	SU1-8	러	러조차도	PA7-1 SU1-9 /* V */
렸거나	렸지요	PA7-2	려	리오	SU1-9
마		SU1-10	만	만이라도	PA1-6
매	프로	SU1-11	사오니	삽고	SU3-19 /*V*/
서	서야말로	PA1-7	세		SU3-20
셔	셨지요	SU1-12	소		SU3-21
소서		SU1-13	쇼		SU1-14
습니까		SU4-1	습니까요		SU2-1
습니다		SU4-1	시거나	시지요	SU1-15
신	신지조차도	SU1-16	실	실텐데도	SU1-17
십니까	십시오	SU1-18	아	았지요	PA3-1
야	야만이	PA1-8	어	었지요	PA4-1
여	였지요	PA5-1	오		SU1-19

오니		SU3-22	와	왔지요	PA8-1
요		PA1-9	으나	우이	PA6-1
운	은지조차도	PA6-2	홀	을텐데도	PA6-3
음세		PA6-4	읍니까요	읍시다	PA6-5
워	웠지요	PA6-6	으나	으이	SU2-2
은	은지조차도	SU2-3	을	을텐데도	SU2-4
음세	읍니다	SU2-5	이		SU1-20
자	잔듯이	SU3-23	잘	잘텐데도	SU3-24
잡고		SU3-25 /* V */	지	지요	SU3-26

Table 2-5. Classification des PPP.

Entrée	PDD	derivé de V/A	Type PPP
거나	PPP	V et A	SU3 .Conj
는구료	PPP	V	SU3 .TmDec
서야말로	PPP	V et A	PV1 .Conj
시어요?	PPP	V et A	SU1 .TmInt

Figure 2-18. Extrait de la liste des PPP.

Les significations de ‘SU1’ et ‘TmInt’ de la dernière entrée “시어요?” dans la figure 2-18 sont les suivantes [30]:

- 1) L’entrée peut être combinée avec les prédicats qui finissent par une voyelle ‘*claire*’.
- 2) L’entrée peut être combinée avec les prédicats qui se terminent par une voyelle ‘*obscure*’.
- 3) L’entrée peut être combinée avec les prédicats qui se terminent par une voyelle ‘*obscure*’ et par une consonne.
- 4) L’entrée peut être combinée avec les prédicats qui finissent par “*니다*”.
- 5) L’entrée est le suffixe terminal au mode *interrogatif*.

FFP est construit facilement avec le dictionnaire PPP et l’information SFR (l’extrait de la liste est donné dans la figure 2-10), parce que PPP et SFR contiennent la même colonne verticale “**Type PPP**” qui connecte les deux.

3. Organisation du dictionnaire

Dans ce chapitre, nous discuterons des techniques de consultation du dictionnaire. Une méthode est celle de l'Automate Fini (AF) [9] [34] [38] qui est la technique appliquée au DELAS et au DELAF du LADL [4] [5] [6] [7] [11], et à Lexique-Compiler de XEROX [16]. L'autre technique est celle qui sera développée dans notre étude. Dans l'état actuel de cette technique, l'utilisation d'automates présente beaucoup d'avantages, alors que l'autre technique que nous appellerons la Consultation Basée sur la Position (CBP) est très compétitive. La méthode CBP a été programmée sur une machine **HP**.

Le technique d'automates ressemble beaucoup à celle des Trie structures [1] [2] [20]. La complexité du temps des deux systèmes est du même ordre. La vitesse de consultation dans les deux structures dépend uniquement de la longueur de la séquence, pas du volume du dictionnaire. L'automate a l'avantage d'occuper moins de place en mémoire que dans la méthode Trie structure [12] [13].

Le Trie représente les entrées du dictionnaire 'ab', 'abc', 'ade', et 'afg' comme dans la figure 3-1-1. Dans ce cas, l'automate fini représente les entrées comme dans la figure 3-1-2. Il y a aussi un autre exemple pour le cas des séquences 'abc' et 'efc'. Le Trie structure est décrit dans la figure 3-2-1, mais l'automate est représenté dans la figure 3-2-2. La différence est nette. Le Trie a besoin d'un plus grand nombre de noeuds que l'automate fini.

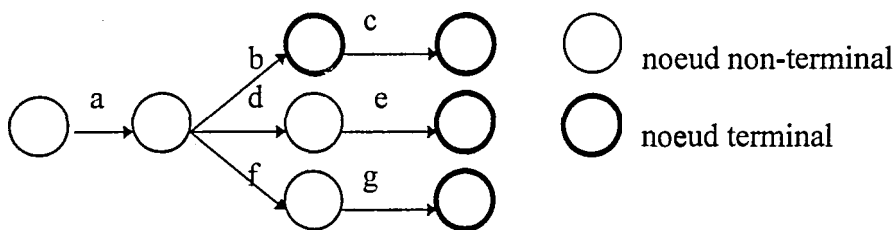


Figure 3-1-1. Diagramme de transition dans Trie structure.

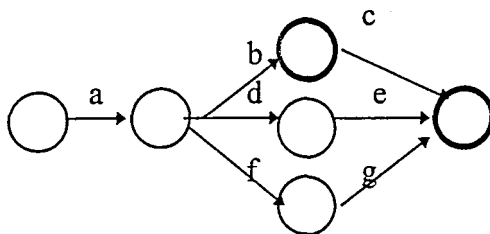


Figure 3-1-2. Diagramme de transition dans Automate Fini.

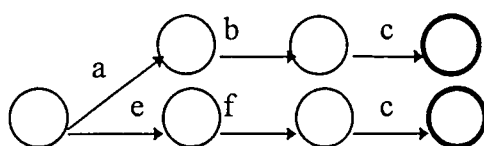


Figure 3-2-1. Représentation par Trie structure pour les séquences 'abc' et 'efc'.

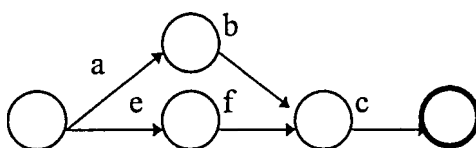


Figure 3-2-2. Représentation par Automate pour les séquences 'abc' et 'efc'.

Dans AF, le caractère 'c' commun en fin de séquence dans la figure 3-2-2 est compressé à la différence de Trie dans la figure 3-2-1. L'Automate permet de compresser les caractères identiques des séquences dans les deux directions. L'Automate est donc plus efficace que Trie structure du point de vue de l'espace mémoire occupé.

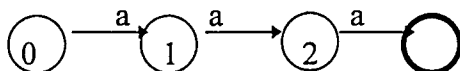

```
while(_gets(input, file_pointer) != NULL) {  
    if(_flm(input) != 0) _sep();  
    _llm(input);  
    if(input != NULL) _create();  
}
```

3.1.2. Exemples de procédures



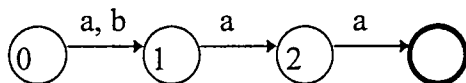
/* **input** : NULL, **current_nodes** : start $\rightarrow$ 0, terminal $\rightarrow$ t */

Figure 3-3. L'état initial.



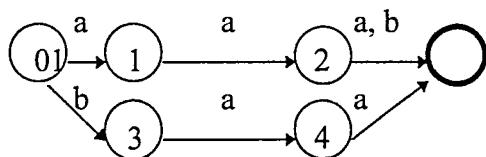
Procédure : /* **input** : "aaa", **current_nodes** : 0, t */
 _flm(input) 0
 _llm(input) 0
 _create(current_nodes, input)

Figure 3-4. Diagramme de transition après l'insertion de la séquence 'aaa'.



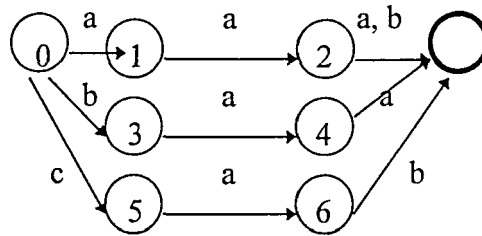
Procédure : /* **input** : "baa", **current_nodes** : 0, t */
 _flm(input) 0
 _llm(input) 1, "aa" /* **input** : "b", **current_nodes** : 0, 1 */
 _create(current_nodes, input)

Figure 3-5. Diagramme de transition après l'insertion de la séquence 'baa'.



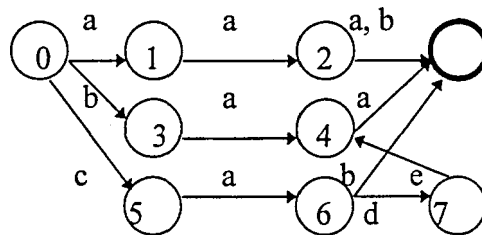
Procédure : /* **input** : "aab", **current_nodes** : 0, t */
 _flm(input) 2, "aa" /* **input** : "b", **current_nodes** : 2, t */
 _sep(0, 2);
 _llm(input) 0
 _create(current_nodes, input)

Figure 3-6. Diagramme de transition après l'insertion de la séquence 'aab'.



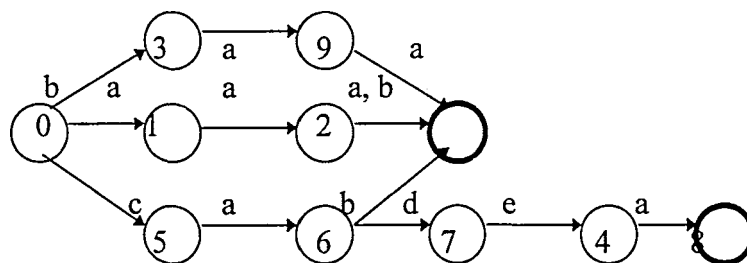
Procédure : /* input : "cab", current_nodes : 0, t */
 _flm(input) 0 /* input : "cab", current_nodes : 0, t */
 _llm(input) 0
 _create(current_nodes, input)

Figure 3-7. Diagramme de transition après l'insertion de la séquence "cab".



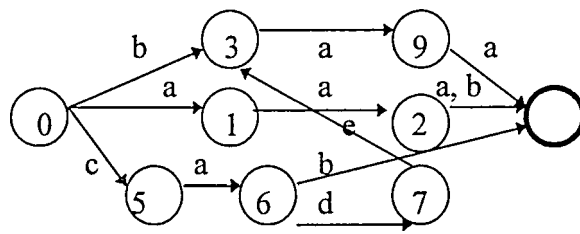
Procédure : /* input : "cadea", current_nodes : 0, t */
 _flm(input) 6, "ca" /* input : "dea", current_nodes : 6, t */
 _sep(0,6)
 _llm(input) 4, "a" /* input : "a" : current_nodes : 6, 4 */
 _create(current_nodes, input)

Figure 3-8. Diagramme de transition après l'insertion de la séquence "cadea".



Procédure : /* input : "cadeaa", current_nodes : 0, t */
 _flm(input) 8, "cadea" /* input : "a", current_nodes : 8, t */
 _sep(0,8)

Figure 3-9-1. Diagramme de transition après l'insertion de la séquence "cadeaa".



```

_llm(input) 3, "aa" /* input : NULL, current_nodes : 4, 3 */
_create(current_nodes, input)

```

Figure 3-9-2. Diagramme de transition après l'insertion de la séquence "cadeaa".

3.1.3. Algorithme pour l'automate fini

Pour construire la structure de l'automate, il est nécessaire de donner certaines définitions et présenter les algorithmes utilisés [13]. Ce sont des algorithmes pour la construction de l'Automate Lexical (qui sera défini ultérieurement). L'optimisation de la structure n'y est pas prise en compte. Ici nous ne discuterons que des types d'opérations nécessaires, c'est-à-dire, nous ne implémentons pas l'Automate lexical, mais nous ne donnons que des spécifications nécessaires pour l'Automate lexical.

Pré-hypothèse

Cet automate sera construit sur la base des hypothèses suivantes:

- Les entrées de l'input sont triées
- Durant l'exécution de la fonction `_sep`, il y a au plus un `nd` dans l'automate.

L'Automate Lexical est l'automate de la description du lexique. Il comporte les 5 structures suivantes.

L'Automate Lexical = (I, Q, P, Q₀, Q_t)

I = série des caractères anglais {a, b, c, ..., z}

Q = série des états, par exemple { 0, 1, 2, ... }

P(Q_i, I) -> Q_{i+1}

Q₀ = état de départ

Q_t = état terminal.

```
char* input;          /* input string, x ∈ input, x ∈ I */      [V]
int position = 0;      /* current position in input */          [V]
int x, y;              /* current start node and terminal node */
typedef int status;
int table[STATUS][ALPHABET] [V]
status top;            /* current_top position of the table */
```

```
_p(state, pos)
status state, char pos;
{
    return table[state][pos - 'a'];
}
```

```
_rp(state, pos)
status state, char pos;
{
    int i;
    for (i = 0; i < STATUS; i++)
        if (table[i][pos - 'a'] == state) return i;
    return 0;
}
```



```

_nd(Qi)
status Qi
{
    status next; static int i;
    for (next = Qi, until next < position, step next = _p(Qi, input[i++] ), do
        { if ∃a, b, (_p(Qn, a) == next) && (_p(Qn, b) == next), if a ≠ b ∈ I,
            then return TRUE; }
    return FALSE
}

```

```

_flm() {
status current, temp;
int i, j;
if(*input == NULL) return 0;
current ← temp ← Q0;
j = _strlen(input); i = position;
while((i < j) && (current ← _p(current, input[i++] ) != 0 &&
    (current != Q0)) temp ← current;
if(current == Q0) return 0;
x ← temp; position = --I; strcpy(input, subcat(input, 0, position));
if( x == Qt)
    x ← table[temp][input[--i] - 'a') ←
        _create_terminal(temp, input[i]); /* undefined_function */
/* _sep(Qs, temp); */
} /* end of the function _flm */

```

```

_llm() {
status current, temp;
int i, j;
i = position;
if(input == NULL) return 0;
current ← temp ← x;
j = _strlen(input); j--;
while((i < j) && (current ← _rp(current, input[j--]) != 0 &&
    (current != Qs)) temp ← current;
if(current == Qt) return 0;
y ← temp; position = ++j;
if(y == Qs) return 0;
if(!_dup(current, y)) {
    y ← current; _set_input(input, position, j);
}
} /* end of the function _llm */

```

```

_sep(start, end)
status start, end;
{
status current;

```

```

current = start;
while(current != end) {
    s'il y a un noeud courant à partir de deux autres noeuds
        then _create(start, current);
}

_create(start, end)
status start, end;
{
    int i, j;
    if(input == NULL) return 0;
    j = strlen(input);
    for(i = 1; i < j; i++) {
        table[start][input[i] - 'a'] = top;
        start = top++;
    }
    table[start][input[i] - 'a'] = end;
}

```

Figure 3-10. Pseudo algorithmes pour l'automate lexical.

3.2. Consultation Basée sur la Position

Pour une consultation efficace du dictionnaire, la structure optimisée du lexique est cruciale. CBP est basée sur la localisation du caractère dans le lexique. Les procédures de consultation ne sont pas basées sur les caractères mais sur les numéros des positions de ces caractères. Bien que cela nécessite plus d'états que AF, le temps et l'espace nécessaires sont aussi efficaces qu'avec AF.

L'automate a beaucoup de cellules vides [35] et pour réduire ces cellules vides, il exige la technique [35] [39]. Mais CBP n'a pas besoin de procédure de réduction comme dans le cas d'AF.

CBP n'a pas non plus besoin de re-construction comme dans la table de transition avec la méthode AF, bien que ce soit un travail **off-line(batch)** [35] [39]. Par exemple, l'input "aab" dans la figure 3-6 détruit la structure de la figure 3-5 et re-construit un nouveau diagramme comme dans la figure 3-6.

CBP est constituée seulement de deux tables; que l'on appelle Table de l'Index (TI) et Table de colonne Horizontalement (TH), avec des procédures liées.

Pour le dictionnaire de 'on-line', nous pouvons supposer que toutes les entrées sont classées, et que certains caractères identiques apparaissent dans les mêmes positions dans les entrées adjacentes. Cela peut nous indiquer qu'il est possible que l'opération oriente la position. La méthode supprime les mêmes caractères lorsqu'ils apparaissent sur la même position dans des entrées adjacents.

Elle est plus efficace que Trie structure. Dans le cas des inputs 'abc' et 'efc' dans la table 3-1 la structure est identique à AF. Autrement dit, CBP a une capacité de mémoire plus efficace que Trie, bien que CBP comprenne une description en table, mais la forme de la représentation est différente de AF. AF utilise la forme de la table 3-1. Chaque cellule de cette structure correspond à une transition de AF, ce qui implique que cette méthode est plus efficace que Trie structure. Chaque cellule de la table correspond aux noeuds de AF ou de Trie.

	Première position	Deuxième pos.	Troisième pos.
Première entrée: abc	a	b	c
Deuxième entrée: efc	e	f	

Table 3-1. Représentation de CBP pour les inputs "abc" et "efc".

Pour d'autres inputs comme "abcd" et "aded", cette méthode les représente comme la table 3-2.

	Première	Deuxième	Troisième	Quatrième
Première entrée: abcd	a	b	c	d
Deuxième entrée: aded		d	e	

Table 3-2. Représentation de CBP pour les inputs “abc” et “efc”.

3.2.1. Procédures de la construction CBP

La première étape de cette méthode est très simple. Si des caractères identiques apparaissent dans la même position que l'entrée précédente, le caractère spécial '#' s'y substitue. Grâce au caractère spécial, la position est réduite.

N° d'entrée	Entrée originale	1 ^{ère} transformation
0	aaron	aaron
1	aas	##s
2	aback	#back
3	abaculus	####ulus
4	abaddon	####don
5	abandon	###n###
6	abandoned	#####ed
7	abandonee	#####e
8	abase	###se
9	abash	####h
10	abask	####k
11	abate	###te
12	eagle	eagl#
13	eaglet	#####t
14	ear	##r
15	eardrop	###drop
16	earl	###l
17	early	####y
18	earmark	###mark
19	earn	###n

Table 3-3. Exemples d'entrées de base et après la première transformation.

Si l'on examine les positions des '#' en lisant ligne par ligne, on peut se rendre compte que les positions marquées sont distribuées relativement au hasard, mais plus souvent à gauche qu'à droite. Mais si on lit colonne par colonne, on peut comprendre facilement l'idée de la compression dans cette méthode. La compression du caractère dans les entrées dépend de la présence du caractère # dans la même position que l'entrée précédente.

Les numéros dans la première colonne de la table 3-3 correspondent aux numéros des entrées originales. Pour la représentation numérique des positions marquées, nous devons passer de cette forme à la forme illustrée par dans la table 3-4.

La dernière ligne contient les numéros des entrées qui finissent par la position correspondance à la colonne. Par exemple, la longueur de l'entrée N° 2, 'aback' correspond à 4 (la longueur est comptée à partir de 0). Cette information est inscrite dans la cellule à l'intersection de la colonne numéro 4 (cinquième colonne) et de la

dernière ligne, comme pour l'information de l'état dans le diagramme de transition (par exemple, le noeud gras dans la figure 3-1). Cette information est gardée en mémoire pour être utilisée dans la fonction **finish_or_not** lors de la consultation du pseudo-code dans la figure 3-11.

	0	1	2	3	4	5	6	7	8
a	0, 11	0, 1 12, 19	2, 11		18, 18				
b		2, 11							
c				2, 3					
d				4, 4 15, 15	4, 7				6, 6
e	12, 19				8, 8 11, 13			6, 7	7, 7
f									
g			12, 13						
h					9, 9				
i									
j									
k					2, 2 10, 10		18, 18		
l				12, 13 16, 17		3, 3			
m				18, 18					
n				5, 7 19, 19	0, 0		4, 7		
o				0, 0		4, 7 15, 15			
p							15, 15		
q									
r			0, 0 14, 19		15, 15	18, 18			
s			1, 1	8, 10				3, 3	
t				11, 11		13, 13			
u					3, 3		3, 3		
v									
w									
x									
y					17, 17				
z									
			1, 14	16, 19	0, 2, 8, 9, 10, 11, 12, 17	13	4, 5, 15, 18	3	6, 7

Table 3-4. Deuxième forme de la table dans CBP.

Les paires de numéros dans les cellules représentent la première et la dernière d'une série d'entrées adjacentes comportant le même caractère dans la même colonne. Par exemple, les deux paires (5,7) et (19, 19) dans la table 3-4 [n] [3] signifient que toutes les entrées entre la 5^e et la 7^e lignes, qui sont 'abandon', 'abandoned', et 'abandonee' comportent le même caractère 'n' dans la colonne n° 3. Mais les deux

entrées adjacentes (18^e et 20^e entrées) de la 19^e entrée 'earn' n'ont pas de caractère 'n' dans la position n° 3.

La taille des cellules de la table 3-4 varie et il y a des cellules vides. La forme n'est pas efficace pour la gestion de la mémoire. Elle peut être remplacée par les tables 3-5 et 3-6. La table 3-5 est appelée TI. La taille maximum de TI est fixée. C'est 26 (nombre de caractères de l'alphabet) * 9 (longueur maximum des entrées). Dans la table 3-3, 'abandoned' et 'abandonee' ont la longueur maximum, soit 9.

	0	1	2	3	4	5	6	7	8
a	1	4	6		22				
b		5							
c				11					
d				13	23				44
e	2				25			42	45
f									
g			7						
h					26				
i									
j									
k					28		38		
l				15		33			
m				16					
n				18	29		39		
o				19		35			
p							40		
q									
r			9		30	36			
s			10	20				43	
t				21		37			
u					31		41		
v									
w									
x									
y					32				
z									

Table 3-5. TI dans CBP.

Chaque numéro de la table 3-5 est le total des nombres accumulés des paires qui sont données dans la table 3-4 dans l'ordre colonne par colonne. Par exemple, le numéro 4 dans TI [a][1] dans la table 3-5 est le nombre cumulé des données (0,11) (la paire dans la table 3-4 [a][0]), (12,19) (la paire dans la table 3-4 [e][0]), (0,1) (table 3-4 [a][1] et

(12,19) (table 3-4 [a][1]). La table 3-6 montre seulement la liste des séquences de la table 3-4 dans l'ordre ligne par ligne. L'information de la table 3-5 est connectée avec celle de la table 3-6. Pour consulter une entrée, les deux tables sont nécessaires. L'information donnée par TI est l'indice de la table 3-6, appelée TH. TH donne uniquement la représentation des séquences de la table 3-4 avec le concept qu'une cellule de TH est une paire de la table 3-4.

Cette méthode a été testée sur une machine **HP**. Pour l'implanter sur un **PC**, TI doit être en mémoire, la taille de TI n'est pas grande. Dans cet exemple, la taille de TI est 234 (= 26 * 9, la longueur maximum des séquences dans cet exemple d'entrées). Mais TH peut exister dans le mécanisme externe du système de fichier. Nous espérons que la méthode CBP figurera dans l'environnement **DOS** dans l'avenir.

0, 11	12, 19	0, 1	12, 19	2, 11	2, 11	12, 13	0, 0	14, 19
1, 1	2, 3	4, 4	15, 15	12, 13	16, 17	18, 18	5, 7	19, 19
0, 0	8, 10	11, 11	18, 18	4, 7	8, 8	11, 13	9, 9	2, 2
10, 10	0, 0	15, 15	3, 3	17, 17	3, 3	4, 7	15, 15	18, 18
13, 13	18, 18	4, 7	15, 15	3, 3	6, 7	3, 3	6, 6	7, 7

Table 3-6. TH.

```

procedure look_up(entry)
begin
  int start, end, from, to, i;
  start=0;
  end= total_number_of_the_entries;
  for (i = 0; i < strlen(entry); i++) {
    if(i == 0) from = 0; else
      from =      TI[entry[i-1] - 'a'] [i]; /* TI dans la table 3-5 */
      to =      TI[entry[i ] - 'a'] [i];
      if(!intersection(TH, start, end, from, to))
/* re-adjust the values start and to,
   start = first value of ( interval(start, end) ∩ interval(from, to)),
   end = last value of ( interval(start, end) ∩ interval(from, to)) */
      return False;
  }
return finish_or_not(start, end);
end.

```

Figure 3-11. Pseudo-algorithme pour CBP.

3.2.2. Evaluation de la performance

Nous avons testé l'espace et le temps nécessaires pour la consultation avec l'exemple du dictionnaire des noms acquis à partir de **WordNet** [22]. L'exemple du dictionnaire possède 20,329 entrées qui ne représentent pas le nombre total des noms. La longueur maximum des entrées est 23. Après exécution du programme, nous avons obtenu 72,345 paires. L'espace nécessaire dans la méthode CBP est estimé ainsi:

$$\begin{aligned}\text{Espace nécessaire} &= (\text{nombre d'entrées} + \text{taille de TI} \\ &\quad + \text{taille de TH}) * \text{taille d'un entier} \\ &= (20,329 + 26 * 23 \\ &\quad + 2 (\text{taille d'une cellule dans TH}) * 72,345) \\ &\quad * 4 \text{ octets (taille du type } \textbf{int} \text{ en langage C).} \\ &= 662,468 \text{ octets} \\ &= \text{soit } 662 \text{ KO.}\end{aligned}$$

L'espace total nécessaire est d'environ 700 KO. La taille de la table TI n'est pas toujours grande et elle est de 2 402 octets ($= 26 * 23 * 4$ octets), parce qu'il n'y a qu'un facteur qui affecte l'espace nécessaire dans CBP. C'est la longueur maximum des entrées dans le lexique. Dans cet exemple, elle est de 23.

La taille TH n'a pas de cellule vide. Elle ne nécessite pas de technique de compression de la table, comme pour l'Automate lexical [35].

Pour programmer la technique de CBP sur **DOS**, on a besoin de deux niveaux de mémoire : TI est chargé au niveau mémoire (par exemple, commande '*malloc*'), et TH est au niveau fichier (par exemple, commande '*fopen*'). En effet, **DOS** ne fournit pas beaucoup de mémoire (il n'a pas de mémoire virtuelle).

Calculons le temps nécessaire avec CBP. Comme avec l'Automate lexical, le temps de consultation dans CBP dépend de la longueur des séquences saisies, mais pas du volume du dictionnaire. Pour mesurer la vitesse de consultation dans CBP, la fonction '*ftime*' a été utilisée sur **HP9000/735**. La mesure est au 1/1000 de seconde près. Nous l'avons testée à plusieurs reprises pour obtenir le temps moyen. La méthode CBP permet de rechercher 18 000 entrées par seconde.

CBP se caractérise par la méthode de la position orientée. Même si elle a plus de paires que l'automate de transition, elle ne nécessite ni la table des procédures de reconstruction, ni la table de compression. Quand un nouvel input est entré dans la méthode de l'automate, tous les états de la transition doivent être reconstruits[8]. Cela demande beaucoup de temps.

Dans l'état actuel de nos recherches, nous ne pouvons pas dire quelle est la méthode la plus efficace. Les deux méthodes possèdent la même complexité de temps pour la consultation. La complexité de mémoire dans CBP est très compétitive.

4. Etiquetage des parties du discours basé sur le lexique

EPL est le marqueur de parties du discours(PDD) qui analyse les phrases entrées par le biais du Lexique-Grammaire [12]. Celui-ci utilise systématiquement cette approche à l'analyser.. EPL, basé sur le système DECO [25], ne couvre pas toutes les PDD de la langue coréenne, car DECO n'est pas un système de dictionnaires complets. L'Analyseur morphologique, EPL, analysera les mots en entrée, s'ils comprennent les types de nom (T1', T2) ou les types de prédicats (T3). D'autres cas, comme (T4), n'ont pas été considérés dans l'analyseur morphologique EPL.

EPL se compose de modules : analyseur lexical, T1', T2, T3 et BDD filtrage. La figure 4-1 montre le déroulement du contrôle.

L'analyseur lexical examine les erreurs syntaxiques et les paires de phrases mal combinées. Il envoie les tokens aux modules T1', T2 et T3.

Ces modules examinent les formes de noms et de prédicats, puis envoient les résultats au module BDD filtrage. Celui-ci gère des règles spéciales comme les règles du *cas vocatif* de SF et du *cas de position* de SF, aussi que l'information BDD. Le mot analysé est finalement retourné avec des informations.

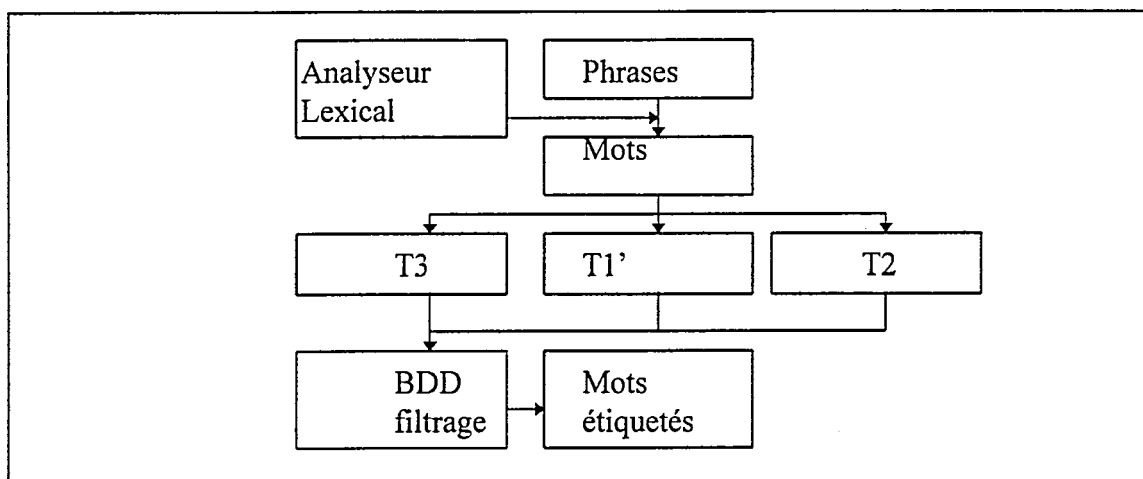


Figure 4-1. Les composantes du système EPL avec le déroulement du contrôle.

4.1. Analyseur lexical

L'analyseur lexical analyse les phrases entrées et les renvoie sous forme de tokens. Le système vérifie aussi les fautes et les erreurs syntaxiques. Il connaît les conditions de fin de phrase. Il y a trois symboles pour marquer la fin d'une phrase : *le point final* (.), *le point d'interrogation* (?), et *le point d'exclamation* (!). Cependant, il n'est pas simple de segmenter un texte en phrases, parce que les symboles spéciaux comme *les guillemets*, *les parenthèses*, *le point final*, *le point d'interrogation*, et *le point d'exclamation* peuvent avoir une incidence sur ce découpage.

Par exemple, certains symboles s'utilisent normalement par paires dans des phrases. Nous les appelons Symboles Pairs. L'état de Symboles Pairs est exprimé par OPEN ou CLOSE comme dans la table 4-1. Durant l'analyse des phrases, si l'analyseur lexical rencontre un nombre impair de symboles, l'état est OPEN, sinon CLOSE. Autrement dit, si l'on rencontre un symbole de fin de phrase à l'état OPEN, la phrase est considérée comme non achevée.

Il existe une exception dans les cas de ` (OPE_Q) et ‘ (CLO_Q) dans la table 4-1. Par exemple, dans une séquence comme *l'expression*, les Symboles Pairs (OPE_Q et CLO_Q) ne sont pas utilisés deux fois. Dans ce cas, si le programme rencontre un symbole de fin de phrase, il est considéré comme terminé, même s'il est OPEN.

	OPEN	NAMING	CLOSE	NAMING
cas 1	(	LOW_O	)	LOW_C
cas 2	[	MID_O	]	MID_C
cas 3	{	HIG_O	}	HIG_C
cas 4	<	COM_O	>	COM_C
cas 5	“	DOU_Q	“	DOU_Q
cas 6	`	OPE_Q	`	CLO_Q
cas 7	‘	CLO_Q	’	CLO_Q

Table 4-1. Les symboles pairs.

4.1.1. Exemples de test

Le programme **snt_chk** est l'analyseur lexical typique. Il vérifie des erreurs syntaxiques et l'agencement des phrases entrées, il reconnaît alors les phrases bien formées.

Les figures suivantes illustrent des exemples de phrases où on a marqué des informations comme '%FROM, %FILE NAME, ...' dans la figure 4-2 et '#TITL' dans la figure 4-3. Elles indiquent la formule du corpus. Par exemple, la ligne de '%FROM' dans la figure 4-2 décrit un texte extrait du journal '한겨레'(HanGurLe). La ligne '#TITL' dans la figure 4-3 montre que le titre du texte est 'Science et Médecine'. Les lignes marquées sont ignorées dans le programme **snt_chk**. Elles s'utiliseront dans d'autres logiciels, comme la recherche d'information'.

Le programme **snt_chk** a analysé et examiné les phrases dans la figure 4-2 et la figure 4-3. La figure 4-4 montre le résultat.

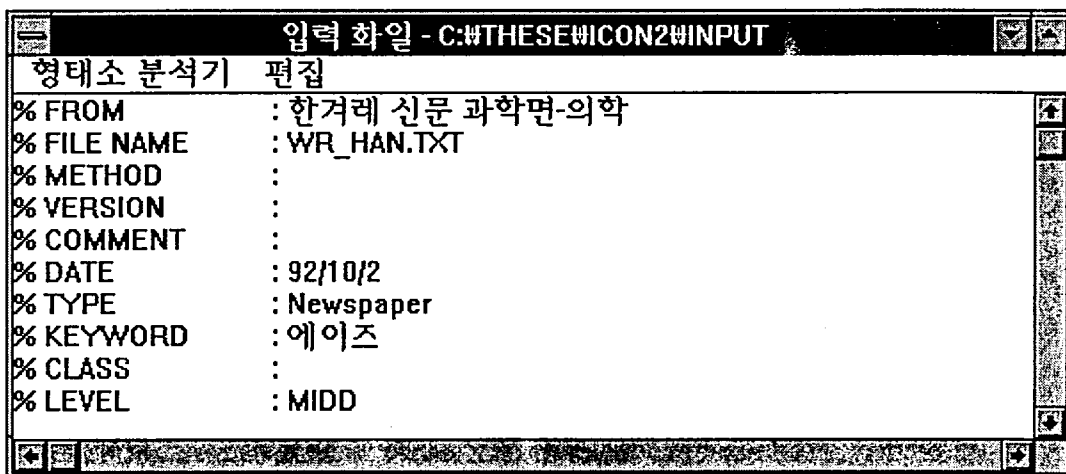


Figure 4-2. Un type de corpus

Les lignes 2 et 4 de la figure 4-3 sont arrangées comme la ligne 2 de la figure 4-4. Les lignes 4 et 8 de la figure 4-3 sont arrangées comme la ligne 3 de la figure 4-4. Les lignes de 9 à 13 de la figure 4-3 sont ajustées comme la ligne 4 de la figure 4-4. Bien qu'il y ait deux *points finals* pour la phrase présentée dans les lignes 9-13, le premier *point final* qui est dans *les guillemets* n'est pas le signal de la fin d'une phrase.

La ligne 14 de la figure 4-3 comporte CLO_Q dans la ligne. Le guillemet est traité comme une marque spéciale dans cette phrase. La ligne 14 de la figure 4-3 est réajustée comme les lignes 5 et 6 de la figure 4-4.

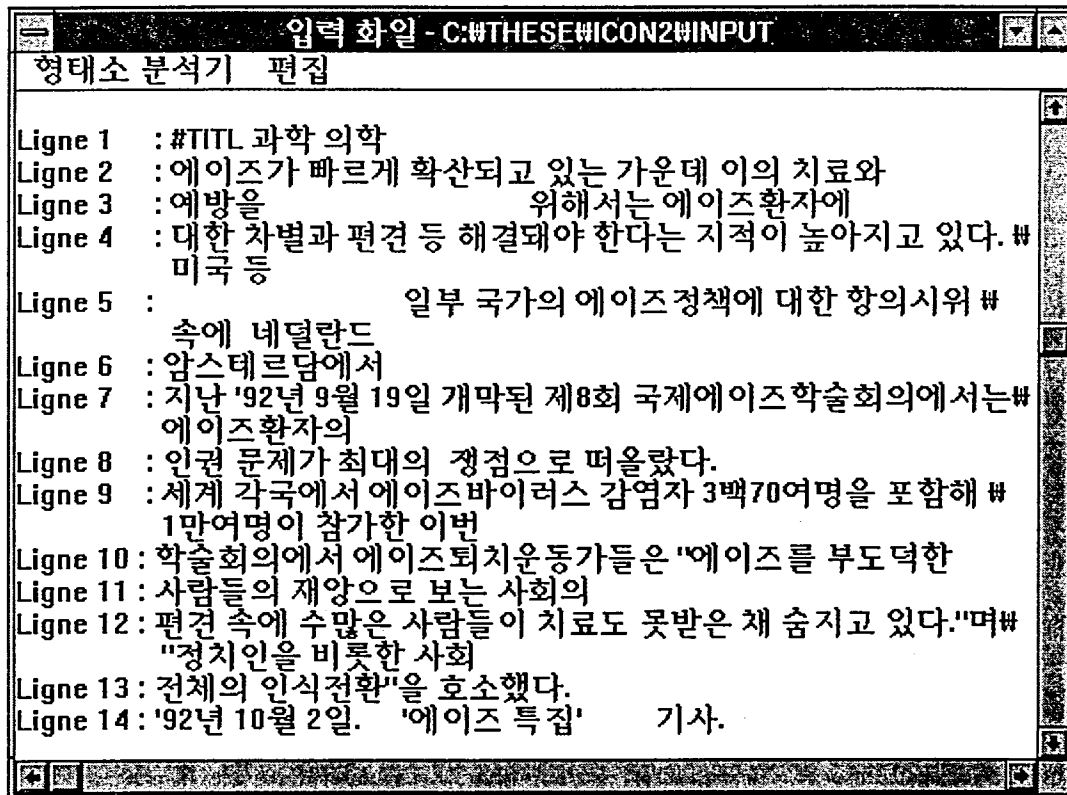


Figure 4-3. Suite de la figure 4-2.

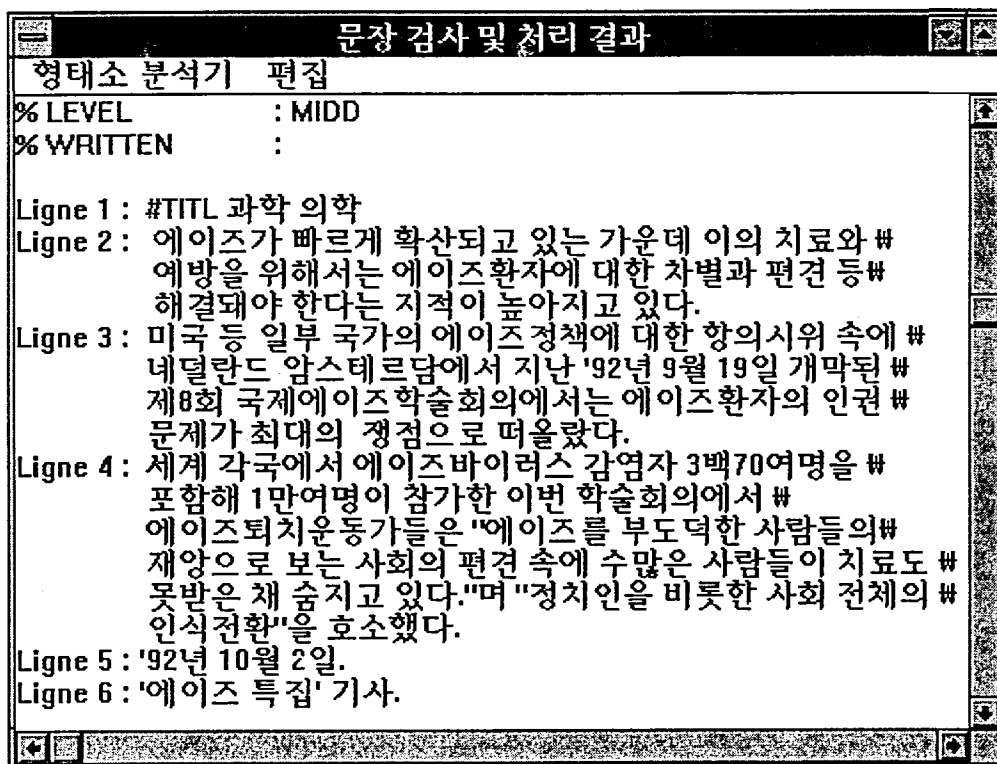


Figure 4-4. Sortie de l'analyseur lexical de la phrase d'entrée de la figure 4-2 et la figure 4-3.

snt_chk met en ordre toutes les phrases en entrée selon le principe “*une phrase sur une ligne*”. Ce principe est utilisé dans l’étape d’analyse morphologique appelée **EPL** dans ce travail.

4.2. Analyseur morphologique

Dans la structure T1', NS peut être situé dans n'importe quelle position du mot. Ainsi, pour trouver les candidats de NS dans le mot, nous n'utilisons pas la recherche exacte, mais une méthode de recherche des sous-ensembles du mot d'entrée.

Supposons, par exemple, que l'entrée est "abc". "a", "b", "c", "ab", "bc" et "abc" sont tous des sous-ensembles de l'entrée. Cela suggère de consulter le dictionnaire dans l'ordre tête vers fin. La méthode de la recherche du sous-ensemble est décrit dans la figure 4-5.

Dans la figure 4-5, la fonction *str_same* guide la recherche dans une position au hasard. Dans les arguments de la fonction, *instr* est un mot en entrée, *str* est un candidat NS, et *i* est un argument de positionnement dans *instr*. Si la recherche entre *instr* et *str* est déclenché (dans ce cas, la fonction *str_same* retourne 1), le candidat de NS(*str*) est enregistré dans les structures *sn->start* et *sn->end* avec les marques de positions *i*. Avec la variable *i*, la fonction *str_same* est appliqué à plusieurs reprises. Par conséquent, tous les candidats NS sont enregistrés dans les structures *sn->start* et *sn->end*. Chaque candidat NS est vérifié avec PF, SF et PPN, puisque le mot est composé de NS, PF, SF et PPN dans la structure T1'.

```
initializing();
j = k = 0;
while(fscanf(dict, "%s", str) != EOF) {    /* dict est NS */
    leng= strlen(str);
    i = 0;
    while ( i < leng) {
        if(str_same(instr, str, i) == 1) {
            sn->start[k][j] = i;
            sn->end[k++][j] = i + leng;
        }
        i = i + 2; /* 2 byte */
    }
    break;
}
```

Figure 4-5. Fonction des sous-ensembles.

Dans les structures T1', T2 et T3, tous les candidats entrées de dictionnaires sont gardés temporairement dans la structure *cand* de la figure 4-6. Dans cette structure, les résultats sont analysés.

Dans T1', NS est la PDD principale dans le mot. Cela veut dire que l'on cherche tout d'abord NS dans le mot, puis les autres particules avec l'information de NS. FN et FFP sont les PDD principaux, respectivement dans T2 et T3. Dans la structure *cand*, la PDD principale est située dans les positions *start[k][0]*, et *end[k][0]*. La variable *k* de

chaque structure dénote les possibilités d'interprétation du mot en entrée. Les autres particules, tels que PF, SF et PPN dans le cas de T1', peuvent présenter des ambiguïtés avec chaque candidat du NS. L'information est gardée dans la position j de la figure 4-5. Le détail est représenté dans la figure 4-6.

```

struct cand {
int    tota;                /* total des candidats */
int    leng;               /* taille de l'entrée */
char   data[WORD];         /* l'entrée */
int    start[NUM][CAND];   /* les premières positions des candidats */
int    end[NUM][CAND];     /* les dernières positions des candidats */
} *sn, *sf, *npp, *is, *fsr, *v, *a, *nis, *np;

struct cant {
int    vert;               /* total des candidats */
int    hori;               /* indicateur de la position courante */
char   data[WORD];         /* l'entrée */
int    flag[WORD];         /* le nombre de l'ambiguïté */
int    type[WORD][WORD];   /* types de PD */
int    loca[WORD][WORD];   /* les positions de la type */
} *result;

```

Figure 4-6. Structure des données de EPL.

Nous allons examiner la procédure d'étiquetage sur un exemple. D'abord, nous avons trouvé tous les candidats possibles pour le NS à partir du mot, concept simple. Le concept n'autorise pas la répétition interne R1 dans le dictionnaire. Si le mot est “정치가는” (le politicien est) (NS), les candidats du NS sont “정치” (la politique) (NS), “가” (NS), et “정치가” (NS-NS). Pour les candidats du NS, la structure *cand* comporte l'information suivante :

<i>sn->tota</i>	= 3, (le nombre de k)
<i>sn->leng</i>	= 8
	(parce que nous utilisons 2 bytes du système de code),
<i>sn->data</i>	= “정치가는”,
<i>sn->start</i> [0][0]	= 0 /* $k = 1$ 정치 */
<i>sn->end</i> [0][0]	= 4
<i>sn->start</i> [1][0]	= 4 /* $k = 2$ 가 */
<i>sn->end</i> [1][0]	= 6
<i>sn->start</i> [2][0]	= 0 /* $k = 3$ 정치가 */
<i>sn->end</i> [2][0]	= 6

1) $k = 1$ cas:

Le premier candidat du NS “정치” est situé dans (0, 4) de la *data* (0, 8). (0, 4) spécifie la forme de la quatrième position à partir de zéro de *sn->data*. La taille de la donnée en entrée de *sn->data* est (0, 8).

Dans une autre position comme (4, 8), on peut trouver SF ou PPN. PF est situé à la tête du NS, mais le candidat NS occupe la première place du mot. Donc, il n’y a pas de PF.

SF diffère de PPN par le fait que SF est toujours en tête du PPN. Cela signifie que les candidats SF, “가는”(4, 8) et “가”(4, 6) sont possibles, mais pas “는”(6, 8). Si SF est utilisé, il doit commencer en position 4. D’ailleurs, “가는”(4, 8) et “는”(6, 8) peuvent être les candidats PPN, qui doit se terminer à la position 8. Les autres formes contreviennent à ce principe dans la structure T1’.

Dans le dictionnaire de SF, on a les entrées “가”(4, 6) et “는”(6, 8). Dans SF, la valeur de *j* est 2 et cela signifie qu’il y a deux interprétations possibles. L’information est enregistrée provisoirement dans la structure *cand*, ainsi :

```
sf->start[0][0] = 4
sf->end[0][0] = 6
sf->start[0][1] = 6
sf->end[0][1] = 8.
```

Dans le dictionnaire PPN, les entrées “가”(4, 6) et “는”(6, 8) existent aussi. Dans *cand*, la forme PPN est la même que dans le cas de SF. Par conséquent, pour la place (4, 8), SF (가) (4, 6) - PPN (는) (6, 8) est seulement autorisé. Les autres combinaisons violent le principe de T1’, “SF se situe toujours à la tête de PPN”.

Pour le premier candidat du NS, nous pouvons proposer la forme NS (정치) (0,4) - SF (가) (4, 6) - PPN (는) (6, 8).

Les résultats positifs de l’analyse sont enregistrés dans la structure *cand* de la figure 4-5 le contenu est comme suit:

```
vert = 1          /* La donnée actuelle est le premier candidat de NS */
hori = 3          /* Le nombre des types est 3 */
data = “정치가는”
flag[0] = 1       /* Pour le premier candidat, il y a seulement une forme,
                  cela signifie qu’il n’y a pas d’ambiguïté
                  Pour le premier candidat de NS */
type[0][0] = 2    /* type 2 dénote NS */
loca[0][0] = 4    /* la dernière position de SN dans data */
type[0][1] = 3    /* type 3 dénote SF */
loca[0][1] = 6    /* la dernière position de SF dans data */
```

<i>type</i> [0][2] = 4	/*	type 4 dénote PPN	*/
<i>loc</i> [0][2] = 8	/*	la dernière position de PPN dans <i>data</i>	*/

2) k = 2 cas:

Le deuxième candidat NS est “가”(4, 6). Pour analyser cette deuxième possibilité, nous devons la rechercher dans le dictionnaire de PF car seule la forme PF peut se situer à la tête du NS. Mais, il n’y a pas d’entrée “정 치” dans le dictionnaire. C’est un cas d’échec.

3) k = 3 cas:

Le troisième cas est une combinaison des deux cas précédents. Le troisième candidat NS, “정 치가”(0, 6) (NS-NS), est la forme répétitive de NS. Cette forme est acceptable dans T1’ comme “정 치” (la politique) (NS) - “가” (l’addition) (NS) - “는” (PPN) (nominatif). Bien que la forme marquée, NS-NS-PPN, n’ait pas d’erreur syntaxique, le deuxième NS (l’addition) “가” est en conflit avec SF et n’est sémantiquement pas acceptable. BDD filtrage dans la figure 4-1 n’a pas d’information pour lever l’ambiguïté P1 entre NS et d’autres interprétations possibles.

T2 est le même que le cas de T1’. La PDD principale de T2 est PN. T2 possède en plus le dictionnaire par rapport à T1’.

T3 utilise aussi les structures *cand* et *cant* pour l’analyse du mot. D’abord PPP, puis SFR, le prédicat et les séquences du nom sont appliqués. La structure T3 est analysée avec les séquences en ordre inverse dans le mot, parce que la particule principale de T3 est PPP qui est placé à la fin du mot. Les procédures d’analyse sont les mêmes que dans les cas T1’ et T2.

Par exemple, la séquence “비 오 는”(dans la pluie) est analysée comme “비”(NS) (0, 2), “오 ” (SFR) (2, 4), et “는” (PPP) (4, 6) dans l’ordre inverse de la structure T3. Un autre exemple comme “고 독 하 게”(solitairement) est composé de “고 독 하”(le radical, 하 est SFR) (0, 6), et “게” (PPP) (6, 8). Même si la forme “고 독” (0, 4)(solitude) est analysée, rappelons que nous avons enregistré que toutes les formes d’adjectifs comportant ‘하 다’(faire) sont dans le dictionnaire des adjectifs.

4.2.1. Exemples de test

Notre méthode détecte beaucoup d'ambiguïtés, mais elles sont presque inévitables au niveau de la morphologie, dans la mesure où nous n'utilisons pas d'heuristiques approximatives. Pour le NS, la séquence “광산은” a deux possibilités d'étiquetage. L'une est “광산”(la mine) (NS) et “은” (l'argent) (NS). L'autre est “광산”(la mine) (NS) and “은” (nominatif) (PPN). L'étiquetage correct dans ce cas dépend du sens ou de la structure de la phrase.

Mis à part NS, les autres particules incluant PF, SF, PPN, FN, PPP, SFR, V et A ne posent pas de problème. Tous les problèmes morphologiques ont été filtrés dans le module BDD filtrage comportant l'information BDD et les règles.

입력화일 - C:\THESEWICON2\INPUT

형태소 분석기 편집

% FROM : 중앙일보 분수대
% FILE NAME : WR_BU.TXT
% METHOD :
% VERSION :
% COMMENT :
% DATE : 92년 04월 12일자
% TYPE : Newspaper
% KEYWORD :
% CLASS :
% LEVEL : HIGH
% WRITTEN :
#TITL 백세시대
건강에 이상을 느껴 병원에 온 사람을 진찰하고 나서 의사가 말했다.
건강이 몹시 나빠지셨군요.
담배를 끊으시고 여색을 피하십시오.
애기를 듣 고난 환자가 탄식하듯 중얼거렸다.

Figure 4-7. Texte entré pour le système EPL.

Plusieurs exemples sont donnés dans la figure 4-7. Les figures 4-8-1, 4-8-2 et 4-8-3 montrent les résultats d'exécution de EPL pour les entrées de la figure 4-7. Mais ces résultats sont obtenus avant de passer BDD filtrage (cf. figure 4-1). Le premier exemple, “건강에” peut produire trois candidats, comme NS-SF-NPP, NS-PPN, et PF-NS-NPP. On ne peut pas choisir l'un des trois candidats, parce que nous n'avons pas de dictionnaire ND (Nom Dérivé). Dans ce cas, nous fournissons les trois candidats.

Un autre exemple est un cas de “의사가” (médecin est) dans figure 4-8-3. Il y a sept possibilités d'analyse. Mais, après passage dans le module BDD Filtrage, les deux interprétations ‘1th’ et ‘4th’ seront supprimées, car cette entrée “의사가” (médecin est) n'est pas trouvée dans le module. C'est-à-dire, on n'interprétera pas certaines entrées comme *noms quelconques* + SF. Si on avait un dictionnaire ND, on pourrait enlever la

première et la dernière interprétation comme “의” (NS) + “사” (SF) + “가” (PPN) et “의” (PF) + “사-가” (NS).

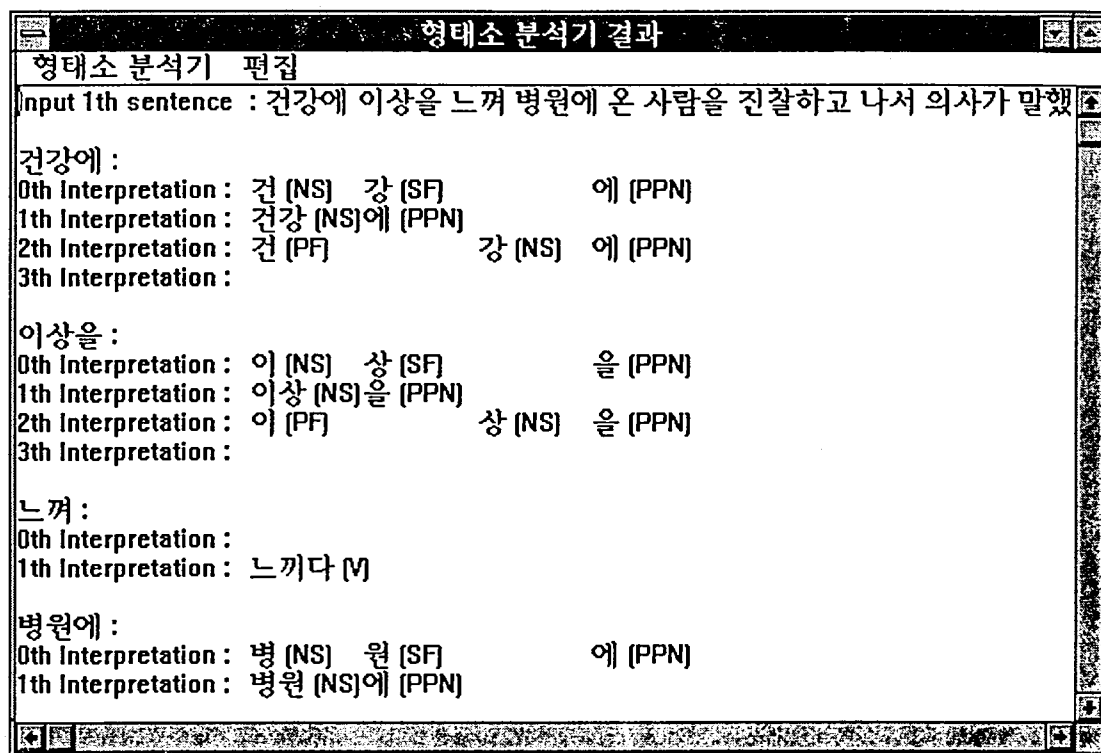


Figure 4-8-1. Résultats d'exécution de EPL des entrées de la figure 4-7.

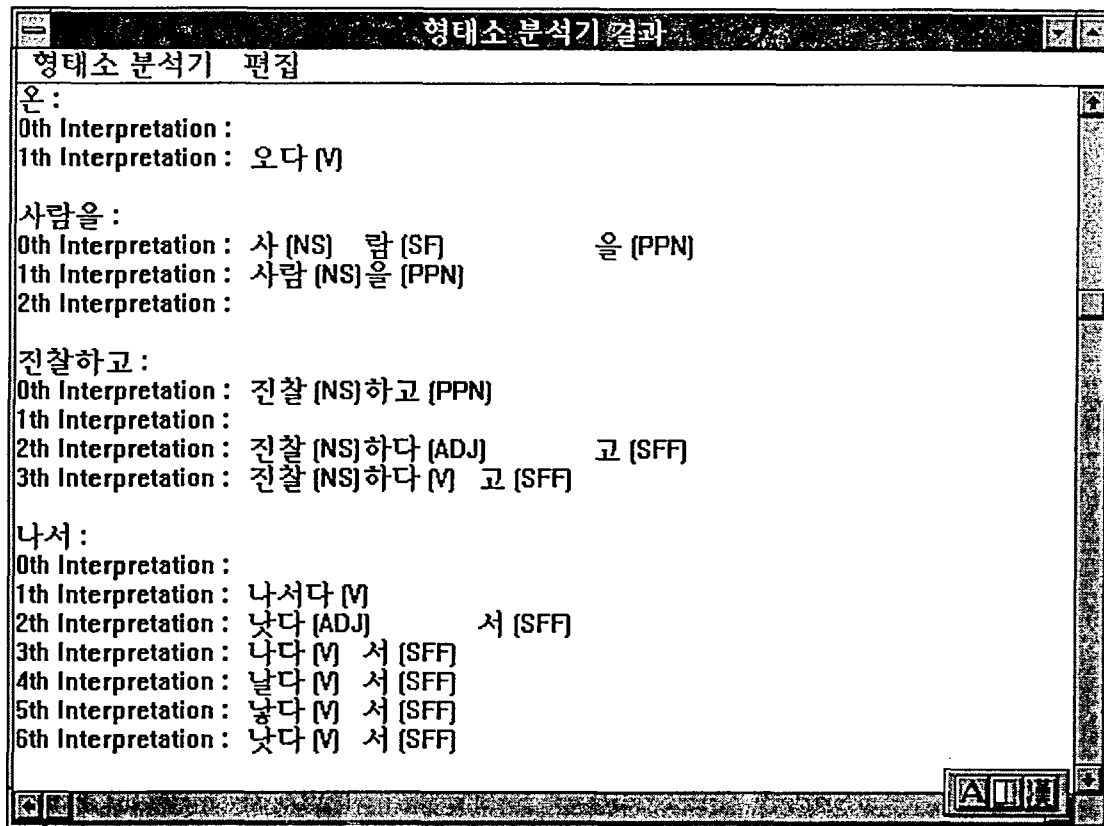


Figure 4-8-2. Suite de la figure 4-8-1.

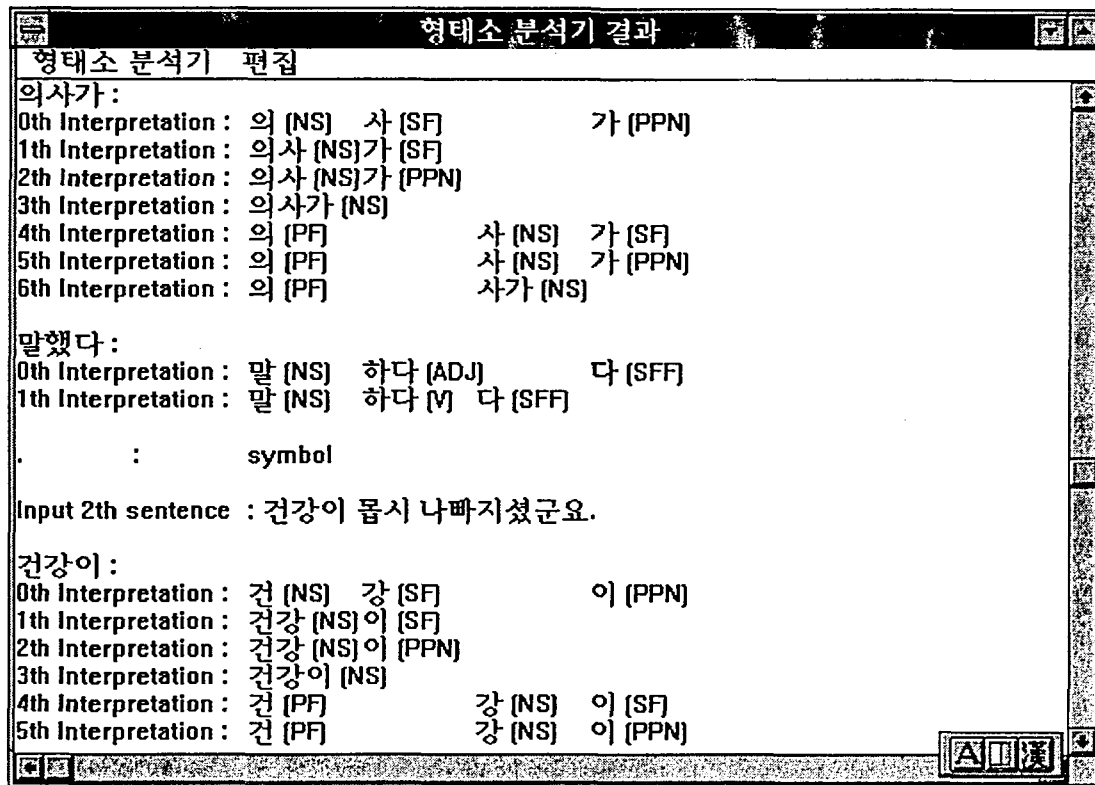


Figure 4-8-3. Suite de la figure 4-8-2.

5. Conclusion

Dans ce travail, nous avons développé un analyseur morphologique du coréen. La particularité de notre système consiste en l'utilisation de dictionnaires construits systématiquement. Etant donné qu'en coréen, les unités délimitées par des séparateurs ne sont que rarement des unités lexicalement simples, mais des séquences plus complexes (e.g. les postpositions casuelles sont soudées aux substantifs, et par ailleurs elles se combinent entre elles), des dictionnaires qui ne comportent que des mots simples ne sont pas suffisants pour reconnaître tous les mots et analyser correctement un texte donné.

Pour ce travail, nous avons engendré des dictionnaires de formes fléchies à partir des dictionnaires des *mots simples*, et des dictionnaires d'*affixes* et de *postpositions* : Dictionnaire des *nominalisations* des termes prédicatifs (FN), et Dictionnaire des formes fléchies des *verbes* et des *adjectifs* (FFP). Pour les cas où il existe des homographes (e.g. entre *suffixe* et *postposition*), nous avons établi des listes complètes de noms dérivés comportant ce type de suffixes. Cette liste est intégrée dans la module BDD. Quand un nom se trouve à gauche d'un morphème qui peut entrer éventuellement dans la liste des suffixes, on cherchera cette séquence dans le dictionnaire des mots dérivés de type N-SF : si on ne le trouve pas dans ce dictionnaire, cette séquence sera analysée comme N-PPN.

En nous basant sur ces dictionnaires, nous avons implémenté un système d'analyse morphologique baptisé **EPL**. Ce système propose toutes les interprétations possibles d'une séquence donnée. Voici un exemple :

<i>janeun</i> : 1>	자 (<i>ja</i>) [V:dormir]	+ 면	(<i>neun</i>) [PPP:Déterminatif]
2>	자 (<i>ja</i>) [NS:règle]	+ 면	(<i>neun</i>) [PPN:Nominatif]
<i>saineun</i> : 1>	새 (<i>sai</i>) [NS:oiseau]	+ 면	(<i>neun</i>) [PPN:Nominatif]
2>	새 (<i>sai</i>) [V:fuir]	+ 면	(<i>neun</i>) [PPP: Déterminatif]
<i>jugessda</i> : 1>	죽 (<i>jug</i>) [V:mourir]+	었다(<i>essda</i>)	[PPP:Déclaratif_Passé]

Figure 5-1. Exemples de l'analyse morphologique du coréen.

Nous devons ajouter quelques données lexicales dans la version actuelle de notre système : il s'agit d'une liste d'éléments relativement courte. Les dictionnaires qui comportent des *adverbes*, des *pronoms*, des *noms incomplets* et des *numéraux* doivent être construits selon les mêmes principes.

Les dictionnaires des *noms propres* et des *noms composés* sont à établir à part : ils sont importants en nombre d'entrées et ils ont leur syntaxe particulière. Dans l'analyse morphologique de textes écrits et la documentation automatique, ces dictionnaires seront indispensables.

Nous nous sommes limité, dans ce travail, à la construction d'un analyseur morphologique du coréen, mais l'analyse syntaxique d'un texte donné sera nécessaire dans tous les systèmes de traitement de documents écrits en langue naturelle. Lors de l'analyse syntaxique et sémantique de textes coréens, le *lexique-grammaire* et des *grammaires locales* doivent être utilisés[18], aussi bien que les dictionnaires que nous construisons.

BIBLIOGRAPHIE

- [1] Cho Young-Hwan, 1992, *Hangul TRIE and Bidirectional Access Method*, RT, Kaist CS Lab.
- [2] Choi Ki-Sun et al., 1993, *Le développement de la documentation informelle*, Rapport du Departement Commercial, Séoul,
- [3] Choi, Sung Woo, 1996, "Hangul Code Manager", A meeting of Korean Linguistics in CERIL : *Mémoires du CERIL* N° 14.
- [4] Courtois, Blandine, 1987, *Dictionnaire Electronique du Laboratoire d'Automatique Documentaire et Linguistique pour les mots simples du francais* (DELAS), RT-17, LADL, Université Paris 7.
- [5] Courtois, Blandine, 1989, *Dictionnaire Electronique du LADL pour les mots flechis du francais* (DELAF), RT-17, LADL, Université Paris 7.
- [6] Crochemore, Maxime et Rytter, Wojciech, 1994, *Text Algorithms*, Oxford University Press.
- [7] Crochemore, Maxime et Hancart, C., 1997, *Automata for matching patterns*, IGM-97-06, Université de Marne-La-Vallée..
- [8] Gazdar, G., Klein, E., Pullum, G., et Sag, I., 1985, *Generalized Phrase Structure Grammar*, Oxford : Blackwell.
- [9] Gazdar G. and Mellish, 1989, *Natural Language Processing in LISP*, Addison Wesley.
- [10] Gross, Gaston, 1994, "Classes d'objets et description de verbes", *Langages* 115.
- [11] Gross, Maurice, 1975, *Méthodes en Syntaxe*, Hermann,.
- [12] Gross, Maurice, 1987, *Electronic Dictionaries and Automata in Computational Linguistics, Lecture Notes in Computer Science of the LITP Spring School on Theoretical Computer Science*, Berlin: Springer-Verlag.
- [13] Gross, Maurice, 1987, The use of finite automata in the lexical representation of natural language, *Electronic Dictionaries and Automata in Computational Linguistics, Lecture notes in Computer Science* 377, Berlin : Springer-Verlag.
- [14] Han, Sun-Hae, 1993, "Sur la construction nominale *NO NI-leul Npréd-leul Hada*", *Mémoires du CERIL*, N°12, Institute Gaspard Monge, Université de Marne-la-Vallée.
- [15] Kang, HyunKyu, LEE, ChangYeol et al, 1994, "An Implementation of an automatic Keyword Extraction System", *The 3rd Pacific Rim International Conference on Artificial Intelligence (PRICAI)*, Beijing, China.
- [16] Karttunen, Lauri, 1993, *Finite-State Lexicon Compiler*, Xerox, Palo Alto Research Center, ISTL-NLTT-1993-04-02.
- [17] Krovetz, Robert et W. Bruce Croft, 1992, Lexical Ambiguity and Information Retrieval, *ACM Trans. on Information Systems* Vol.10 N° 2.
- [18] Laporte, Eric, 1994, "Experiences in lexical disambiguation using local grammars", *In COMPLEX 94, Papers in Computational Lexicography*.
- [19] Lee, ChangYeol, 1996, "Ambiguity Processing in DECO system", A meeting of Korean Linguistics in CERIL : *Mémoires du CERIL* N° 14.

- [20] Liang , Franklin Mark, 1983, *Word Hy-phen-a-tion by Com-put-er*, PhD. thèse, Department of Computer Science, Report N° STAN-CS-83-977 : Stanford University.
- [21] Longman Language Activator, 1993, *The World's First Production Dictionary*, Longman Group UK.
- [22] Miller, George A. et al., 1990, *Introduction to WordNet: An On-line Lexical Database*, Cognitive Science Lab., Université Princeton, CSL Report N° 43.
- [23] Nam, Jee-Sun, 1994, *Classification syntaxique des constructions adjectives en Coréen*, Vol II-annexes, Ph.D. Université Paris 7.
- [24] Nam, Jee-Sun, 1994, *Dictionnaire Des Noms Simples Du Coréen*, RT-46, LADL.
- [25] Nam, Jee-Sun, 1996, *Dictionary of Korean Simple Verbs*, RT-49, LADL.
- [26] Nam, Jee-Sun, 1996, *Building A Korean On-Line Dictioanry Adequate to Parsing Systems*, RT-50, LADL.
- [27] Nam, Jee-Sun et Lee, ChangYeol, 1996, *Korean Electronic Dictionary I/II*, IGM96-3/4, Université de Marne-La-Vallée.
- [28] Nam, Jee-Sun, 1996, *Lexicon of Korean predicative terms classified by mophological ending forms*. Vol I., IGM 96-8, Université de Marne-La-Vallée.
- [29] Nam, Jee-Sun, 1996, *Lexicon of Korean predicative terms classified by mophological ending forms*. Vol II., IGM 96-9, Université de Marne-La-Vallée.
- [30] Nam, Jee-Sun, 1996, *Dictionary of N-Postpositions, A-Postpositions, and V-Postpositions in Korean*. RT-51, LADL.
- [31] Nam, Jee-Sun, 1997, *Traitement des séquences comportant hada dans un lexique électronique du coréen*, RT-IGM 97, Université de Marne-la-Vallée.
- [32] Park Dong-Sun, 1989, "La standardisation de Hangul et Hanza", *OA Lecture*, Séoul, p98~ p106.
- [33] Pereira, Fernando C.N. et Barbara J. Grosz, 1994, *Natural Language Processing*, Cambridge, Mass. : The MIT Press.
- [34] Perrin, Dominique, 1989, "Automates et algorithmes sur les mots", *Annales des télécommunications*, p20- p33.
- [35] Revuz, Dominique, 1991, *Dictionnaires et Lexiques Methodes et Algorithmes*, Université Paris 7, LITP 91-44.
- [36] Shim Jae-Ki, 1980, "La fonction de la nominalisation", *Langue*, Vol. 5, N° 1, Séoul, p79 ~ p102.
- [37] Shin Ki-Chul, Shin Yong-Chul, 1990, *Nouveau Dictionnaire de la Langue Coréenne*, Séoul, *Samsung Chulphansa*.
- [38] Silberztein, Max, 1993, *Le système INTEX*, Masson.
- [39] Tarjan, Robert Endre et Andrew Chi-Chih Yao, 1979, "Storing a Sparse Table", *CACM* Vol. 22 N°11, Stanford University.
- [40] Walker, Donald E., et al., 1995, *Automating The Lexicon: Research and Practice in a Multilingual Environment*, Oxford : Clarendon Press

ANNEXE A. Les entrées ambiguës P1 entre PF et NS

(de PF) (252) 가 가장 각 간 감 갑 강 개 객 건 격 결 경 계 고 동 곡 골 곶 공 과
 관 광 광대 구 정 국 군 귀 군 극 극대 극소 극한 금 급 기 길 낙 난
 날 남 날 내 냉 노 녹 논 농 다 다갈 단 단장 담 당 대 대접 도 도마
 독 동 되 뒤 득 들 등 때 마 막 만 말 망 매 면 명 모 모시 목 묘 무
 무당 묵 문 미 밀 박 반 발 방 배 백 범 법 변 별 보라 복 복사 본
 봉 부 분 불 비 빗 사 사시 산 살 삽 상 상고 상품 새 선 설 성 세
 소속 술 수 수양 순 술 쉬 시 식 신 실 쌍 악 안 알 암 애 액 약 양
 양재 얼 얼룩 역 역순 연 연두 열 염 영 예 오 지 옥 온 난 올 왕 요
 욕 용 용수 ^으 ^으 두 운 원 위 윤 융 은 음 의 이 인 입 자 잔 장 재
 쟁 저 적 전 ^천 ^천 절 점 접 정 제 비 조 조약 좀 종 좌 주 죽 중 중품
 진 진 진도 질 징 짝 차 참 창 채 처 천 철 청 체 초 축 춘 최 단 추
 축 축사 탈 태 토 통 투 파 판 패 편 평 폐 포 표 품 풍 피 하 품 학
 한 함 함박 합 해 향 허 현 형 호 호로 호리 혹 혼 황 회 효 훈 흥

de PF de T1' 극대 극소 극한 다갈 상품 역순 온난 중품 최단 하품
 (10)

(262 au total)

ANNEXE B-1. Les candidats de l'ambiguïté P2 entre PF et NS

Les candidats suivants de l'ambiguïté P2 sont produits automatiquement à partir de la liste de 131 PF et 2,739 NS. Ces candidats sont les formes acceptables dans T1'.

가장(PF) +	가골관군기난남내단담대도마면모문물미발복본인부비사삼성 소손수식신악안애원인자전정죽천티푸스파품학해화(NS).
강남(PF) +	매바위발비색자장창포(NS).
강남(PF) +	군만비색인패(NS).
고동(PF) +	감갑강결경계공기면명무문물반백복봉산상설안업 원위의이 인자전정조참창체침태토티포품한향호인화(NS).
고품(PF) +	격계귀명목별사성위절(NS).
곱사(PF) +	가각감건격계골공과관교복교성기나이내다리단담당대도돈동 마귀막망면명모무문물발방범변별복본부분비사살상색생자생화 선설시식신실양업역연열영육용원위의이인입자재적전절정 조죄주죽중진질창채체촌춘기태통투팔표학형화활회(NS).
곱수(PF) +	간감갑건검공관구가국금급난납뇌다단달당도동라마면모목 미박반배복부분비사사학산상색성세소속수술식신액양양녀업 염영예요육용용액은음의인입장재저절정제비제품종질집차첩청 초촉탈태평포표피하물학호화화기화인화자(NS).
광대(PF) +	가가리각강개결경계공과관국궐금기농담담대도독마망면모목 물변본비사상성세소수양역용음위의인입장적전접조좌중진 질차책첩청체추파판패포폭표품피학한함합해화회(NS).
구공(PF) +	간감개격경공관국금급기납막매명모무문물박방백별범복부비 사산상설세소수시식안약양업역연영예용인입자장저적전정주 중채책천치사토통판평포표학한지해화제회(NS).
구정(PF) +	각간경계곡공관글기당도독돈무물미소밀도박방별범변복본부 비사산선설성세수수리식신액약양열염예원의인자장쟁적전절 정조종진책처청체초탐토통품학합해형혼화(NS).
극단(PF) +	가결계골기도독문발복산상소속식신열원위일자장전정조종좌 죄주청체촉풍합화(NS).
극대(PF) +	가가리각강개결경계공과관국궐금기농담담대도독마망면모목 물변본비사상성세소수양역용음위의인입장적전접조좌중진 질차책첩청체추파판패포폭표품피학한함합해화회(NS).
극소(PF) +	각감강개경계관국국금급기대독동등등라마맥명모묘문박반방 복산설속식신심성심증액양요원위의이인일입장정조주진질집 추침통파포품풍피(NS).
극한(PF) +	계과극기대도문방복설수역육의자정참창천탄파풍학화(NS).
급회(PF) +	개계관국귀금담도동복부비사상선수순식신안연오리 우원의자 장전조중진피한함향화(NS).
나막(PF) +	내대동이바지사(NS).
나박(PF) +	격대봉사살색수애자장진차탈학해(NS).
냉난(PF) +	간관국독동무발봉사산세소입자정투파포해황(NS).
누비(PF) +	가결계관단등라마면명문밀방상성소속화약용운위자장재전촉탈 토판판평품호화황(NS).
다갈(PF) +	고리구리기대등망매기무리비삼수피(NS).
단장(PF) +	가골관군기난남내단담대도마면모문물미발복본인부비사삼성 소손수식신악안애원인자전정죽천티푸스파품학해화(NS).
대접(PF) +	경골대변속수시영전종촉합(NS).
도마(PF) +	감개고자구간귀당마모무리부비사회소수술음의적중(NS).
도요(PF) +	가강건결금기도동망변부사세소약양원인절정철체통해(NS).
동아(PF) +	가가리가미국귀기낙내들명미범부비사성씨악양연우이쟁전집 침침편(NS).
딱정(PF) +	각간경계곡공관글기당도독돈무물미소밀도박방별범변복본부 비사산선설성세수수리식신액약양열염예원의인자장쟁적전절 정조종진책처청체초탐토통품학합해형혼화(NS).

마제(PF) +
 맨헛(PF) +
 맹혹(PF) +
 모시(PF) +
 무당(PF) +
 배내(PF) +
 보금(PF) +
 복사(PF) +
 부시(PF) +
 사시(PF) +
 삼살(PF) +
 상고(PF) +
 상품(PF) +
 세양(PF) +
 생안(PF) +
 선들(PF) +
 설농(PF) +
 송골(PF) +
 수양(PF) +
 순헛(PF) +
 신접(PF) +
 심심(PF) +
 안건(PF) +
 양제(PF) +
 어금(PF) +
 열간(PF) +
 역순(PF) +
 연두(PF) +
 오막(PF) +
 오목(PF) +
 오지(PF) +
 온난(PF) +
 용수(PF) +
 우두(PF) +
 원두(PF) +
 월계(PF) +

가 감 계 공 과 관 국 기 답 도 독 등 막 매 명 목 물 반 방 법 복 본 분 비 사 세
 소 시 약 염 위 의 제 전 정 조 주 창 초 패 한 해 회 (NS).
 간 (NS).
 사 상 세 자 정 한 (NS).
 가 각 골 공 국 골 내 달 도 등 무 발 범 비 사 사 실 사회 상 설 세 술 식 신 역 위
 음 가 의 인 인 장 접 정 조 중 중 청 체 초 주 침 판 한 해 현 효 (NS).
 국 나 귀 도 면 목 번 부 분 사 선 위 의 일 좌 질 초 (NS).
 각 계 공 국 기 당 막 복 본 색 성 세 속 수 시 실 알 연 용 의 자 장 계 재 인 접 정
 조 중 진 처 통 포 피 한 향 홍 (NS).
 기 납 단 물 수 실 액 옥 용 주 혼 (NS).
 가 각 감 건 건 격 계 골 공 과 관 교 복 교 상 기 나 이 내 다리 단 담 당 대 도 든 동
 마 귀 막 망 면 명 모 무 문 물 발 방 범 변 별 복 본 부 분 비 사 실 상 색 생 자 생 화
 선 설 시 식 신 실 양 업 역 연 열 옥 용 원 위 은 의 이 인 임 자 재 제 전 질 정
 조 좌 주 주 죽 중 국 골 내 달 도 등 무 발 범 비 사 사 실 사회 상 설 세 술 식 신 역 위
 음 가 의 인 인 장 접 정 조 중 중 청 체 초 주 침 판 한 해 현 효 (NS).
 음 기 모 사 무 사 상 의 인 포 해 (NS).
 가 개 객 고 학 공 관 국 군 금 기 난 뇌 대 도 독 등 명 모 무 문 물 발 배 백 별
 분 비 사 사 터 성 소 속 수 시 안 역 옥 용 원 의 인 자 자 질 장 저 적 전 정 조 중
 주 질 접 참 체 초 후 투 풍 학 함 향 혹 (NS).
 격 등 계 귀 명 목 별 사 성 위 질 (NS).
 금 개 경 광 내 대 도 마 면 목 무 부 산 식 염 위 전 정 주 중 질 (NS).
 통 (NS).
 가 간 경 공 기 노 답 도 목 본 부 사 산 상 성 악 액 업 장 축 토 학 한 기 (NS).
 격 등 목 무 반 상 상 수 질 초 통 (NS).
 가 각 계 국 국 공 난 단 도 체 든 동 이 말 명 모 미 간 병 복 봉 부 분 산 상 서 류
 성 수 식 악 어 장 옥 용 위 은 인 자 장 계 계 기 주 처 키 품 풍 학 해 회 (NS).
 간 (NS).
 경 골 대 변 속 수 시 영 전 중 축 합 (NS).
 계 금 기 뇌 도 독 문 복 부 사 상 성 술 신 심 증 연 열 의 장 저 정 천 통 해 회
 홍 (NS).
 결 계 국 난 단 담 만 목 발 범 부 사 선 소 소 화 수 식 신 은 음 접 조 주 질 책 접
 축 통 투 사 관 호 (NS).
 가 간 개 계 귀 기 난 단 담 독 목 무 물 미 발 범 변 봉 사 산 상 색 선 수 수 생
 실 액 연 옥 원 의 인 임 자 적 정 주 질 천 축 축 침 판 패 편 학 함 향 혼 화 회 홍
 (NS).
 기 납 단 물 수 실 액 옥 용 주 혼 (NS).
 축 통 투 사 관 호 (NS).
 간 결 국 대 도 라 방 백 산 수 시 열 위 장 중 차 풍 화 회 (NS).
 건 독 목 미 발 부 상 통 (NS).
 내 대 동 이 바 저 사 (NS).
 가 각 격 금 기 단 등 마 사 수 옥 이 장 계 적 전 조 질 차 축 침 탄 포 향 화 회
 (NS).
 가 각 개 계 경 관 구 성 구 전 국 금 기 단 대 도 명 목 물 상 반 방 배 부 분 불 사
 성 세 속 시 식 애 양 역 연 우 원 위 은 인 자 장 계 적 정 정 화 조 주 진 질 청 체
 침 탄 평 표 화 향 회 (NS).
 간 관 국 독 등 무 발 봉 사 산 세 소 입 자 정 투 파 포 해 황 (NS).
 간 감 감 건 건 검 공 관 구 가 국 금 금 난 납 뇌 다 단 달 당 도 동 라 마 면 모 목
 미 박 반 배 복 부 분 비 사 사 화 산 상 색 성 세 속 수 술 식 신 액 양 양 녀 업
 염 영 예 요 옥 용 액 은 음 의 인 입 장 재 저 질 정 제 비 제 품 중 질 접 차 침 청
 초 축 탁 탈 테 평 포 피 하 물 학 호 화 화 기 화 인 화 자 (NS).
 건 독 독 목 미 발 부 상 통 (NS).
 건 건 독 목 미 발 부 산 속 수 시 약 장 질 정 진 접 책 통 피 획 (NS).

이복(PF) +	각강귀도면부사상사선수역용원잡화장종창채판화합화술 (NS).
이중(PF) +	가강결국기다리단대말목표무반복비사손식신씨양업용이자적전주창합 (NS).
자봉(PF) +	전급기분양오리인창책토투합화황 (NS).
재중(PF) +	가강결국기다리단대말목표무반복비사손식신씨양업용이자적전주창합 (NS).
저품(PF) +	적제귀명목별사성위절 (NS).
적늑(PF) +	각내장말음초화 (NS).
전전(PF) +	각개격정골과관광국기농단달담당대도도금동돌라말망매모무반방법범병복부분불비사설성소속수술시신실업역용운원위유물의이인임입자장제쟁전정조주진집체체초축토통폐폐학합해향화황 (NS).
제비(PF) +	가결제관단등라만명문밀방상성소속화약용운위자장재전축탈토판평품호화황 (NS).
조약(PF) +	간골과관도방본사소속술식자정조진질체탈호혼 (NS).
주마(PF) +	감개고자구간귀당마모무리부비사회소수술음의적중 (NS).
주안(PF) +	개정광내도마면목무부산식염위전정주중질 (NS).
중품(PF) +	적제귀명목별사성위절 (NS).
진도(PF) +	강공금전주중처체태통포피해화화지회 (NS).
청늑(PF) +	읍입장적전주중처체태통포피해화화지회 (NS).
초생(PF) +	각내장말음초화 (NS).
초승(PF) +	가각강제기도동감득명도물부불사산선수식신에원자장질체청체태판포화활 (NS).
최단(PF) +	가강구강장객계낙마무복부상선세소용운인전진패 (NS).
축사(PF) +	가결제골기도득문발복산상소속식신열원위일자장전정조종좌좌주형체축공합화 (NS).
투명(PF) +	가각감전격계골공과관교복교성기나이내다리단담당대도돈등마귀망면명도무문물발방법변별복분부분비사살상색생자생화선설시식신실양업역연열영육용원위의이인임자재적전절정조좌주중진질향채체촌춘기태통투팔표화형화활회 (NS).
표주(PF) +	공기단당도망맥명목문물사상색성수시암의인장저절주중창태합 (NS).
하품(PF) +	가간객점곡과관군기단대도등막먹목무문물미반발범벽변부불사사위산상선성소술식악업연요위인임일자자학장재저저리전점점정조찬창책체춤과판형화황 (NS).
한냉(PF) -	적제귀명목별사성위절 (NS).
함박(PF) -	기대동소수이장전정조채 (NS).
보라(PF) +	적대봉사살색수에자장진차탈학해 (NS).
애벌(PF) +	디오이트일락 (NS).
응달(PF) +	저송이레목이초 (NS).
자물(PF) +	개비갈관구지래러변음박질음질인팽이필 (NS).
종달(PF) +	가감건고권량테력리망문물상성심육자정주중질체품 (NS).
정검(PF) +	개비갈관구지래러변음박질음질인팽이필 (NS).
	저맹등이류문불사색술시식안역열인정증찰침토 (NS).

ANNEXE B-2. Désambiguïsation entre PF et NS pour les données de P2

PF - NS	Etiquetage correct	PF - NS	Etiquetage correct
가장물(S)	가장물	고품명	고(PF) - 품명
가장비(S)	가장비	고품목	고(PF) - 품목
고동기(S)	고동기	고품위(S)	고(PF) - 품위
고품격	고(PF) - 품격	광대가(S)	광대(PF) - 가
고품계	고(PF) - 품계	광대역	광대
광대파	광대(PF) - 파	구공식	구(PF) - 공식
광대폭	광대(PF) - 폭	구공약	구(PF) - 공약
구공관	구(PF) - 공관	구공장	구(PF) - 공장
구공문	구(PF) - 공문	구공채	구(PF) - 공채
구공사	구(PF) - 공사	구정관	구(PF) - 정관
구정물	구정(PF) - 물	극대파	극(PF) - 대파
구정체(S)	구정(PF) - 체	극소설	극(PF) - 소설
극단가(S)	극단(PF) - 가	극소액	극(PF) - 소액
극단주	극단(PF) - 주	극한대	극 - 한대
극대조	극(PF) - 대조	급회전(S)	급(PF) - 회전
누비관(S)	누비관	대접시	대접시
단장사(S)	단장사	도요지(S)	도요지
단장정(S)	단장정	마제초(S)	마제초
단장품(S)	단장품	맹혹사	맹(PF) - 혹사
단장화(S)	단장화	모시가(S)	모(PF) - 시가
무당질(S)	무당(NS) - 질	복사도	복사(PF) - 도
배내기(S)	배내(PF) - 기	복사면	복사(PF) - 면
복사각	복사(PF) - 각	복사본	복사(PF) - 본
복사건	복사(PF) - 건	복사비	복사(PF) - 비
복사대	복사(PF) - 대	복사색	복사(PF) - 색
복사선(S)	복사(PF) - 선	복사자	복사(PF) - 자
복사시	복사(PF) - 시	복사계(S)	복사(PF) - 계
복사실	복사(PF) - 실	복사파(S)	복사(PF) - 파
복사업	복사(PF) - 업	복사기(S)	복사(PF) - 기
복사열(S)	복사(PF) - 열	복사담(S)	복사담
복사무(S)	복사무	사시도(S)	사시도
복사상(S)	복사상	상고공(S)	상고공
복사체(S)	복사체	상고대(S)	상고대
부시인(S)	부시인	상고배(S)	상고배
부시장(S)	부시장	상고사(S)	상고사

PF - NS	Etiquetage correct	PF - NS	Etiquetage correct
상고인(S)	상고인	상품명	상품(NS) -
상고모	상고모	상품목	상품(NS) -
상고참	상(PF) -고 참	상품별	상품(NS) -
상품격	상(PF) - 품격	상품성	상품(NS) -
상품계	상품(NS) -	생안경	생안경
수양모(S)	수양모	수양재	수양재
수양부(S)	수양부	양재기(S)	양재기
수양산(S)	수양산	양재사(S)	양재사
수양수(S)	수양수	역순산	역(PF) -순산
수양자(	수양자	역순위	역(PF) -순위
오목장(S)	오목장	용수난	용수(NS) -난
오지대(S)	오지대	용수비	용수(NS) -비
온난관	온난(NS) -관	용수액	용수(NS) -액
용수공	용수(NS) -공	용수정	용수(NS) -정
용수관	용수(NS) -관	용수초(S)	용수(NS) -초
우두목	우(PF) -두목	저품계	저(PF) -품계
월계수(S)	월계(PF) -수	저품명	저(PF) -품명
이종손	이종(PF) -손	저품목	저(PF) -품목
재종손(S)	재종손	전전관(S)	전전관
저품격	저(PF) -품격	전전기(S)	전전기
전전부(S)	전전부	중품격	중(PF) -품격
전전파(S)	전전파	중품계	중(PF) -품계
제비가(S)	제비가	중품명	중(PF) -품명
조약도(S)	조약도	중품목	중(PF) -품목
주마감(S)	주마감	중품별	중(PF) -품별
최단가	최단(NS) -가	투명문	투명문
최단기	최단(NS) -기	투명색(S)	투명색
축사도(S)	축사도	투명성	투명성
축사주(S)	축사주	하품격	하(PF) -품격
투명도(S)	투명도	하품계	하(PF) -품계
하품명	하(PF) -품명	함박살(S)	함박살
하품목	하(PF) -품목		
하품별	하(PF) -품별		
한냉기	한냉기		
한냉대	한냉대		

ANNEXE C-1. Les candidats de l'ambiguïté P2 entre NS et SF

Les candidats suivants de l'ambiguïté P2 sont produits automatiquement à partir de la liste de 2271 NS et 253 SF. Ces candidats sont les formes acceptables dans T1'.

감개관국귀낙천농단대매목무물미배부분비사선세시애양연염영요	가락(SF)
원인자장재조종주진찬참창첨초추촉패폐향허호호전화홍(NS)+	
감개관국귀낙천농단대매목무물미배부분비사선세시애양연염영요	가량(SF)
원인자장재조종주진찬참창첨초추촉패폐향허호호전화홍(NS)+	
감개관국귀낙천농단대매목무물미배부분비사선세시애양연염영요	가리(SF)
원인자장재조종주진찬참창첨초추촉패폐향허호호전화홍(NS)	
폐(NS)+	갱이(SF)
검과군기논대동별선소수은의전점정주집천철추(NS)+	거리(SF)
검과군기논대동별선소수은의전점정주집천철추(NS)+	거미(SF)
개경다대등목백병부사상성수정약인전주청초통투폐한핵효(NS)	과리(SF)
가강개공극내노농당도마문미방배백불사살선속식안역연영요육인	구덕(SF)
입자장장신전절정창천첨청촉추촉침탐투폐포표학함해허호화(NS)+	
가강개공극내노농당도마문미방배백불사살선속식안역연영요육인	구리(SF)
입자장장신전절정창천첨청촉추촉침탐투폐포표학함해허호화(NS)+	
가강개공극내노농당도마문미방배백불사살선속식안역연영요육인	구원(SF)
입자장장신전절정창천첨청촉추촉침탐투폐포표학함해허호화(NS)+	
가각감관국궐기막사시안암옥원이인장체홍(NS)+	내기(SF)
가각감관국궐기막사시안암옥원이인장체홍(NS)+	내미(SF)
경광극기냉대독등막무박부분사상선성세소순식악안연.의인임입	대기(SF)
장적전절접종주중차천초추촉침편포피학한함허위현화(NS)+	
가감강경과국군궐기노농다단당대도면명무무인목반복부사선세	도리(SF)
속수순시식신안애약양역영요용위.명의인장적전절정주중진집천	
촉파판편포학한해향호황회(NS)+	유
가감격골기난냉노목무미반발변부사상선소수시신악요우운이자전	동아(SF)
주진찬책청초태과폭피학함해향훈활황회(NS)+	
가감격골기난냉노목무미반발변부사상선소수시신악요우운이자전	동이(SF)
주진찬책청초태과폭피학함해향훈활황회(NS)+	
가감격골기난냉노목무미반발변부사상선소수시신악요우운이자전	동자(SF)
주진찬책청초태과폭피학함해향훈활황회(NS)+	
가리감강경기낙대도마목무병부수안에얼연.이장처천테해(NS)	마님(SF)
가리감강경기낙대도마목무병부수안에얼연.우이장처천테해(NS)+	마리(SF)
가감국당대도동만모방복불사상수안영장체화(NS)+	면피(SF)
가강골공국난동문병분사상시실쌍악안애업연열용운의임잔재전정	무리(SF)
종주집책책총판편플학해허훈(NS)+	
강결공논도두레면반선소속수.임정포함호(NS)+	박이(SF)
강결공논도두레면반선소속수.우임정포함호(NS)+	박질(SF)
가개건공국낙만모문변비사선소순악연예.이일전정처추탐토표한	방면(SF)
해향훈(NS)+	우
건담도분선성세수시정연위진집참첨초촉패(NS)+	배기(SF)
가간갑결공과광국기납농당마망매명발배백범복본부분사상선세수	부리(SF)
신악안양열요육음의인임임산자장전정조조산주중질창천첨청초포	
풍피학현화회홍(NS)+	
가간갑결공과광국기납농당마망매명발배백범복본부분사상선세수	부림(SF)
신악안양열요육음의인임임산자장전정조조산주중질창천첨청초포	
풍피학현화회홍(NS)+	
가간갑결공과광국기납농당마망매명발배백범복본부분사상선세수	부지(SF)
신악안양열요육음의인임임산자장전정조조산주중질창천첨청초포	
풍피학현화회홍(NS)+	

가간계계골공관골내도들무변불사상소염요응전정진체총호화 (NS)+	통수(SF)
가간경구미국기노문방백변부비산수신암에약우원절점참토포표 (NS)+	호지(SF)
부사자페(NS)+	활량(SF)
깨다래도미이장조토(NS)+	끼리(SF)
사이수(NS)+	다귀(SF)
사이수(NS)+	다리(SF)
무(NS)+	덜불(SF)
가운(NS)+	대기(SF)
가귀늑대만박백부선소수쌍양연염우운이자절중초태호화황(NS)+	두리(SF)
기무천(NS)+	둥이(SF)
기무천(NS)+	둥치(SF)
감강곤급미분비상소열차평꼭(NS)+	등신(SF)
때본(NS)+	때기(SF)
가리감강경기낙대도마목무병부수안에얼연우이장처천테해(NS)+	마름(SF)
유우타이(NS)+	머리(SF)
과낙상정(NS)+	오리(SF)
반사장꼭(NS)+	죽지(SF)
가가락간간감객경공골금기낙늑논농단대도화동먹명승무바라반방 백벽변분분분비사선소수실악양연영에요응세음의이인인입자	지기(SF)
장저적전전정중주죽중차처천첩초춘축털탐토통후파패팬페이인편평 폐함해현화황후간(NS)+	
가별후(NS)+	
조(NS)+	침지(SF)
건과도도동명모배변사상세수용자추해(NS)+	카락(SF)
대판부성실연위재경전참호홍(NS)+	태기(SF)
대판에(NS)+	패기(SF)
경광극기냉대독등주중중차천초주축침편포피함한합허위현화(NS)+	꾸머리(SF)
입장장적전절전절사선쇠신실열요육원인전절조책축축패피홍(NS)+	대가리(SF)
유우타이(NS)+	망태기(SF)
가가락각간간감객경공골금기낙늑논농단대도화동먹명승무바라반방 백벽변분분분비사선소수실악양연영에요응세음의이인인입자	머저리(SF)
장저적전전정중주죽중차처천첩초춘축털탐토통후파패팬페이인편평 폐함해현화황후간(NS)+	지저리(SF)
고주(NS)+	
방(NS)+	망태기(SF)
	부성이(SF)

ANNEXE C-2. Désambiguïsation entre NS et SF pour les données de P2.

NS - SF	Etiquetage correct	NS - SF	Etiquetage correct
첨가량(S)		사가리(S)	
초가량(S)		대거리(S)	
추가량(S)		등거리(S)	
대가리(S)		소거리(S)	
비가리(S)		전거리(S)	
개구녁	- 구녁(SF)	공구리	- 구리(SF)
문구녁	- 구녁(SF)	내구리(S)	
안구녁	- 구녁(SF)	방구리(S)	
창구녁	- 구녁(SF)	자구리(S)	
개구리(S)		투구리(S)	
청구원	- 원(SF)	화구원(S)	
연구원(S)		감내기(S)	
화구원(S)	- 원(SF)	사내기	
감내기(S)		장내기(S)	
화구원(S)		극대기	-기(SF)
기대기	-기(SF)	막대기(S)	-기(SF)
냉대기	-기(SF)	부대기(S)	-기(SF)
대대기	-기(SF)	분대기	-기(SF)
독대기	-기(SF)	소대기(S)	-기(SF)
등대기	-기(SF)	연대기(S)	-기(SF)
입대기	-기(SF)	현대기	
중대기	-기(SF)	가도리(S)	
편대기(S)	-기(SF)	당도리(S)	
포대기(S)	-기(SF)	신도리(S)	
함대기(S)	-기(SF)	장도리(S)	
중도리(S)		선동이(S)	
중도리(S)		수동이(S)	
변통수(S)		악동이	-이(SF)
방호지		난동자	-자(SF)
무동이	-이(SF)	노동자(S)	-자(SF)
반동자(S)	-자(SF)	진동자(S)	-자(SF)
변동자	-자(SF)	찬동자(S)	-자(SF)
선동자(S)	-자(SF)	책동자(S)	-자(SF)
이동자	-자(SF)	폭동자(S)	-자(SF)
주동자(S)	-자(SF)	피동자	-자(SF)

NS - SF	Etiquetage correct	NS - SF	Etiquetage correct
부동자(S)	-자(SF)	결박질	-질(SF)
수동자(S)	-자(SF)	논박질	-질(SF)
전동자(S)	-자(SF)	도박질	-질(SF)
소박이(S)		두레박질	-질(SF)
정박이	-이(SF)	선박질(S)	
가방면(S)		매부리	
이방면	-방면(SF)	가부리(S)	
시정배기	-배기(SF)	발부리(S)	
집배기	-기(SF)	광부림	
도배기(S)		기부지	-지(SF)
납부囚	-지(SF)	각서리	
명부지(S)	-지(SF)	각서방	-서방(SF)
발부지	-지(SF)	독서방	-방(SF)
자부지(S)	-지(SF)	비서방	-방(SF)
갱생원		대서방(S)	-방(SF)
전서방(S)		수호지(S)	
정서방(S)		반신료(S)	-료(SF)
독선생(S)		발신료	-료(SF)
갱신료	-료(SF)	전신료(S)	-료(SF)
산호지(S)		추신료	-료(SF)
통신료	-료(SF)	소이전(S)	
투신료	-료(SF)	물자위(S)	
회신료	-료(SF)	가장이(S)	-이(SF)
놀이개	-개(SF)	만장이(S)	-이(SF)
동이개	-개(SF)	목장이(S)	-이(SF)
미장이(S)	-이(SF)	종주부(S)	
염장이(S)	-이(SF)	도주사(S)	
통장이(S)	-이(SF)	진주사(S)	
경쟁이(S)	-이(SF)	부주의	
정쟁이(S)	-이(SF)	각지개	
감지개	-개(SF)	저지방	
기지개(S)		가지방(S)	
무지개		함지방(S)	
자지개(S)		강짜기	
각지방(S)		말초리	-초리(SF)

NS - SF	Etiquetage correct	NS - SF	Etiquetage correct
문초리	-초리(SF)	소치기	
식충이(S)	-이(SF)	양치기(S)	-치기(SF)
날치기(S)	-치기(SF)	자치기(S)	-치기(SF)
등치기(S)	-치기(SF)	재치기(S)	
배치기(S)	-치기(SF)	차치기	-치기(SF)
합치기		계통수(S)	
백치기(S)		폐활량(S)	
새치기(S)			
장치기(S)			
골통수			

ANNEXE D. La list de des-ambiguïtés P1 entre SF et PPN

Les données suivantes sont extraits manuellement à partir de la liste de l'ambiguïté P1 entre SF et PPN. Elles sont marquées comme NS-SF.

가공가 각본가 각색가 감독가 감상이 감시가 감식이 감정이 강가 개량가 개발가 개사가 개선가
 개울가 개척가 개혁가 건축가 검사가 검색가 검술가 검시가 검안이 검역가 검열가 검증가
 검침가 격양가 경력이 경매가 경연가 경영가 경작가 경제가 경험가 경호가 계약가 계몽가
 계약가 계획가 고문가 고시가 고용가 고유가 고전가 고지가 곡예가 공격가 공급가 공론가
 공매가 공명가
 공모가 공사가 공상이 공시가 공약가 공업가 공연가 공예가 공장가 공정이 공증가 공지가
 공출가 관상가 관찰가 관청가 관측가 광고가 광신가 광업가 괴력이 교역가 교육가 교정이
 교차가 교화가 교훈가 구매가 구원가 구입가 구제가 구조가 구직가 구축가 구출가 국사가
 국악가 국역가 국졸가 국학가 군가 권농가 권력이 권세가 권투가 궤변가 균열가 극작가 금연가
 금욕가 금융가 기고가
 기공가 기교가 기술가 기업가 기계가 기입가 기준가 기획가 길가 깎연가 낙농가 낙찰가 낙향가
 난동가 납세가 납품가 내정가 냉혈가 노동가 노력이 노예가 노획가 농경가 농민가 농성이 눈가
 능력이 다식가 다작가 달변가 답가 대농가 대변가 대소가 대역가 대체가 대출가 대필가 대학가
 대행가 도덕가 도매가 도박가 도배가 도벌가 도산가 도식가 도심가 도안가 도입가 도화가
 독립가
 독살가 독선가 독식이 독재가 독점가 독주가 독학가 동결가 동맹가 동심가 등반가 마술가
 만담가 만화가 망명가 망상가 매각가 매매가 매수가 매입가 면담가 멸문가 명가 명망가 명문가
 명상이 명언가 명창가 명필가 모반가 모사가 모색가 모시가 모험가 목동가 목축가 못가 몽상가
 무당가 무도가 무명가 무산가 무속가 무술가 무심가 무역가 무용가 무장가 문벌가 문예가
 문장가 문필가
 문학가 미신가 미장가 미필가 밀매가 박제가 번역가 반출가 발권가 발명이 발주가 발행가
 방범가 방어가 방청가 방탕가 배급가 배당가 배상가 백가 번역가 벌목가 법가 변론가 변절가
 병가 보급가 보상이 보유가 복사가 봉인가 봉합가 부대가 부력이 부속가 부식가 부양가 분석가
 분쇄가 분양가 분업가 분절가 불가 불교가 불참가 불평가 불허가 비원가 비평가 비행가 빈농가
 사대가
 사랑가 사발가 사상이 사색가 사업가 사육가 사제가 사진가 사창가 사친가 사학가 사행가 삼가
 삼가 삼화가 상담가 상속가 상식이 상업가 상점가 상품이 생명이 생산가 서도가 서예가 서화가
 석별가 선동가 선지가 선협가 설계가 섭외가 성가 성악가 성조가 세도가 세력이 소국가 소농가
 소설가 소액가 소작가 송별가 송환가 쇼핑가 수가 수저가 수색가 수송가 수심가 수완가 수용가
 수입가 수집가 수채화가 수출가 수탈가 수평가 수필가 수혈가 숙련가 순시가 순찰가 시술가
 시중가 시추가 시판가 식당가 실력이 실무가 실습가 실업가 실종가 실지가 실질이 실천가 안가
 안무가 암벽가 암흑가 애견가 애국가 애도가 애연가 애조가 애주가 애창가 애향가 애호가
 액션가 야심가 약정가 양도가 양돈가 양봉가 양산가 양상이 양생가 양식이 어부가 언론가
 여행가 역사가 역설가
 연구가 연락가 연설가 연쇄가 연주가 연출가 영농가 영양가 영업가 예술가 예시가 예언가
 예지가 예탁가 예편가 완력이 왕가 외환가 요람가 유가 유물가 운동가 운반가 운송가 운변가
 원예가 원자가 원작가 원정가 유람가 유랑가 유림가 유통가 유행가 육상가 육식이 윤락가
 은행가 음양가 응결가 응원가 의술가 의창가 이까 이농가 이론가 이별가 이주가 인권가 일가
 임대가 입양가
 입주가 입찰가 입하가 자력이 자본가 자산가 자선가 자작가 자찬가 자탄가 작명가 작사가
 장식가 재력이 재봉가 재산가 재상이 재정가 저력이 저술가 저작가 적가 적정가 전도가 전략가
 전문가 전범가 전술가 전용가 절약가 접견가 정력이 정상가 정열가 경찰가 정책가 정치가
 정탐가 제련가 제작가 제조가 조각가 조난가 조력가 조언가 종교가 종횡가 주문가 주술가
 주연가 주제가 주택가
 중심가 증권가 증산가 지도가 지식이 지정가 지평가 직류가 직행가 진료가 집필가 찬미가
 찬불가 찬송가 참전가 창업가 창작가 채권가 채무가 채식이 채집가 채취가 책가 책동가 책략가
 처가 처분가 처사가 천황가 철거가 철학가 첩보가 청춘가 청취가 초빙가 초심가 촉망가 총가
 총원가 최고가 최신가 최저가 추도가 추산가 추정가 축조가 출고가 출발가 출입가 출자가
 출하가 취득가
 침몰가 컬러가 쾌속가 탁송가 탈출가 탐방가 탐사가 탐색가 탐정이 탐험가 탐승가 택지가
 통과가 통상가 통술가 통신가 통역가 통제가 투시가 투자가 특급가 특매가 특지가 판매가

편도가 편력이 평행가 폐품가 포획가 폭등가 폭락가 표류가 표준가 풍수가 풍수가 피난가 피랍가 하늘가 하역가 하천가 학가 학도가 학원이 한복가 한정이 할인이 합작가 합창가 항해가 해변가 해설가 해안가
 해학가 행상가 행진가 향락가 혁명가 현지가 현학가 협상가 협잡가 협정가 호객가 호사가 호송가 호식이 화전가 화정가 환락가 환승가 환전가 활동가 활용가 황가 황도가 화가 화도가 홈연가

가공과 가물치와 가상과 가스과 가오리와 가자미와 가제과 가전학과 가정과 가지과 각과 간호와 갈매기와 갈치와 감리와 감시와 감식과 감쇠와 감찰과 개구리와 개미와 개발과 거머리와 거미와 거위와 건과 건설과 건축과 검사와 검색과 검시와 검역과 검열과 검증과 검찰과 검침과 거자와 격과 경제와 계획과 고고학과 고등어와 고사리와 고슴도치와 고양이고용과 곰과 공매과 공모과 공미리와 공보과
 공시과 공업과 공연과 공예과 공증과 판과 구매과 국문과 국사와 국세과 국악과 국어과 국학과 국화와 글과 귀뚜라미과 글과 기계와 기린과 기상과 기술과 기획과 까마귀과 꼴뚜기와 공치와 피꼬리와 꿩과 곤꾼이와 나리와 나방과 나비와 낙과 낙지와 낙타과 날치와 남생이와 내과 내외과 너구리와 넓치와 논과 농공과 농과 느타리와 다람쥐과 다래과 다랭이과 다슬기와 다시마과 단과 단풍과 단화과
 달팽이과 담당과 대구과 대승과 대출과 도와 도라지와 도롱뇽과 도루묵과 도미과 독수리와 돌고래과 두꺼비과 두터지고 두투미과 등어과 따개비과 따오기와 딱다구리와 뽕부기와 마취과 망과 망둥이과 매과 매미과 뿔뿔이과 멧개과 메기와 메뚜기와 멧돼지고 멸구와 멸치와 명과 모기와 모텔과 목련과 무과 무궁화와 무역과 무용과 문과 문란과 문리와 문어과 문학과 물리와 물상과 미나리와 미역과 민어과 밀과 밀봉과
 발권과 백과 백부와 백시와 백합과 뱀과 뱀어과 벌레과 벌목과 법과 배짱이과 벼과 벼룩과 벌과 병어과 보육과 보장과 복과 본과 봉선화와 부사와 부유과 부인과 부평초과 분과 불가사리와 불과 불문과 비노기와 비둘기와 비어과 비자와 비취와 빈대와 빈사와 사마귀와 사생과 시슴과 사회와 사회와 산과 산실과 산업과 산호와 살모사와 살무사와 삼치와 생강과 생원과 서무와 석류와 석이과
 선과 선인장과 성악과 세제과 소과 소라과 소아과 송백과 송사리와 송이과 수공과 수국과 수리와 수선과 수선화와 수의과 수학과 숭어과 신경과 실과 실업과 썩과 아귀과 아욱과 악어와 안과 아자와 양귀비과 양미리와 양조과 여치와 연수와 연어과 연작과 영문과 영사와 영어와 예과 예능과 예술과 오락우탄과 오리과 오미자와 오징어과 옥돔과 올빼미과 외래과 원숭이와 위와 으서과 으원과 으화과
 윤리과 음양과 의과 이파 인문과 인사와 잉어과 자라와 자분과 자연과 자유와 잠자리과 장과 장미과 전갈과 전경이과 전문과 전자와 정보과 정신과 정치과 조과 조사와 조선과 족제비과 종다리와 종사와 종친과 주목과 지네과 지빠귀과 진과 진달래과 진드기와 진디과 진료과 진사와 진찰과 질경이과 저르레기와 참별과 철쭉과 철학과 침단과 청각과 청어과 청정과 체육과 축수와 춘충과 총무과 칠면조과
 철엽수와 카멜레온과 켄거투우과 코브라과 동과 타조과 토끼과 토란과 토목과 통계과 통신과 통역과 투어과 파레과 파리와 파인에플과 편매과 평근과 편충과 포도와 표피과 풍뎡이과 피부과 하마과 학과 학무과 학생과 한의과 해과 해동과 해사와 해양과 해파리와 행정과 향과 홍어과 홍합과 화공과 화학과 환경과 회양목과

가공관 가스관 가정관 가족관 가치관 감독관 감리관 감사관 감시관 감식관 감지관 감찰관 강관 개봉관 개신관 개원관 개찰관 개표관 검사관 검색관 검시관 검안관 검역관 검열관 검인관 검증관 검찰관 검침관 검토관 결재관 결혼관 경리관 경연관 경영관 경찰관 계몽관 계수관 계엄관 계획관 고무관 고문관 고시관 고압관 고용관 공매관 공보관 공사관 공시관 공업관 공예관 과학관
 관계관 관리관 관찰관 관측관 관할관 교도관 교련관 교양관 교육관 교정관 교직관 교화관 구리관 국가관 국사관 군관 굴관 굴착관 귀관 귀빈관 근로관 금관 금수관 금유관 기구관 기념관 기록관 기상관 기술관 기업관 기포관 기획관 나무관 나팔관 내관 내세관 냉각관 냉관 냉수관 네온관 노관 노무관 노예관 뇌관 달팽이관 담당관 대독관 대사관 대청관 도덕관 도서관
 독립관 돌관 동맥관 드림관 매관 매물관 매설관 메시아관 명령관 명사관 목관 목사관 목회관 무관 무역관 문무관 문서관 문정관 문학관 문화관 미관 미태관 민족관 밀관 박물관 반응관 발광관 발사관 발행관 방역관 방적관 방진관 배관 배기관 배설관 배수관 배심관 배전관 배출관 번역관 법관 변호관 병관 병기관 보급관 보도관 보안관 보좌관 보호관 복합관 봉관
 부검관 부모관 부부관 분류관 분수관 불교관 비닐관 비서관 사교관 사랑관 사령관 사료관 사무관 사물관 사생관 사서관 사업관 사진관 사화관 산란관 삼림관 상무관 상설관 상연관

운항선 운전자선 원정선 위도선 람선 음선 물선 방선 치선 윤곽선 용기선 이민선
 이운선 이주선 이중선 이하선 인출선 일선 임신선 입사선 입양선 입자선 입하선 자력선 자매선
 잡역선 장선 재단선 재봉선 재수선 저속선 저압선 저지선 저항선 전기선 전라선 전력선 전류선
 전선 전송선 전신선 전압선 전용선 전위선 전이선 전자선 전철선 전투선 전화선 전개선 절선
 절연선
 점선 점경선 점속선 점착선 점촉선 정박선 정지선 정찰선 정탐선 제련선 제어선 제작선
 조난선 조사선 조준선 종선 주사선 주유선 주파선 주행선 중간선 중심선 중앙선 증기선 지도선
 지평선 직각선 직류선 직행선 진로선 진주선 진해선 차단선 차양선 착륙선 참치선 채집선
 채취선 천이선 철거선 첩보선 초선 측각선 쇄고선 최상선 최선선 최전선 추적선 출발선 출격선
 측후선
 침몰선 패속선 패속선 탁송선 탄소선 탈선 탈출선 탈피선 탐방선 탐사선 탐색선 탐지선 탐험선
 탐승선 토사선 통과선 통진선 통운선 통제선 통행선 통화선 투사선 투시선 투영선 특급선
 판매선 판복선 편도선 평행선 평화선 폐색선 폐품선 포경선 포선 포위선 포획선 폭등선 표류선
 표선 표적선 피난선 피랍선 피복선 피부선 피아노선 하역선 하행선 한계선 한정선 함몰선 합선
 합작선
 항공선 항포선 항해선 해골선 해안선 해적선 해태선 혁미선 호르몬선 호송선
 합판선 혼선 혼용선 화물선 화개선 회귀선 회유선 회조선 후송선 혼련선 휴전선 흉선
 호위선 홀수선

건아 고립아 기린이 냉혈아 만삭아 망아 문란아 문제아 반향아 방탕아 불구아 비만아 사산아
 성숙아 야생아 양아 람아 장애아 재능아 정박아 정상아 지체아 천재아 촉망아 출산아 출생아
 태아 태아 폐물아 풍운아 학대아 혼혈아

가결의 개업의 개원의 구역의 무성의 소아의 수련의 실습의 양의 임상의 전문의 전문의 지구의
 천구

간난이 거북이 겸댕이 겨드랑이 갑장이 겸정이 누렁이 능청이 맹추이 빨간이 식충이 열륙이
 염장이 우렁이 울이 이월이 지은이 황진이

ANNEXE E. Classes de SFR dans le verbe

CLASSE	SFR	Information
1 감다	감	SU3 : 1 2 3 4 5 7 8 12 13 14 17 18 23 24 25 26 SU1 : 5 SU2 : 1 2 3 4 5 SU4 : 1 PV1 : 1 3 PV3 : 1 PV4 : 1
2 녹다	녹	SU3 : 1 2 3 4 5 7 8 12 13 14 17 18 23 24 26 SU1 : 5 SU2 : 1 2 3 4 5 SU4 : 1 PV1 : 1 3 PV3 : 1
3 익다	익	SU3 : 1 2 4 5 7 8 9 10 12 13 14 15 17 21 23 24 SU2 : 1 2 3 4 5 SU4 : 1 PV4 : 1 SU1 : 5 6 PV1 : 1 3
4 듣다	듣	SU3 : 1 2 4 5 7 8 9 10 12 13 14 15 17 21 23 24 SU2 : 1 SU4 : 1 SU1 : 5 6 PV1 : 1 3
	들	SU2 : 2 3 4 5 PV4 : 1
5 곱다	고	PV6 : 1 2 3 5 6 PV8 : 1
	곱	SU1 : 4 5 SU3 : 1 2 3 4 5 6 7 8 9 10 11 12 14 15 16 17 18 26 SU2 : 1 SU4 : 1 PV1 : 1
6 읊다	오	SU3 : 8 12 SU1 : 5 6 9 15 16 17 18 PV1 : 1
	읍	SU3 : 1 2 4 5 7 14 17 18 21 23 26 SU1 : 6 9 13 SU4 : 1 SU2 : 3 4 5
7 갈다	가	SU3 : 8 9 12 13 20 SU1 : 5 6 12 13 14 15 16 17 18 PV1 : 1 3
	간	SU1 : 1
	갈	SU1 : 2 9 10 11 SU2 : 2 3 4 5 SU3 : 1 2 3 4 5 6 7 14 15 16 17 18 23 24 26 PV3 : 1 PV7 : 1 2

8 낳다	나	SU3 : 1 4 7 8 SU1 : 5 6 11 12 13 15 16 17 18 19 PV1 : 1 7 8 9 PV3 : 1
	난 날 낳	SU1 : 1 SU1 : 2 SU3 : 1 2 3 4 5 7 8 9 12 13 14 15 17 18 21 23 26 SU1 : 5 6 SU2 : 1 2 3 4 5 SU4 : 1 PV1 : 1 3 8 9 PV3 : 1
9 낳다	나	SU1 : 6 7 12 13 15 16 17 18 PV3 : 1 SU3 : 4 5 7 17 18 PV1 : 7
	난 날 낳	SU1 : 1 SU1 : 2 SU3 : 1 2 3 4 5 6 7 8 9 10 12 13 14 15 17 18 26 SU1 : 5 SU2 : 1 2 4 5 SU4 : 1 PV1 : 1 7
10 뵈다	뵈 뵈 뵈	SU1 : 1 SU1 : 2 SU3 : 1 2 4 5 7 8 12 13 17 SU1 : 5 6 7 8 9 15 16 17 18 PV4 : 1 PV1 : 1 3 PV7 : 1
	건 거	SU1 : 1 SU3 : 8 12 13 20 SU1 : 5 6 12 15 16 17 18 PV1 : 1 2 3
11 결다	결	SU1 : 2 7 8 9 10 11 SU2 : 2 3 4 5 SU3 : 1 2 3 4 5 7 14 15 17 18 23 26 PV1 : 5 PV4 : 1 PV7 : 1
	순 술 술	SU1 : 1 SU1 : 2 SU3 : 1 2 3 4 5 7 8 12 13 14 15 17 18 19 SU2 : 1 SU4 : 1 SU1 : 5 6 PV1 : 1 3
12	숫다	SU1 : 9 19 11 12 13 15 16 17 18 23 SU3 : 20 21 PV4 : 1
	수	

13 넣다	넌 널 너	SU1 : 1
		SU1 : 2
		PV1 : 5
		SU1 : 7 8 9
	넣	PV7 : 1
		SU3 : 1 2 3 4 5 6 8 12 13 14 17 18 23 26
		SU1 : 5 6 9 10 11 12 13 15 16 17 18
		SU2 : 1 2 3 4 5
		SU4 : 1
		PV4 : 1
		PV1 : 1 2 3 7 8
		PV7 : 1
		SU1 : 1
		SU1 : 2
		SU3 : 1 2 4 5 7 8 12 13 14 17 18 23 26
		SU1 : 5 6
		PV1 : 1 3
14 굿다	근 글 굿	SU1 : 9 10 11 12 13 15 16 17 18
		SU3 : 20 21
		SU2 : 2 3 4 5
		PV4 : 1
	그	SU1 : 1
		SU1 : 2
		SU3 : 1 2 3 4 5 6 7 8 9 12 14 15 17 18 23 24 26
		SU1 : 5 6 9 11 12 13 14 15 16 17 18
		PV1 : 1
		PV7 : 1
15 가르다	가른 가를 가르	SU3 : 1 4 5 6 7 8 9 12 21 26
		SU1 : 5 6 12 13 14
		SU2 : 3 4 5
		PV1 : 1
	갈랐	PV4 : 1
		SU1 : 1
		SU1 : 2
		SU3 : 1 2 4 5 7 8 12 13 14 17 18 19 20 21 26
		SU1 : 5 6 7 8 9 10 11 12 13 15 16 17 18
		SU2 : 2 3 4 5
		PV3 : 1
		PV5 : 1
		PV1 : 1 3 5
		PV7 : 1
16 꼬다	꼰 꼰 꼬	SU3 : 1 4 5 7 12 14 15 17 18 21 26
		SU2 : 1 2 4 5
		SU4 : 1
		PV4 : 1
	꼴	SU1 : 1
		SU1 : 2
		SU3 : 1 2 4 5 7 8 12 13 14 17 18 19 20 21 26
		SU1 : 5 6 7 8 9 10 11 12 13 15 16 17 18
		SU2 : 2 3 4 5
		PV3 : 1
		PV5 : 1
		PV1 : 1 3 5
		PV7 : 1
17 개다	갸 갸 개	SU3 : 1 4 5 7 12 14 15 17 18 21 26
		SU2 : 1 2 4 5
		SU4 : 1
		PV4 : 1
	갸	SU1 : 1
		SU1 : 2
		SU3 : 1 2 4 5 7 8 12 13 14 17 18 19 20 21 26
		SU1 : 5 6 7 8 9 10 11 12 13 15 16 17 18
		SU2 : 2 3 4 5
		PV4 : 1
		PV5 : 1
		PV1 : 1 3 5 8
		PV7 : 1
		SU3 : 1 4 5 7 12 14 15 17 18 21 26
		SU2 : 1 2 4 5
		SU4 : 1
		PV4 : 1

18 굽다	구	PV1 : 5
		SU1 : 7 8 9 10 11 12 13 15 16 17 18
		SU2 : 2 3 4 5
		PV4 : 1
		PV6 : 1 2 3 4 5 6
		PV7 : 1
		SU1 : 1
	군 굴 굽	SU1 : 2
		SU3 : 1 2 3 4 6 7 8 9 10 11 12 13 14 15 17 18 23 26
		SU2 : 1
		SU4 : 1
		SU1 : 5 6
		PV1 : 1 3
19 가다	간 갈 가	SU1 : 1
		SU1 : 2
		SU3 : 1 2 3 4 5 7 8 9 10 12 13 14 15 17 18 22 23 24 26
		SU1 : 5 6 7 9 10 12 13 15 16 17 18 19
		PV1 : 1 3 7 8 9
		PV7 : 1
		SU3 : 1 3 4 5 7 8 12 13 14 15 17 18 21 26
	갓	SU1 : 5 6 13
		SU2 : 1 2 4 5
		SU4 : 1
		PV4 : 1
		PV1 : 1 3
20 읊다	우	SU3 : 1 4 5 8 12 13 17 18 29 21 26
		SU1 : 5 6 9 10 11 12 13 15 16 17 18 19 20
		PV7 : 1
		PV1 : 1 3
		SU1 : 2
		SU1 : 1
		SU3 : 1 2 4 5 7 8 9 12 13 14 17 18 26
	울 운 읍	PV1 : 1 3
		SU1 : 5 13 15
		SU1 : 5 6
		SU2 : 1 2 4 5
		SU4 : 1
		SU3 : 1 2 3 4 6 7 8 9 12 14 15 16 17 18 26
		PV1 : 1 2
21 하다	하	SU3 : 1 2 3 4 5 7 8 10 11 12 13 14 15 17 18 21 23 24 26
		SU1 : 5 6 8 9 10 11 12 13 15 16 17 18 19 20
		PV5 : 1
		PV1 : 1 3
		PV7 : 1
		SU1 : 1
		SU1 : 2
	한 할 해 했	PV1 : 7
		SU3 : 1 2 3 4 5 7 8 9 10 12 13 14 15 16 17 18 19 26
		SU1 : 5 6
		PV1 : 1 3 7
		PV4 : 1
		SU2 : 2 4 5
99 기타		

ANNEXE F. Classes de SFR dans l'adjectif

CLASSE	SFR	Information
1 감다	감	SU3 : 1 2 4 5 7 8 9 12 14 17 18 23 24 26 SU1 : 5 6 SU2 : 1 2 3 4 5 SU4 : 1 PA1 : 1 3 4 PA3 : 1 PA4 : 1
2 녹다	녹	SU3 : 1 2 4 5 7 8 9 12 14 17 18 23 24 26 SU1 : 5 6 SU2 : 1 2 3 4 5 SU4 : 1 PA1 : 1 3 4 PA3 : 1
3 검다	검	SU3 : 1 2 3 4 5 6 7 8 9 12 14 15 17 18 21 26 SU1 : 4 5 6 SU2 : 2 3 4 5 PA1 : 1 PA4 : 1
4 없다	없	SU3 : 1 2 3 4 5 6 7 8 9 12 14 15 17 18 21 24 SU1 : 3 4 5 6 19 20 SU2 : 1 2 4 5 SU4 : 1 PA1 : 1 3 PA4 : 1
5 갑다	가	PA6 : 1 2 3 5 6 PA8 : 1
	갑	SU1 : 4 5 SU3 : 1 2 3 4 5 6 7 8 9 10 11 12 14 15 16 17 18 26 SU2 : 1 SU4 : 1 PA1 : 1
6 돕다	두	PA6 : 1 2 3 5 6 SU1 : 4 5
	돕	SU3 : 1 2 3 4 5 6 7 8 9 10 11 12 14 15 16 17 18 26 SU2 : 1 SU4 : 1 PA1 : 1
7 달다	다	SU1 : 5 6 SU3 : 8 12
	단	SU1 : 1
	달	SU1 : 2 SU2 : 2 3 4 5 SU3 : 1 2 4 5 7 15 17 18 26 PA3 : 1
8 걸다	거	SU1 : 5 6 9 11 12 13 14 15 16 17 18 19 SU3 : 12 20 21 22 PA1 : 1 3
	건	SU1 : 1
	걸	SU1 : 2 9 11 SU2 : 2 3 4 5 SU3 : 1 2 3 4 5 6 7 14 15 17 18 23 24 26 PA4 : 1

9 되다	되	SU3 : 1 2 3 4 5 7 8 9 12 14 15 17 18 26 SU1 : 5 6 7 9 10 11 12 13 14 15 16 17 18 PA1 : 1 3 5
	된	SU1 : 1 SU3 : 14 15 PA1 : 4
	될	SU1 : 2
10 기다	기	SU3 : 1 2 3 4 5 6 7 8 9 10 12 14 17 18 20 21 23 24 26 SU1 : 5 6 9 11 12 13 15 16 17 18 PA1 : 1 3 PA4 : 1 PA6 : 1 2 3 4 5 6
	긴	SU1 : 1
	길	SU1 : 2
	졌	SU3 : 1 4 5 6 7 8 9 10 12 14 15 17 18 23 24 26 SU1 : 5 SU2 : 1 2 4 5 SU4 : 1 PA1 : 1 PA4 : 1
	나	SU3 : 1 2 3 4 5 7 8 9 10 12 14 15 17 18 21 26 SU1 : 5 6 13 14 15 16 17 18 SU2 : 2 3 4 5 PA1 : 1 3 8 9 PA3 : 1
	난	SU1 : 1
	날	SU1 : 2
	났	SU3 : 1 4 5 6 7 8 9 10 12 14 15 17 18 21 23 26 SU1 : 5 13 14 19 SU2 : 1 2 4 5 SU4 : 1 PA1 : 1 PA4 : 1
12 차다	차	SU3 : 1 2 3 4 5 6 7 8 9 10 12 14 17 18 21 26 SU1 : 5 6 8 9 10 11 12 13 15 16 17 18 SU2 : 3 4 5 PA1 : 1 3 5 7 8 PA3 : 1 PA8 : 1
	찬	SU1 : 1
	찰	SU1 : 2
	찻	SU3 : 1 4 5 6 7 8 9 10 12 14 15 17 18 21 26 SU1 : 5 13 14 19 SU2 : 1 2 SU4 : 1 PA1 : 1 PA4 : 1
	쌔	SU1 : 6
	쌔	SU1 : 1
13 켜다	켜	SU1 : 2
	켜	SU3 : 1 4 5 6 7 8 12 14 17 18 26 SU2 : 1 2 4 5 SU1 : 5 19 PA1 : 1 3 SU4 : 1 PA4 : 1
	켜	

14 낮다	나	SU1 : 6 9 12 13 15 16 17 18 PA3 : 1 SU3 : 4 5 7 17 18 PA1 : 7
	난	SU1 : 1
	날	SU1 : 2
	낮	SU3 : 1 2 3 4 5 6 7 8 9 10 12 14 15 17 18 26 SU1 : 5 SU2 : 1 2 4 5 SU4 : 1 PA1 : 1 7
15 하다 (동사와 같음)		
16 그르다	그르	SU3 : 1 2 3 4 5 6 7 8 9 12 14 15 17 18 23 24 26 SU1 : 5 6 9 11 12 13 14 15 16 17 18 PA1 : 1
	그른	SU1 : 1
	그를	SU1 : 2
	글렀	SU3 : 1 4 5 6 7 8 9 12 21 26 SU1 : 5 6 12 13 14 SU2 : 3 4 5 PA1 : 1 PA4 : 1
17 짝다	가	SU3 : 8 17 SU1 : 5 6 PA1 : 1
	간	SU1 : 1
	갈	SU1 : 2
	짱	SU3 : 1 2 3 4 5 6 7 8 9 10 12 14 15 21 22 26 SU1 : 3 4 5 6 19 SU2 : 1 SU4 : 1 PA1 : 1
	갸	SU3 : 1 4 6 7 8 9 10 12 14 15 17 18 21 26 SU1 : 5 19 SU2 : 1 2 4 5 SU4 : 1 PA1 : 1 PA4 : 1
18 ETC.		

