

Université Paris 7

**Centre d'études et de recherches
en informatique linguistique**

Eric Laporte

13, rue Charles-Friedel

Paris 20^e

Université Paris 7

Section 24-01, emploi 1906

n° 5

**METHODES ALGORITHMIQUES ET LEXICALES
DE
PHONETISATION DE TEXTES**

Applications au français

Eric LAPORTE

_ Volume I

Thèse de doctorat en informatique
Directeur de thèse : Maurice GROSS
31 mai 1988

Jury :

M. NIVAT
R. CARRE
M. GROSS
J.S. LIENARD
D. PERRIN
M.P. SCHUTZENBERGER

Université Paris 7
Centre d'études et de recherches
en informatique linguistique

METHODES ALGORITHMIQUES ET LEXICALES
DE
PHONETISATION DE TEXTES

Applications au français

Eric LAPORTE

Volume I

Thèse de doctorat en informatique
Directeur de thèse : Maurice GROSS
31 mai 1988

Jury :

M. NIVAT
R. CARRE
M. GROSS
J.S. LIENARD
D. PERRIN
M.P. SCHUTZENBERGER

Remerciements

Mes remerciements vont d'abord au professeur Marcel-Paul Schützenberger. Sa présence dans le jury de soutenance m'honore et évoque pour moi la permanence d'une tradition scientifique et humaine dans l'informatique.

Je remercie le professeur Maurice Nivat de s'être intéressé de près à mes travaux et d'avoir accepté de présider le jury. Je suis également reconnaissant au professeur Dominique Perrin d'avoir bien voulu faire partie du jury.

Je tiens à remercier Jean-Sylvain Liénard, directeur du LIMSI, pour l'accueil qui m'a été réservé en 1987-1988 dans son laboratoire, pour l'appui dont j'ai bénéficié au moment de la rédaction de la thèse, et aussi pour ses ultimes recommandations en matière de bibliographie.

Je suis heureux que le professeur René Carré ait bien voulu participer à la soutenance, et je l'en remercie.

Que les membres du LADL et du CERIL soient remerciés de m'avoir initié à l'informatique linguistique et de m'avoir donné de précieux conseils, et notamment : Christian Leclère, qui a fait une lecture particulièrement attentive de mon texte, Laurence Danlos, Alain Guillet, Jean-Paul Boons, Annie Meunier, Mireille Piot, Gaston Gross, Robert Vivès, Morris Salkoff, Blandine Courtois, Max Silberstein, Jean-Roger Vergnaud, Jacqueline Giry-Schneider, ainsi que Gaby Klarsfeld, qui participé à la saisie des données informatiques avec générosité. Je remercie aussi Amid Nehmé, Mirella Conenna, Elisabeth Paul et Françoise Emerard pour leur aide et leur soutien.

Enfin, j'ai le plus grand plaisir à dire ici ma reconnaissance envers le professeur Maurice Gross. Il m'a formé et accompagné dans ce chemin difficile, il m'a accordé sa confiance et a été là à toutes les étapes.

Alphabet phonétique

Dans les transcriptions, nous utilisons certains symboles spéciaux avec une valeur qui est, chaque fois que possible, la même que dans l'alphabet phonétique international. Voici les principales conventions qui peuvent poser un problème de lecture :

[ã] représente la voyelle nasale de *sentir* [sãtir] ;

[e] le é fermé de *céder* [sede] ;

[ɛ] le è ouvert de (*il*) *cède* [sɛd] ;

[ɛ̃] la voyelle nasale de *plainte* [plɛ̃t] ;

[g] le g de *gare* [gar], y compris dans [ge] *gué* ;

[G] le ng de *ring* [riG] ;

[j] le y de *yaourt* [jaurt] et le ll de *filles* [fij]. Nous l'utilisons également dans des mots d'origine étrangère, tels que *iceberg* lorsqu'il est prononcé [ajsberg] et *geisha* lorsqu'il est prononcé [gejSa]. En effet, on ne sait pas distinguer d'une manière reproductible ces séquences [aj] et [ɛj] de celles de *hail* [baj] et *veille* [vej] ; de plus, cette notation reflète la différence entre [aj] et la séquence dissyllabique [ai] de *hair* [air], ou entre [ɛj] et [ɛi] dans *pays* [pɛi].

[n] représente le n de *cane* [kan] ;

[N] le gn de *signe* [siN] ;

[o] le o fermé de *peau* [po] ;

[ɔ] le o ouvert de *sol* [sɔ] ;

[õ] la voyelle nasale de *conte* [kõt] ;

[O] le eu fermé de *jeu* [ZO] ;

[œ] le eu ouvert de *jeune* [Zœn] ;

[r] le r de *ré* [re], bien que ce symbole phonétique corresponde au r roulé apical dans l'alphabet phonétique international. Le signe phonétique [ʁ] aurait mieux représenté le r du français standard, qui est généralement une fricative uvulaire, mais il aurait posé des problèmes informatiques et typographiques.

[s] représente le s de *sel* [sɛl], y compris dans [pase] *passer* ;

[S] le ch de *chat* [Sa] ;

[u] le ou de *flou* [flu] ;

[w] le ou de *fourine* [fwɪn]. Nous l'utilisons également dans la séquence monosyllabique [aw] pour des onomatopées et des mots d'origine étrangère, comme *out* lorsqu'il est prononcé [awt], en une syllabe. Cette notation rend explicite la distinction entre [aw] et le [au] dissyllabique de *Raoul* [raul] ; elle est parallèle à la notation [aj], évoquée ci-dessus à propos de [j].

[y] représente le u de *sur* [syr] ;

[ɥ] le u de *lui* [lɥi] ;

[z] le z de *gaz* [gaz] ;

[Z] le j de *jus* [Zy].

Alphabet phonémique

Les transcriptions phonétiques sont insuffisantes pour certains usages. Nous utilisons alors des transcriptions phonémiques exprimées dans un alphabet légèrement différent du précédent. Pour des raisons de commodité technique, cet alphabet phonémique ne comprend que des caractères représentés dans le code ASCII strict. Voici les principales conventions qui peuvent poser des problèmes à la lecture :

- /e/ représente le *é* fermé de *céder* /sede/ comme le *è* ouvert de (*il*) *cède* /sed/ ;
- /g/ le *g* de *gare* /gar/, y compris dans /ger/ *guerre* ;
- /G/ le *ng* de *ring* /riG/ ;
- /h/ signale un *h* aspiré, c'est-à-dire une interdiction de la liaison et de l'élision, comme dans *hérisser* /herise/ ;
- /j/ représente le *y* de *yaourt* /jaurt/ ;
- /n/ se lit d'une façon spéciale :
 - (1) /an/, /en/, /on/, etc., placés devant consonne, représentent les voyelles nasales de *sentir* /santir/, *plainte* /plent/ et *conte* /kont/ ;
 - (2) /an*/, /en*/, /in*/, /on*/, /yn*/ représentent les voyelles nasales finales de *plan* /plan*/, *mien* /mien*/, *fin* /fin*/, *bon* /bon*/ et *commun* /komyn*/ ;
 - (3) /an/, /en/, /in/, /on/, /yn/, placés à la fin d'un mot, représentent les finales de *panne* /pan/, *benne* /ben/, *mine* /min/, *tonne* /ton/, *lune* /lyn/.
- /N/ représente le *gn* de *signe* /siN/ ;
- /o/ le *o* fermé de *peau* /po/ comme le *o* ouvert de *sol* /sol/ ;
- /O/ le *eu* fermé de *jeu* /ZO/ comme le *eu* ouvert de *jeune* /ZOɲ/ ;
- /s/ le *s* de *sel* /sel/, y compris dans /pase/ *passer* ;
- /S/ le *ch* de *chien* /Sien*/ ;
- /u/ le *ou* de *flou* /flu/ ;
- /y/ le *u* de *sur* /syr/ ;
- /z/ le *z* de *gaz* /gaz/ ;
- /Z/ le *j* de *jus* /Zy/.

Cinq phonèmes spéciaux s'ajoutent à cette liste.

- (1) /*/, placé après une consonne à la fin d'un mot, signale l'effacement de cette consonne : *sot* /sot*/, *grand* /grand*/, *gras* /gras*/, *gris* /griz*/, *minier* /minier*/...
- (2) /:/, placé après un /o/ ou un /O/, signale que ce /o/ ou ce /O/ est fermé dans un contexte où ces phonèmes sont généralement ouverts : *jaune* /ZO:n/, *meute* /mO:t/.

(3) / + /, placé après /i/, /u/ ou /y/, signale que ces voyelles peuvent avoir une valeur syllabique : *lier* /li+e/, *rouage* /ru+aZ/, *fluide* /fly+id/ (chacun de ces trois mots peut se prononcer en deux syllabes ou en une syllabe).

(4) /' / et /" / représentent deux variétés de *e* dits muets : *sonnerie* /son'ri/, *pelouse* /p"luz/.

Introduction

La langue naturelle, et en particulier la parole, est appelée à jouer un rôle de premier plan dans la communication entre l'homme et la machine. Ses utilisations industrielles constituent un objectif important, pour lequel des moyens techniques considérables sont mis en jeu dans des laboratoires d'informatique publics et privés.

Dans l'état actuel des techniques, l'industrie exploite la parole enregistrée et la parole compressée, qui sont adaptées à quelques applications, mais n'emploie pas encore la parole synthétique. Celle-ci sera cependant utile dans des applications telles que le terminal à sortie vocale, qui aura à produire des messages variés et nombreux. On sait maintenant que la variété et la souplesse dans l'intonation, dans les constructions syntaxiques, dans les mots utilisés, sont des conditions préalables pour rendre commercialisables les applications de ce type. Or ces trois problèmes sont des défis pour la recherche fondamentale.

En ce qui concerne la reconnaissance de la parole, plusieurs systèmes commercialisés reconnaissent des langages de commande qui comportent quelques dizaines d'éléments. Les systèmes de reconnaissance de la parole continue auront à prendre en compte une variété de constructions et de mots sans commune mesure avec ces langages de commande. Cela demandera, entre autres, des connaissances formalisées sur les principales variations phonétiques des mots, car la variabilité de la parole est évidemment un problème crucial pour ce type d'applications.

Les éléments d'un système de traitement automatique de la parole, tel qu'on peut les prévoir d'après les travaux actuels, manipuleront les mots ou les textes sous différentes formes : le signal de parole ; le résultat d'une analyse spectrale de ce signal ; des chaînes phonétiques ; des chaînes orthographiques... Cette liste de niveaux n'est d'ailleurs pas limitative. Ces niveaux correspondent à des domaines de recherche complémentaires, qui ne peuvent être totalement indépendants.

A un certain niveau, les mots sont représentés sous forme de chaînes de caractères, phonétiques et orthographiques. Les problèmes de transcription sont parmi les plus connus dans ce domaine. L'orthographe et les chaînes phonétiques sont deux modes de représentation des mots et on est amené à passer de l'un à l'autre. Ainsi, une des étapes de la reconnaissance de la parole est de retrouver automatiquement l'orthographe des textes lorsqu'ils ne sont connus que sous la forme de chaînes phonétiques. On rencontre de sérieuses difficultés pour trouver des solutions réalistes à ce problème. Toutefois, le problème inverse ne nous a pas paru hors d'atteinte : transcrire phonétiquement des textes orthographiques (phonétisation automatique). Ce problème est connu dans le traitement automatique de la parole. Il a également des applications dans le domaine de la correction orthographique. Rappelons que le problème de la correction est crucial : la fiabilité de la communication entre l'homme et la machine est conditionnée par l'existence de procédures efficaces de correction des textes, le plus souvent entrés de façon approximative.

L'objectif de notre travail était de constituer des méthodes, d'élaborer des outils informatiques et d'obtenir des résultats sur ce problème de transcription phonétique. De nombreux algorithmes de phonétisation automatique ont été implantés et diffèrent par leurs performances. Mais l'utilisation industrielle d'un algorithme de phonétisation impliquerait des contraintes importantes : d'une part, une contrainte de fiabilité ; d'autre part, les transcriptions phonétiques obtenues doivent être utilisables par le système auquel elles sont transmises, quel qu'il soit. Des algorithmes de phonétisation ne pourront être exploités à grande échelle que s'ils satisfont à ces contraintes.

Or, dans le cas du français, comme d'ailleurs de l'anglais, la phonétisation automatique se heurte à trois types de difficultés.

1. Le problème algorithmique consiste à effectuer la transcription malgré les ambiguïtés de l'orthographe française. Formellement, la phonétisation automatique est une transduction, c'est-à-dire une application d'un ensemble de chaînes dans un ensemble de chaînes. Les méthodes de la théorie des langages formels lui sont donc applicables : ainsi, la transduction peut être effectuée soit par un système de règles, soit par un automate qu'on appelle alors un transducteur, les deux solutions étant équivalentes en puissance. En français, la transduction est locale dans une certaine mesure : cela facilite sa description sous forme d'un système de règles. Cette heureuse propriété a toutefois ses limites, car les règles locales ont des exceptions. Nous verrons que le problème algorithmique se ramène à déterminer ces limites et à les prendre en compte. L'enjeu est la qualité du traitement des exceptions, et donc la fiabilité du phonétiseur. Nous avons considéré les règles et les exceptions qui se dégagent de l'examen de 64.000 entrées de dictionnaire du français courant, et nous avons écrit en langage C un logiciel de phonétisation.

2. Au problème algorithmique s'ajoute un problème purement descriptif : il n'existe pas de descriptions formalisées, précises et exhaustives sur la prononciation des mots français, même si l'on s'en tient à l'usage courant dans le français standard contemporain. En effet, les seules études sur le lexique disponibles sont constituées par les transcriptions phonétiques des dictionnaires de langue, français et bilingues, ainsi que celles des dictionnaires phonétiques, qui ne vont pas plus loin dans leurs analyses. Aussi des données beaucoup plus complètes sur les variations phonétiques sont-elles indispensables. Nous avons apporté des éléments de solution à ces besoins en constituant de façon semi-automatique un dictionnaire électronique consacré à la prononciation, le DELAP¹, qui comporte 64.000 entrées. Un programme de flexion appliqué à ces entrées les développe en plus de 500.000 formes.

3. Le troisième problème est celui de la représentation formelle à donner à toutes ces informations phonétiques, et donc aux sorties du phonétiseur, ainsi qu'aux articles du DELAP. Les variations phonétiques s'ajoutent aux variations morphologiques pour former un système complexe dans lequel, étant donné la variété des mots à considérer, on trouve des cas fréquents et des cas plus isolés. Les phonologues se sont consacrés à ces problèmes linguistiques, mais sans l'outil

¹ Dictionnaire Electronique du LADL pour la Phonémique.

informatique que nous pouvons utiliser aujourd'hui. Le dictionnaire DELAP et les programmes associés nous ont permis d'aborder des problèmes de représentation formelle avec à notre disposition des données concrètes et explicites. En effet, les informations pertinentes à une question donnée figurent dans les listes de mots concernés, et ces listes peuvent être extraites automatiquement du DELAP. Comme un dictionnaire de cette taille recèle une masse d'informations importante, nous avons utilisé un langage d'interrogation qui combine le langage C et les fonctionnalités de l'éditeur programmable EDT.

Cette thèse est organisée en six parties.

Le chapitre I expose les méthodes de phonétisation de textes mises au point au cours des vingt dernières années, et présente les applications pour lesquelles elles ont été conçues.

Le chapitre II introduit le cadre général de notre travail : le système de données linguistiques du LADL. Le phonétiseur et le DELAP, qui en font partie, ont été développés parallèlement et progressivement par des méthodes qui seront exposées.

Le chapitre III présente le DELAP et plusieurs des programmes associés : les programmes de comparaison et de contrôle qui servent à l'entretien et au développement du dictionnaire, et le langage d'interrogation qui permet d'extraire du DELAP les informations pertinentes à une question donnée.

Le chapitre IV décrit le système de flexion automatique qui traite les variations morphologiques et permet de produire le dictionnaire de formes fléchies.

Le chapitre V concerne les variations phonétiques systématiques. En particulier, la description et la représentation formelle d'un ensemble de variations phonétiques sont traitées en détail.

Le chapitre VI est consacré aux applications informatiques : la production automatique de textes phonétiques, la correction orthographique, la synthèse de la parole et la reconnaissance de la parole.

Le second volume regroupe les annexes : l'algorithme de phonétisation ainsi que des listes de mots extraites du DELAP.

Chapitre I. Les méthodes en phonétisation automatique

Ce chapitre est destiné à introduire les concepts fondamentaux ainsi que les principales méthodes et les données qui ont été utilisées pour la phonétisation automatique de textes.

Un des objectifs de la phonétisation automatique est de construire des systèmes qui permettront des entrées et sorties vocales en français et en d'autres langues. Soulignons que la phonétisation automatique n'est utile que si les entrées et sorties vocales qui sont visées sont exprimées en langue naturelle, et non dans un langage limité défini par des conventions spécifiques à un système. Cette restriction exclut de nos investigations les systèmes à commande vocale qui reconnaissent un ensemble de quelques dizaines de commandes prononcées par l'utilisateur : dans le cas de ces systèmes, les transcriptions phonétiques utilisées sont suffisamment peu nombreuses pour pouvoir être élaborées à la main. Quant aux sorties vocales, nous excluons, pour les mêmes raisons, les systèmes capables de produire un jeu de quelques dizaines de messages oraux. Les sorties vocales en langue naturelle constituent un objectif plus ambitieux : produire automatiquement des messages oraux nombreux et variés, de sorte que leur contenu sémantique puisse être complexe et précis. C'est à cet objectif que nous faisons référence lorsque nous parlons de synthèse de la parole, car il suppose la production automatique des textes phonétiques et l'utilisation de la parole artificielle. Remarquons toutefois que la phonétisation automatique a d'autres applications à plus court terme, comme l'aide à la correction orthographique.

1. Les débuts de la phonétisation automatique

Les premiers logiciels de phonétisation automatique de textes en français furent conçus pour être incorporés dans des systèmes de synthèse de la parole à partir de textes. Ainsi, au LIMSI, le phonétiseur de D. Teil (1969), associé à un synthétiseur de parole par diphtongues, permet de constituer une machine à lire, à laquelle on pouvait soumettre des textes orthographiés et qui en donnait une version en parole synthétique (D. Teil, 1974). Lors de ces premières réalisations, la transcription phonétique semblait facile à automatiser. On prévoyait, à terme, l'apparition de logiciels simples capables de produire automatiquement, à partir de textes écrits, des transcriptions phonétiques qui puissent faire l'objet d'une synthèse vocale, et ce avec une qualité suffisante pour permettre la commercialisation de ces systèmes. En ce qui concerne l'étape de la phonétisation, on pensait pouvoir tirer parti des régularités observées dans la correspondance entre l'orthographe et la prononciation, et compenser ainsi les irrégularités et les ambiguïtés de cette correspondance.

Ainsi, le phonétiseur de D. Teil (1969) comportait environ 200 règles locales du type *oi* → [ua]. Par "règles locales", nous entendons celles qui permettent de transcrire phonétiquement une lettre ou une séquence de quelques lettres, en tenant compte éventuellement de leur contexte droit et gauche. Ainsi, mis à part un certain nombre de cas irréguliers, la lettre *g* peut être transcrite [Z] devant *e, é, è, ê, i, y*, comme dans *gel*,

et [g] dans la plupart des autres contextes, comme dans *gare*. En français, le contexte à prendre en compte pour exprimer ces règles générales ne dépasse guère 5 caractères de part et d'autre du caractère à transcrire, d'où le terme de règles locales. Presque toutes ces règles ont des exceptions, qui constituent une des difficultés de la phonétisation automatique du français, et il en est de même pour l'anglais.

Outre les 200 règles locales, ce phonétiseur comportait une liste d'une vingtaine de mots particulièrement difficiles à transcrire, et pour lesquels les règles locales étaient insuffisantes. Depuis, de nombreux autres phonétiseurs par règles ont été construits sur le même principe, et avec le même objectif : la synthèse de la parole à partir de textes écrits. Ils diffèrent les uns des autres par le nombre de mots qu'ils phonétisent correctement et par diverses améliorations qui ont été apportées. Nous ne les citons pas tous ici, mais nous rappelons dans quelles directions principales les auteurs de ces phonétiseurs par règles ont porté leurs efforts.

Dans trois réalisations différentes, trois équipes françaises ont tenté d'apporter des solutions à un important problème pratique : la mise au point et la mise à jour des phonétiseurs par règles. Au CEA, L. Poirot et X. Rodet (1976) ont élaboré un outil d'aide à la construction et à la modification des règles locales de leur phonétiseur. Cet outil effectuait des recherches de chaînes de caractères dans une liste de 65.000 mots orthographiés, ceci afin d'obtenir la liste des mots susceptibles d'être concernés par une règle donnée. L'extraction de listes de mots dans des dictionnaires électroniques est d'ailleurs devenue une méthode de base que nous utilisons régulièrement, aussi bien dans des buts théoriques que pour réaliser des applications informatiques. M. Divay et M. Guyomard (1977), au CNET, ainsi que B. Prouts (1979), au LIMSI, ont apporté des contributions dans le même sens en programmant leurs phonétiseurs de façon à ce que les règles soient faciles à exprimer, à lire et à modifier. Aux yeux de ces auteurs, la souplesse de mise au point et de mise à jour était en effet un élément capital de la qualité d'un phonétiseur par règles.

Ce souci de souplesse et de lisibilité s'explique aisément si l'on considère la hiérarchie complexe que forment les règles et les exceptions. Soit un phonétiseur par règles, qui comporte par exemple les règles suivantes :

Règle R_1 .

La lettre *s* se prononce [s].

Exemple : *piste* [pist].

Exception E_1 à la règle R_1 .

Un *s* entre voyelles se prononce [z].

Exemple : *rasoir* [razwar].

Exception E_2 à l'exception E_1 .

Après certains préfixes tels que *para-*, *s* entre voyelles se prononce [s] et non [z].

Exemple : *parasol* [parasɔl].

Exception E₃ à l'exception E₂.

Dans quelques mots comme *parasite* et ses dérivés, *s* se prononce [z] et non [s] : [parazit].

Dans un système de règles et d'exceptions, les éléments modifiables minimaux ne sont pas les éléments lexicaux, mais les règles et les exceptions. Une règle ou une exception donnée concerne généralement de nombreux éléments lexicaux. De plus, les règles et les exceptions constituent une hiérarchie complexe puisque, par définition, toute exception a un caractère restrictif vis-à-vis d'une ou plusieurs règles. Même si la programmation du système ne reflète pas explicitement cette structure hiérarchique, une règle et ses exceptions ne peuvent pas être modifiées indépendamment. Dans un tel système, lorsqu'on modifie, qu'on corrige, qu'on ajoute ou qu'on supprime une règle ou une exception, cela modifie le traitement de plusieurs mots : tout d'abord, ceux qui relèvent de la règle touchée ; et éventuellement, ceux qui relèvent de règles ou d'exceptions en relation hiérarchique avec cette règle. Or le système n'indique pas la liste des mots concernés par une règle donnée. C'est pourquoi, d'un point de vue pratique, il est souvent difficile d'évaluer les conséquences de la modification apportée, et celle-ci peut produire des effets secondaires inattendus. Par exemple, l'exception E₂ ci-dessus doit comporter une liste de préfixes après lesquels *s* se prononce [s] et non [z] :

<i>para-</i>	<i>parasol</i>
<i>aéro-</i>	<i>aérosol</i>
<i>ortho-</i>	<i>orthosympathique</i>
<i>a-</i>	<i>asymétrique</i>
<i>déca-</i>	<i>décasyllabe</i>
etc.	

L'établissement de cette liste est une opération complexe, ainsi que toute modification qu'on veut lui apporter. Par exemple, si on introduit *dé-* parmi ces préfixes, pour tenir compte de *désolidariser*, où *s* se prononce [s], on doit prendre en considération les exemples tels que *désembuer*, où *s* se prononce [z], et les ajouter à la liste mentionnée dans l'exception E₃, faute de quoi on glisse dans le système une erreur systématique. Ces difficultés, en dernière analyse, résultent du fait que la structuration en règles et en exceptions entraîne des dépendances dans le traitement des différents éléments lexicaux. Dans un dictionnaire phonétique électronique, au contraire, les données sont structurées d'une façon qui reflète le fait que la prononciation de *s* dans *désolidariser* et dans *désembuer* sont des informations indépendantes à considérer indépendamment.

2. L'apparition de la phonétisation par dictionnaire

La phonétisation par dictionnaire phonétique informatisé est apparue pour la première fois pour l'anglais, et ce toujours en vue de la synthèse de la parole. Dans le système de C.H. Coker, N. Umeda et C.P. Browman (1973), la production du texte phonétique est simple dans son principe : la consultation d'un dictionnaire phonétique informatisé permet de retrouver la transcription phonétique de chaque mot du texte.

Le système de phonétisation développé au MIT (J. Allen, 1973 ; J. Allen, M.Sh. Hunnicutt et D. Klatt, 1987) utilise un dictionnaire de parties de mots. En anglais,

si l'on connaît la décomposition d'un mot en ses éléments morphologiques, la transcription phonétique automatique en est facilitée. Par exemple, après avoir analysé l'adjectif anglais *changeable* en *change* + *able* à l'aide du dictionnaire d'éléments de mots et d'un ensemble de règles morphologiques, ce système phonétise correctement ce mot sans tomber dans les pièges tendus par la séquence orthographique *ea* en anglais. En allemand, la connaissance des limites entre les éléments morphologiques des mots est même nécessaire à la phonétisation automatique, car les mots composés sont soudés : les noms *Kinder* et *Sprache*, par exemple, forment un composé orthographié *Kindersprache*. Il en est de même pour quelques composés français. Ainsi, si on analyse l'adjectif *antiatomique* en *anti* + *atomique*, il devient facile à phonétiser correctement, alors que la séquence orthographique *-tia-* est ambiguë en français : elle se prononce de trois façons différentes dans *antiatomique* [tia], *étiage* [tja] et *initiation* [sja]. Le nom *interaction* est un exemple analogue à *antiatomique*, mais ce type de composés est marginal en français : dans la grande majorité des cas, la composition et la dérivation n'introduisent pas d'obstacles supplémentaires à la phonétisation, que les éléments soient soudés, comme dans *perçage* ou *électroacoustique*, ou séparés, ce qui est courant, comme dans *corps de chauffe*. C'est pourquoi, dans le cas du français, l'utilisation d'un dictionnaire d'éléments de mots ne réaliserait probablement pas une économie substantielle par rapport à un dictionnaire de mots.

La phonétisation par dictionnaire débute en France au CERFIA avec G. Tep (1979), qui utilise un dictionnaire phonétique et grammatical informatisé, également en vue de la synthèse de la parole à partir de textes écrits. Pour les raisons que nous venons d'exposer, le dictionnaire comporte des mots, et non des éléments de mots ; il s'apparente donc à celui de C.H. Coker et al. (1973). La consultation du dictionnaire permet de retrouver la transcription phonétique de chaque mot du texte. Le dictionnaire fournit également des informations grammaticales simples : les catégories grammaticales des mots - verbe, nom, adjectif... - et leurs traits flexionnels - temps, personne, genre, nombre. Ces données grammaticales sur les mots du texte sont essentielles dans le système de G. Tep, car elles compensent en partie les ambiguïtés de l'orthographe. Considérons deux exemples qui illustrent cette désambiguïssation : les liaisons et les homographes non homophones.

Les liaisons constituent un des obstacles à la phonétisation de textes en français, car elles ne sont pas indiquées dans l'orthographe. Des observations simples permettent de dresser la liste des facteurs qui contribuent à déterminer si on fait une liaison ou non. L'initiale du mot qui suit est bien sûr un de ces critères. C'est celui qui est déterminant dans la paire suivante¹ :

*C'est son ([], *[n]) jouet favori*

C'est son ([], [n]) amusement favori*

¹ Les prononciations sont données entre crochets [] dans un alphabet phonétique. L'astérisque * indique une prononciation ou une phrase inacceptables. Quand plusieurs prononciations sont envisagées dans un même contexte, elles sont séparées par des virgules et entourées d'une paire de parenthèses : ([...], [...]).

mais ce sont les relations syntaxiques entre les deux mots qui jouent dans la paire suivante :

C'est son ([], [n]) amusement favori

*Il a fait un son ([], *[n]) amusant*

Ici ce sont les traits flexionnels du premier mot :

*Luc a un travers ([], *[z]) insupportable*

Luc a des travers ([], [z]) insupportables

Il y avait là un corps ([], ?[z]) inerte*

Il y avait là des corps ([], [z]) inertes

Ces trois facteurs concernent des informations phonétiques, grammaticales et syntaxiques sur les mots : ce sont des données formalisables et donc manipulables dans un système informatique. On peut leur opposer un quatrième facteur, nettement plus difficile à formaliser : le niveau de langue. Suivant le niveau de langue, on fait plus ou moins de liaisons facultatives, comme dans la phrase suivante :

Ce projet met ([], [t]) en jeu un investissement ([], ?[t]) important

La prise en compte de plusieurs niveaux de langue différents dans un phonétiseur par règles n'est ni indispensable ni facile à concrétiser. En revanche, le système de G. Tep met à profit les informations grammaticales fournies par le dictionnaire. Ainsi, lorsqu'on lui soumet la phrase *Les oiseaux volent*, il accède à la catégorie grammaticale de *oiseaux* qui est un nom ; cette information contribue à décider qu'une liaison sera faite entre *les* et *oiseaux*. Au contraire, dans *Donne-les à Guy*, le mot qui suit *les* est une préposition, ce qui détermine le système à ne pas faire de liaison.

On appelle homographes non homophones des mots qui s'écrivent de la même façon mais se prononcent différemment, comme les deux mots *poster* dans les phrases suivantes :

Guy a acheté un poster

Guy va poster une lettre

Pour des raisons évidentes, les homographes non homophones sont difficiles à phonétiser. Ici encore, la connaissance des catégories grammaticales des mots permet au système d'en désambigüiser certains. Considérons la phrase *Les poules du couvent couvent*, dans laquelle les deux mots *couvent* sont des homographes non homophones. Le dictionnaire donne la correspondance entre les deux formes phonétiques et les deux catégories grammaticales :

<i>couvent</i>	[kuvã]	nom
<i>couvent</i>	[kuvə]	verbe

Il donne également les catégories des autres mots de la phrase, ce qui permet de faire

quatre hypothèses pour l'ensemble de la phrase :

1. [le]	[pul]	[dy]	[kuvā]	[kuvā]
article	nom	article	nom	nom
2. [le]	[pul]	[dy]	[kuvā]	[kuvə]
article	nom	article	nom	verbe
3. [le]	[pul]	[dy]	[kuvə]	[kuvā]
article	nom	article	verbe	nom
4. [le]	[pul]	[dy]	[kuvə]	[kuvə]
article	nom	article	verbe	verbe

L'examen des séquences de catégories grammaticales permet au système d'éliminer les hypothèses 1, 3 et 4 à partir des règles suivantes, qui sont généralement vérifiées :

- Une phrase correcte ne comprend pas deux noms de suite.
- Un article n'est pas suivi d'un verbe.

Il ne reste plus alors que l'hypothèse 2, qui est la bonne.

Le dictionnaire de G. Tep était rudimentaire : il comprenait seulement 2.000 mots. Cependant, sa contribution constitue, à nos yeux, une étape importante pour trois raisons.

Premièrement, le choix d'un dictionnaire comme structure de données permettait, pour la première fois, une évaluation du nombre de mots correctement phonétisés par le système. Dans le cas d'un système de phonétisation par dictionnaire, ce nombre de mots est le nombre d'entrées correctes du dictionnaire. Dans le cas d'un phonétiseur par règles et en l'absence de tout dictionnaire, cette évaluation est impossible. Or la fiabilité d'un phonétiseur dépend du nombre de mots qu'il traite correctement. L'introduction d'un dictionnaire constituait donc un pas vers des systèmes plus fiables que les précédents.

Deuxièmement, le remplacement d'un système de règles de phonétisation par un dictionnaire phonétique apportait une solution radicale aux problèmes de mise au point, de mise à jour et de maintenance des phonétiseurs précédents. En effet, nous avons vu que dans un dictionnaire phonétique, l'élément modifiable minimal est l'entrée lexicale : la modification, la correction, l'ajout ou la suppression d'un article n'oblige pas à remettre en question d'autres articles. En raison de cette propriété des dictionnaires, il est relativement facile d'en faire la maintenance et la mise à jour. Le prix à payer pour cet avantage était l'encombrement et le temps d'accès d'un dictionnaire. Mais, si ces désavantages pouvaient sembler importants à l'époque, l'évolution de la technologie les a rendus mineurs aujourd'hui.

Troisièmement, l'utilisation d'un dictionnaire qui comportait des informations grammaticales permettait d'utiliser ces informations facilement et d'une manière explicite, comme l'illustre le traitement des liaisons et des homographes non homophones. D'autres réalisations, à la même époque, font état de la nécessité de considérations grammaticales pour la phonétisation en vue de la synthèse de la parole.

Ainsi, le phonétiseur par règles de B. Prouts (1979) fait intervenir quelques données grammaticales sur les mots du texte pour contribuer à la détermination des liaisons et de la prosodie, et celui développé par l'équipe HESO (N. Catach, 1984) en utilise aussi pour déterminer des liaisons et pour désambiguïser des homographes non homophones.

Remarquons toutefois, à propos des phonétiseurs par règles, que l'exploitation de données grammaticales ne peut y être aussi aisée que dans un phonétiseur par dictionnaire tel que le précédent. En effet, à l'aide d'un dictionnaire, de nombreuses informations grammaticales sur les mots d'un texte peuvent être réunies, mais en l'absence d'un dictionnaire, le problème est différent. Les catégories grammaticales des mots : nom, verbe, adjectif... et leurs traits flexionnels : temps, personne, genre, nombre, ne sont pas indiqués explicitement dans les textes orthographiés. Ils ne se déduisent pas de la forme des mots par des règles simples, ni même par des règles fiables. Dans leurs phonétiseurs par règles, B. Prouts (1979) et N. Catach (1984) ont donc recours à diverses solutions approchées et à des astuces. Par exemple, ils utilisent des listes de mots-outils : articles, pronoms personnels, prépositions. B. Prouts utilise en outre une petite liste d'adjectifs susceptibles d'être placés avant le nom, comme *petit* ; pour désambiguïser les homographes non homophones tels que *couvent*, il utilise un critère approché : il décide si le mot est précédé d'un groupe nominal sujet au pluriel en testant s'il est précédé d'un mot terminé par *s*. Même en accomplissant des prouesses, les procédures qui recueillent des informations sur les mots par ces méthodes réunissent des données grammaticales fragmentaires, et les procédures qui exploitent ensuite ces données peuvent difficilement compenser ce caractère fragmentaire. De plus, les règles approchées qui permettent de rassembler des informations grammaticales sur les mots, ainsi que celles qui exploitent ces informations, constituent un nouvel ensemble de règles qui s'ajoutent aux règles locales de transcription phonétique. Ces diverses règles sont en relation : elles ne peuvent pas, en général, être modifiées indépendamment les unes des autres. Il en résulte de nouveaux problèmes de mise au point et de mise à jour de ces logiciels.

3. Les dictionnaires phonétiques récents

A la suite de G. Tep, plusieurs équipes françaises construisent des dictionnaires phonétiques en vue d'applications informatiques. Au CERFIA, G. Pérennou dirige l'élaboration de la base de données phonétiques et grammaticales BDLEX, qui comprenait environ 150.000 formes dans sa version 0 (G. Pérennou et M. de Calmès, 1986) et 350.000 dans sa version 1 (G. Pérennou et M. de Calmès, 1987). Au LADL, nous travaillons sur le dictionnaire électronique DELAP qui inclut également des données phonétiques et grammaticales. Il totalisait environ 300.000 formes dans sa version 3 (E. Laporte, 1986) et 500.000 dans sa version 4 (février 1988). Ces deux systèmes de données linguistiques ont plusieurs points communs.

- (1) Ils représentent la prononciation des mots par des transcriptions.
- (2) Ils offrent des informations sur les variations phonétiques que subissent les mots, en distinguant les variations libres et les variations conditionnées. Les variations libres sont par exemple celle du nom *ananas*, qu'on prononce couramment [anana] ou [ananas], ces deux variantes étant interchangeables. On

trouve un exemple de variation conditionnée dans la conjugaison du verbe *manier* : le *i* du radical est syllabique (*il manie*) ou non syllabique (*manier*), et la distribution entre ces variantes dépend du suffixe de conjugaison.

(3) Les deux systèmes tiennent compte du fait que certaines variations phonétiques sont systématiques, ou du moins très étendues dans le lexique. Pour représenter ces variations systématiques, ils utilisent non pas des transcriptions phonétiques traditionnelles, mais des transcriptions phonémiques (au sens de Z.S. Harris, 1951), ou phonologiques, qui sont des représentations plus abstraites. Celles-ci constituent des formes de base à partir desquelles des transductions permettent de produire les transcriptions phonétiques des variantes. Ces représentations abstraites et ces transductions s'appuient sur des travaux théoriques. Les transductions sont implantées sous forme d'algorithmes.

(4) Des informations grammaticales sur les mots s'ajoutent aux informations phonétiques : les catégories grammaticales - verbe, nom, adjectif... - et les informations nécessaires à la flexion, c'est-à-dire à la conjugaison des verbes et aux variations en genre et en nombre des noms et des adjectifs. Des algorithmes s'appliquent aux formes canoniques des mots, par exemple aux verbes à l'infinitif, et en calculent les formes fléchies, ce qui produit des dictionnaires de formes fléchies. Ce calcul peut avoir lieu aussi bien sur les formes orthographiques que sur les formes phonémiques.

(5) Chacun des deux dictionnaires est associé à un système d'aide à la correction orthographique. Toutefois, ces deux systèmes d'aide fonctionnent sur des principes différents. Nous y reviendrons dans le chapitre VI, section 1, p. 141.

Ces deux systèmes de données constituent donc des ensembles complexes qui renferment une importante quantité d'informations. On trouve de plus dans le système BDLEX un essai de représentation formelle des relations de dérivation, telles que celles entre *égal* et *inégal* ou entre *égal* et *également*. Une liste de suffixes et de préfixes a donc été établie ; dans cette liste, certains suffixes homophones et homographes ont été distingués et représentés séparément. Ainsi, le préfixe *in-* de *inégal* est considéré comme distinct d'un ou plusieurs autres suffixes *in-*.

Par ailleurs, deux systèmes phonologiques parallèles ont été construits : l'un pour représenter la prononciation standard, et l'autre pour la prononciation marseillaise. Cette démarche est indispensable à l'élaboration de bases de données phonétiques susceptibles de couvrir un échantillon de la prononciation française plus étendu que le seul français standard.

De plus, le système d'aide à la correction orthographique associé à la base de données comporte d'une part des procédures adaptées à la correction des fautes qui ne modifient pas la prononciation, comme *rythme* pour *rythme*, et d'autre part des procédures adaptées à la correction des fautes dues à l'utilisation d'un clavier, comme *clacul* pour *calcul* (G. Pérennou, F. Lahens, P. Daubèze et M. de Calmès, 1986). En effet, ces deux types d'erreurs, qui ont souvent des causes distinctes, ne doivent être négligés ni l'un ni l'autre dans l'élaboration d'aides efficaces à la correction orthographique.

Mais l'élément central de BDLEX comme du DELAP en est le système phonémique, ou composante phonologique, qui permet de rendre explicites la prononciation et les variations des mots à l'aide de représentations phonémiques.

Dans BDLEX, ce système (F. Dell et M. Plénat, 1985 ; Lambert et M. Rossi, 1986) représente une contribution notable à l'étude des variations phonétiques conditionnées, en particulier en ce qui concerne les variations des huit mots *six*, *huit*, *dix*, *plus*, *tous*, *boeuf*, *oeuf*, *os*, les assimilations consonantiques, et les chutes de liquides finales. Ce système permet ainsi de produire, à partir des formes phonémiques, la transcription phonétique [apsã] du mot *absent* et [prãndytã] pour *prendre du temps*. Certaines variations libres particulièrement répandues dans le lexique ont également été étudiées, et le système produit les transcriptions phonétiques [ijɛr] et [jɛr] de l'adverbe *hier*, ainsi que l'effacement facultatif du *e* dit muet dans *la fenêtre*.

Toutefois, on peut regretter que l'extension lexicale de certaines variations libres ne soit pas spécifiée. Prenons un exemple. La prononciation variable de *hier*, [ijɛr] ou [jɛr] se retrouve dans de nombreux autres mots, comme l'infinitif *lier* : [lije] ou [lje] ; néanmoins, ces mots s'opposent à d'autres comme le nom *pied* et l'infinitif *manier*, qui ne subissent pas cette variation libre. Celle-ci concerne donc un ensemble de mots assez étendu, et caractérisé par certaines limites hors desquelles elle n'a pas lieu. En ce sens, il ne s'agit pas d'une variation systématique, bien qu'elle concerne de nombreux mots. Le français offre d'ailleurs d'autres exemples de variations libres qui sont assez répandues sans être systématiques, comme l'alternance entre consonnes simples et consonnes géminées : ainsi, le nom *allusion* se prononce soit [allyzjɔ̃] soit [alyzjɔ̃], alors que le verbe *aller* se prononce toujours avec une consonne simple. Une variation non systématique n'est complètement spécifiée que lorsque l'ensemble des mots qui la subissent a été caractérisé et distingué des autres mots². Un dictionnaire phonétique informatisé a plusieurs raisons d'être, mais une des plus importantes, du point de vue théorique comme du point de vue pratique, est de rendre possible cette caractérisation de l'ensemble des mots concernés par un phénomène non systématique. En effet, de tels systèmes de données constituent un matériau de choix pour examiner méthodiquement les entrées lexicales, comme nous le verrons dans les chapitres II et V. C'est même le seul outil qui rende cette tâche réalisable. Il est donc regrettable que, malgré une analyse détaillée des variations libres des types *hier*, *tuer* et *louer*, et une implantation informatique des résultats obtenus, l'extension lexicale de ces variations soit négligée dans l'analyse comme dans l'implantation. Dans le formalisme choisi, qui est celui de la phonologie générative, l'élément qui permet de représenter ces variantes libres est le caractère facultatif de certaines règles. L'application facultative d'une règle produit deux variantes libres pour un mot, l'une à laquelle la règle s'est appliquée, et l'autre à laquelle elle ne s'est pas appliquée. Avec ce mécanisme, les mots représentés comme sujets à la variation sont ceux auxquels la règle peut s'appliquer. Les conditions d'application des règles facultatives permettraient donc de spécifier l'extension lexicale de ces variations, notamment par l'entremise des marques de limites de morphèmes : les

² Les variations systématiques, comme certaines assimilations consonantiques, sont moins embarrassantes, car elles sont complètement spécifiées par la donnée des différentes variantes et des conditions dans lesquelles elles alternent.

règles mentionnent de telles marques, présentes dans certains mots tels que *tuer* représenté /ty-er"/, et non dans d'autres tels que *suite* /syit /. Mais, d'une part, si l'on examine les transcriptions phonémiques données par BDLEX, on n'y trouve aucune marque de limite de morphème. Par ailleurs, la règle qui s'applique aux mots comportant une limite de morphème, comme *tuer*, est une règle obligatoire, et donc ne produit pas de variantes libres. Enfin, nous verrons dans le chapitre V que le marquage de limites de morphèmes ne suffit pas à spécifier cette variation libre, car la variation n'est qu'approximativement corrélée à la présence d'une limite de morphème dans le cas des timbres [u] et [y] ; quant au timbre [i], plusieurs centaines de mots, comme *manier* et *étudier*, ne subissent pas la variation en dépit de l'existence d'une limite de morphème.

Un autre type de variations appelle une seconde critique. Ici encore, il s'agit de variations qui ne peuvent être spécifiées qu'à la suite d'un examen systématique du lexique : les variations libres isolées. Par exemple, la finale *-ct* de l'adjectif au masculin *exact* est prononcée ou muette, les deux variantes étant interchangeables. Cette variation libre est limitée à quelques mots, c'est pourquoi on peut la considérer comme isolée, par opposition aux variations répandues ou systématiques. Beaucoup de mots français d'origine étrangère sont sujets à des variations libres isolées, comme *charter*, dont la prononciation [Sarter] est en concurrence avec au moins une autre, moins francisée : [tSarter]. Ces variations libres isolées, celles des mots d'origine française comme celles des mots d'origine étrangère, semblent peu représentées dans BDLEX, de même que les mots d'origine étrangère en général.

4. Un intérêt croissant pour les dictionnaires phonétiques

L'élaboration de dictionnaires phonétiques et grammaticaux au CERFIA et au LADL manifeste un intérêt et un investissement croissants en ce qui concerne les dictionnaires phonétiques. Cette tendance s'explique par plusieurs raisons.

Une première explication est en relation avec les phonétiseurs par règles. Nous avons vu que les premiers systèmes de phonétisation ont été conçus en vue de la synthèse de la parole à partir de textes écrits, en d'autres termes pour la machine à lire. C'est en particulier le cas de presque tous les phonétiseurs par règles qui ont été réalisés entre 1969 et 1980, et ils sont nombreux. Il n'en est plus de même aujourd'hui : en France comme en Grande-Bretagne et aux Etats-Unis, la synthèse de la parole est de moins en moins conçue comme une application de la phonétisation par règles, alors qu'elle est mieux appropriée à d'autres applications.

En effet, il apparaît qu'une synthèse de qualité correcte requiert des matériaux plus complets que ceux qui peuvent être réunis sur un texte à l'aide d'un phonétiseur par règles, et même par une analyse grammaticale rudimentaire des mots du texte. Pour le constater, reprenons l'exemple du phonétiseur de G. Tep (1979), qui exploite des informations grammaticales recueillies à l'aide d'un dictionnaire. Ce système est capable de déterminer correctement les liaisons dans les phrases suivantes, en tenant compte de

la catégorie grammaticale du mot qui suit *les* :

Les oiseaux volent

Donne-les à Guy

Ce traitement des liaisons est fondé sur un étiquetage des mots simples par leurs catégories grammaticales, et sur les règles suivantes :

- On fait généralement une liaison entre *les* et un nom.
- On n'en fait généralement pas entre *les* et une préposition.

Nous ne contestons pas la nécessité de considérations grammaticales pour phonétiser des textes d'une façon fiable. Toutefois, les méthodes qui ont été implantées font appel à des règles approximatives telles que ces deux règles. Ces méthodes auraient pu être systématisées si de telles règles donnaient une réponse dans tous les cas, et une réponse fiable. Mais, par exemple, doit-on faire la liaison entre le mot *un* et un adverbe ? Dans les phrases suivantes, *un* est suivi de l'adverbe *aussi* :

Guy a un aussi grand projet qu'Anne

Luc en a un aussi

Dans la première, la liaison après *un* est obligatoire ; dans la seconde, elle est facultative. Ces exemples montrent qu'un simple étiquetage des mots par leurs catégories grammaticales ne suffit pas à placer les liaisons d'une façon sûre : ce dernier objectif demande une prise en compte globale de la phrase et l'utilisation de données syntaxiques formalisées plus complètes. Il en va de même pour la désambiguïsation des homographes non homophones.

On sait par ailleurs que la détermination des valeurs des paramètres prosodiques en vue de la synthèse de la parole demande également un traitement syntaxique approfondi. Une séquence phonétique est un élément important pour caractériser un message parlé, mais elle ne suffit pas. En effet, outre les informations phonétiques que nous venons d'évoquer, un message parlé a une "mélodie", ou plus précisément une prosodie, qui est cruciale pour la compréhension, car elle permet notamment de trouver les limites des mots et des groupes syntaxiques. Elle n'est pas représentée dans les transcriptions phonétiques : elle constitue un ensemble d'informations supplémentaires, dites informations prosodiques. La prosodie d'un message parlé peut être définie comme l'ensemble des phénomènes accentuels, rythmiques et intonatifs, qui se manifestent par l'intermédiaire de trois paramètres mesurables : la durée, l'intensité et la fréquence fondamentale.

- Chaque syllabe a une certaine durée, en général différente de celle des syllabes voisines, ce qui donne un rythme au discours. Par exemple, la dernière syllabe de certains mots d'une phrase a souvent une durée supérieure à la moyenne. La parole est interrompue par des pauses, également caractérisées par leur durée.
- L'intensité sonore d'un message parlé varie rapidement au cours du temps entre des valeurs réduites et des valeurs plus accentuées.
- La voix parlée connaît également des variations au cours du temps entre le grave et l'aigu. Ces variations sont bien perceptibles à l'oreille. D'un point de vue

acoustique, il s'agit des variations de la fréquence fondamentale de la voix ; on peut les mesurer après une analyse spectrale de la parole.

Comme les caractéristiques prosodiques d'un message parlé jouent un rôle important dans la compréhension, elles doivent être prises en compte pour produire des messages synthétiques d'une qualité correcte. Si l'on synthétise un message oral dans lequel on donne aux paramètres prosodiques des valeurs constantes dans le temps, on obtient un discours monocorde, sans rythme et plus difficilement intelligible.

Les paramètres prosodiques d'un discours oral ne sont représentés ni dans l'orthographe, ni dans les transcriptions phonétiques, si l'on excepte la ponctuation et le découpage en mots, qui sont insuffisants pour spécifier les valeurs de ces paramètres. Leur détermination repose sur des informations syntaxiques formalisées qui ne se limitent pas à un simple étiquetage des mots. Or l'obtention de telles informations suppose une analyse syntaxique du texte, et donc la consultation d'un dictionnaire qui comporte des informations grammaticales et syntaxiques. Dès lors, la phonétisation par règles n'apporte plus d'économie par rapport à la phonétisation par dictionnaire, puisque celle-ci n'implique pas un accès supplémentaire au dictionnaire. C'est d'ailleurs cette solution qui est utilisée au MIT. Le système décrit par J. Allen, M.Sh. Hunnicutt et D. Klatt (1987) comprend un module d'analyse syntaxique et un module qui produit le texte phonétique. Ces deux modules appellent un même module d'analyse morphologique et d'identification des mots. Ce dernier recueille des informations phonétiques, grammaticales et syntaxiques sur les mots à l'aide d'un dictionnaire d'éléments de mots. Ce dictionnaire comporte, outre les éléments de mots, des mots entiers tels que *bargain* et *commonplace*, pour lesquels la décomposition automatique en éléments morphologiques donnerait des résultats erronés du point de vue phonétique ou grammatical. Les auteurs précisent que la prosodie des messages produits pourrait être encore améliorée par une analyse syntaxique plus poussée.

Les abréviations et les sigles utilisés dans les textes requièrent également l'utilisation de dictionnaires spécifiques. Nous ne parlerons pas ici de la phonétisation de textes écrits en typographie pauvre, c'est-à-dire sans accents, sans cédilles et sans trémas, mais il est clair que l'usage de cette typographie introduit de nombreuses ambiguïtés supplémentaires, qui ne pourraient être levées qu'à l'aide d'un dictionnaire.

Ainsi, l'intérêt de la phonétisation par règles pour la synthèse de la parole s'estompe ; en revanche, cette technique a maintenant deux autres applications majeures :

- (1) La constitution de dictionnaires. Un phonétiseur par règles permet de construire un dictionnaire phonétique à partir d'une liste de mots ou d'expressions. Ainsi, le phonétiseur par règles de B. Prouts (1979) a été modifié au LIMSI par F. Néel, M. Eskenazi et J.J. Mariani (1986) puis A. Dujour (1987) pour produire des variantes phonétiques, l'un des objectifs étant l'aide à la constitution de dictionnaires adaptés à des locuteurs particuliers. De même, nous avons réalisé au LADL un phonétiseur par règles qui sert exclusivement à la mise au point et à la mise à jour du dictionnaire DELAP.

(2) L'aide à la correction orthographique. A partir du phonétiseur par règles précédent, nous en avons réalisé un autre, destiné spécifiquement à cette application. Ces deux phonétiseurs mettent en oeuvre des règles locales analogues à celles des phonétiseurs par règles. Ils sont décrits dans le chapitre II, sections 4 et 5, et le système d'aide à la correction dans le chapitre VI, section 1, p. 141. Soulignons que chacun des deux phonétiseurs est adapté à une application informatique précise. Ils diffèrent essentiellement en ce que le premier opère sur des mots connus, catalogués dans un dictionnaire de particularités orthographiques décrit dans le chapitre III, alors que le second opère sur des mots inconnus et propose une ou plusieurs solutions par mot. Enfin, le premier sert à la maintenance du second.

En raison de cette évolution, une synthèse de la parole plus fiable dépend de la constitution de dictionnaires phonétiques, grammaticaux et syntaxiques, qui pourront prendre le relais de la phonétisation par règles pour assurer la production des textes phonétiques. Etant donné l'enjeu que représente ce type d'applications, on comprend l'importance que prennent les dictionnaires phonétiques.

5. L'insuffisance des données phonétiques disponibles

Cette importance trouve une seconde explication dans la constatation suivante : l'utilisation informatique des transcriptions et textes phonétiques est confrontée à un problème descriptif. On manque de descriptions formalisées précises et exhaustives sur la prononciation des mots français, même si l'on s'en tient à l'usage courant dans le français standard contemporain. Si l'on excepte les systèmes de données évoqués plus haut, les systèmes de transcriptions phonétiques disponibles sont ceux des dictionnaires de langue, français et bilingues, ainsi que ceux des dictionnaires phonétiques, qui ne vont pas plus loin dans leurs analyses. Le dictionnaire d'A. Martinet et H. Walter (1973) a eu le mérite de mettre en évidence l'ampleur de la variabilité phonétique en français, et de montrer que des données beaucoup plus complètes sur les variations phonétiques sont indispensables. Le cadre naturel pour la représentation systématique des variations phonétiques est le dictionnaire. Cette structure de données permet en effet de distinguer les variations isolées des variations systématiques, et, dans les cas intermédiaires entre ces deux extrêmes, de préciser l'extension lexicale des variations.

Citons toutefois une tentative pour prendre en compte les variations phonétiques en l'absence d'un dictionnaire : il s'agit des travaux de F. Néel, M. Eskenazi et J.J. Mariani (1986) et d'A. Dujour (1987), mentionnés ci-dessus. Leur phonétiseur par règles produit une ou plusieurs transcriptions phonétiques par mot, afin de prendre en compte trois types de variations phonétiques :

- (1) un grand nombre de variations libres, telles que celles du mot *immédiat* prononcé [imedja] ou [immedja] ;
- (2) certaines liaisons ;
- (3) certaines variations conditionnées autres que les liaisons : des assimilations consonantiques telles que celle qui produit la prononciation [apsã] de *absent*, et l'effacement de consonne liquide qui produit [katpa] pour *quatre pas*.

(4) En outre, pour la majorité des paires d'homographes non homophones, telles que *fil* dans *Guy a un fils* et *Guy enroule les fils sur des bobines*, le phonétiseur produit la liste des transcriptions qui correspondent aux mots homographes, ici [fis] et [fil].

A partir d'une liste de mots, ce phonétiseur permet d'établir une liste de transcriptions phonétiques qui correspondent aux différentes variantes phonétiques des mots. Il est donc capable de fournir des informations comparables au contenu d'un dictionnaire phonétique qui traite des variations phonétiques. Toutefois, on ne saurait le considérer comme un équivalent peu encombrant d'un dictionnaire phonétique. Un tel logiciel a en effet nécessairement des limitations qui constituent autant de différences avec un dictionnaire.

- Les informations grammaticales sur les mots ne sont disponibles que dans le cas des quelques mots qui subissent un traitement spécifique et sont listés dans le programme : des mots-outils et des mots exceptionnels comme *fil*.

- L'extension grammaticale des variations libres ne peut pas toujours être cernée précisément. Nous avons vu que certaines variations libres ne sont pas systématiques : ainsi, la consonne double orthographique du nom *allusion* se prononce couramment comme une consonne simple ou comme une consonne double, alors que celle du verbe *aller* est toujours prononcée simple. Certaines variations libres sont même isolées, c'est-à-dire qu'elles ne concernent qu'une petite liste de mots : ainsi, les mots *dompter* et *exempter* comportent un *p* orthographique qui est muet ou prononcé, alors que le *p* est toujours prononcé dans *symptôme* et toujours muet dans *compter*. Dans ces exemples, l'orthographe des mots n'indique pas clairement s'ils sont sujets à la variation considérée. Un phonétiseur par règles, en l'absence d'un dictionnaire, ne peut alors offrir que des solutions approchées ; une solution par excès, qui consiste à produire des variantes inusitées en plus des variantes existantes, est alors généralement préférable à une solution qui reviendrait à omettre des variantes existantes. C'est le parti qu'ont pris les auteurs en ce qui concerne les consonnes doubles (*allusion*), le *p* muet ou prononcé (*dompter*), le *h* aspiré (*honte*), le *e* dit muet (*petit*), la séquence orthographique *gn* prononcée [nj] ou [gn] (*ignifuger*), le *l* final muet ou prononcé (*persil*), le *t* final muet ou prononcé (*août*)...

La mise au point de ce phonétiseur par règles témoigne d'un intérêt pour la description des variations phonétiques. Cet intérêt est partagé par les équipes qui mettent au point des dictionnaires, et les a amenées à opter pour des dictionnaires non pas phonétiques à proprement parler, mais phonémiques, c'est-à-dire dans lesquels des représentations phonémiques constituent des formes de base à partir desquelles des algorithmes produisent les transcriptions phonétiques des diverses variantes des mots. Le choix de ces représentations abstraites et la conception de ces algorithmes s'appuient sur des notions théoriques qui ont été introduites par les phonologues. Elles mettent en jeu des travaux de recherche linguistique, parmi lesquels les études dialectales et historiques ne sont pas exclues.

Mentionnons enfin deux raisons qui s'ajoutent aux précédentes pour justifier notre intérêt pour les dictionnaires phonétiques : d'une part, le développement de l'aide à la correction orthographique par des méthodes phonétiques, et d'autre part, les possibilités d'évaluation que recèle un dictionnaire électronique. Le premier sujet est traité dans le chapitre VI, section 1, p. 141. L'autre a déjà été évoqué : dans un dictionnaire phonétique, la structuration en entrées lexicales vérifiables fait qu'on peut évaluer le nombre d'entrées correctes. Cette évaluation fournit des indications sur la fiabilité avec laquelle on peut exploiter le contenu d'un dictionnaire dans une application informatique. Dans un phonétiseur par règles qui opère sur des mots, la seule façon d'évaluer le nombre de mots correctement phonétisés par les règles est de le faire par référence à un dictionnaire phonétique lui-même évalué.

Toutes ces raisons nous ont incité à élaborer un dictionnaire phonétique qui réunit un grand nombre d'entrées, et qui offre un traitement des variations, en relation avec les faits grammaticaux pertinents, tels que la flexion. Les différentes catégories de variations phonétiques sont distinguées. Notre travail a débouché sur deux applications informatiques décrites dans le chapitre VI : la génération de messages oraux à partir de représentations sémantiques et l'aide à la correction orthographique.

Chapitre II. Cadre général

Ce chapitre présente le système de données linguistiques qui a servi de cadre général à notre travail, et expose les méthodes utilisées pour le développer et pour l'exploiter. Ce système de données regroupe des algorithmes et des dictionnaires morphologiques et phonémiques du français. Il est désigné sous le nom de DELA¹. Notre expérience concerne principalement les programmes de phonétisation et le dictionnaire phonémique qui fait l'objet du chapitre III, mais ces éléments sont en relation étroite avec les autres éléments du système DELA et une vision d'ensemble est souhaitable avant de les aborder plus spécifiquement.

1. Dictionnaires électroniques et dictionnaires usuels

Le système DELA comprend des algorithmes et des dictionnaires électroniques ; ces derniers sont utilisés par les algorithmes. Nous appelons dictionnaires électroniques les dictionnaires conçus pour être utilisés par des programmes. Nous les opposons aux dictionnaires usuels, destinés à l'usage des personnes. Les premiers se présentent généralement sur support informatique, et les seconds sous forme imprimée. Mais ils présentent des différences plus profondes : on peut imprimer sur papier un dictionnaire électronique, et on peut introduire en machine le contenu d'un dictionnaire usuel, il n'en reste pas moins qu'ils diffèrent essentiellement par leur contenu.

Dans un dictionnaire électronique, toutes les informations sont explicitées, puisque les informations implicites ne sont pas utilisables en machine. Dans un dictionnaire usuel, de nombreuses informations sont laissées implicites : celles qui n'intéressent que peu de lecteurs, par exemple le fait que le participe passé d'un verbe qui refuse tout objet direct, comme *flâner*, est invariable ; celles que le lecteur rétablira de lui-même, par exemple la prononciation du mot *car*, qui se lit sans difficulté d'après l'orthographe ; celles qui sont douteuses, par exemple le pluriel de nombreux noms composés. De nombreuses informations sont données sous forme de textes : définitions, exemples, citations, explications, remarques. On peut lire, par exemple, à l'article *beau* d'un dictionnaire : "*bel* devant un mot commençant par une voyelle ou un *h* muet", ou, à l'article *fleurir* : "au sens figuré, l'imparfait et le participe présent ont des formes en *floriss-* à côté des formes en *fleuriss-*". Pour des raisons évidentes, dans un dictionnaire électronique, ces informations doivent être codées d'une façon plus formalisée.

Par ailleurs, l'exhaustivité est plus importante dans le cas d'un dictionnaire électronique que dans le cas d'un dictionnaire imprimé : lorsqu'un dictionnaire électronique donne un certain type d'informations sur les mots, ces informations ne sont utilisables que si elles sont fournies pour tous les mots concernés. Ainsi, les dictionnaires du système DELA indiquent si les noms peuvent se mettre au féminin. Pour certains noms, la réponse est évidente : ainsi, il est clair que *acteur* peut se mettre au féminin, et que *ordinateur* ne le peut pas. Pour d'autres, comme *zingueur*, le doute est permis.

¹ Dictionnaires Electroniques du LADL.

Cependant, pour que cette information soit utilisable par des programmes, elle doit figurer pour tous les noms. Le problème est différent dans le cas des dictionnaires imprimés : rien n'impose qu'ils soient soumis à une règle aussi absolue. D'ailleurs, même les dictionnaires excellents ne se prononcent pas dans certains cas. En effet, peu de lecteurs, par exemple, consulteront leur dictionnaire pour savoir si le nom *zingueur* peut se mettre au féminin.

Enfin, si des informations grammaticales, syntaxiques ou phonétiques peuvent être formalisées d'une façon assez naturelle, il n'en est pas de même, par exemple, des définitions sémantiques des mots. Celles-ci n'ont donc pas leur place dans un dictionnaire électronique, alors qu'elle constituent un élément fondamental de nombreux dictionnaires usuels.

Pour résumer ces différences, le contenu d'un dictionnaire électronique est formalisé et directement exploitable par des programmes. Cette propriété a une conséquence importante : on peut effectuer des manipulations sur le contenu d'un dictionnaire électronique par des traitements automatiques, et par exemple construire automatiquement un second dictionnaire à partir du premier ; nous en verrons de nombreux exemples dans ce chapitre. Ce genre de traitement automatique à grande échelle n'est pas envisageable sur un dictionnaire usuel, dont le contenu ne peut guère être modifié que mot par mot.

La notion de dictionnaire électronique, telle que nous cherchons à la définir, peut être caractérisée d'une part par le contenu des dictionnaires, comme nous venons de le faire, et d'autre part par leur structure. De ce point de vue, un dictionnaire électronique peut être implanté en machine de multiples façons. Ainsi, les entrées peuvent être de longueur fixe ou de longueur variable. Toutefois, quelques propriétés de base semblent être des caractéristiques essentielles et générales de la structure d'un dictionnaire :

- Un dictionnaire électronique est un ensemble d'articles.
- Chaque article comporte un champ qui sert de clé d'entrée au dictionnaire. Si la présentation est proche de celle des dictionnaires usuels, cette "clé d'entrée" est constituée par la forme orthographique des mots. Lors de la recherche d'un mot dans le dictionnaire, le mot est comparé aux clés d'entrée, et si le dictionnaire comporte un article sur le mot, l'une des clés donne accès à cet article. Les articles sont généralement rangés par ordre alphabétique des clés.
- Chaque article comporte, outre la clé, des données sur le mot auquel est consacré l'article. Un cas limite est celui où le dictionnaire consiste en une simple liste de mots : les clés d'entrée sont alors les mots de la liste, et les données fournies par le dictionnaire se résument à la forme des mots.

Dans la formulation de ces trois propriétés, nous avons supposé qu'il s'agissait de dictionnaires de mots, mais les caractéristiques d'un dictionnaire de mots composés ou d'expressions figées sont analogues.

Nous avons vu une conséquence de ces propriétés dans le chapitre I, section 1, p. 10, en faisant une comparaison entre un système de règles et un dictionnaire conçus pour effectuer une même opération. Rappelons nos conclusions. Dans un dictionnaire

dont chaque article est consacré à un mot distinct, les articles sont indépendants les uns des autres : chacun peut être considéré et modifié indépendamment des autres. Dans un système de règles, une règle ou une exception donnée concerne généralement de nombreux éléments lexicaux ; de plus, les règles et les exceptions constituent une hiérarchie complexe ; cette structuration des données entraîne des dépendances dans le traitement des différents éléments lexicaux. C'est pourquoi un dictionnaire a une plus grande facilité d'évolution qu'un système de règles.

Un phonétiseur par règles et un dictionnaire phonétique sont équivalents en puissance : ils donnent des formes phonétiques censées correspondre aux formes orthographiques qu'on leur soumet ; tous deux peuvent comporter des erreurs. Les différences entre ces deux formalismes se situent à d'autres niveaux :

- La facilité d'évolution.
- La facilité d'évaluation (cf. chapitre I, section 2, p. 14).
- En revanche, un phonétiseur par règles occupe un encombrement plus réduit qu'un dictionnaire.

Dans la suite, lorsque nous ferons référence à des dictionnaires électroniques, nous parlerons plus simplement de dictionnaires.

2. Le dictionnaire DELAS²

Le système DELA est un ensemble de fichiers et de programmes destinés à assurer une couverture lexicale du français aussi complète que possible. Le dictionnaire DELAS en est l'élément central. C'est une base de données orthographiques, dont la spécificité est le codage grammatical et morphologique associé à chaque entrée du dictionnaire. Les grammairiens ont observé des variations de forme des mots et les ont cataloguées. C'est probablement dans ce but qu'ils ont défini les parties du discours, ou catégories grammaticales :

- (1) les verbes,
- (2) les noms,
- (3) les adjectifs et participes,
- (4) les adverbes, invariables,
- (5) les parties résiduelles, en général invariables : conjonctions, prépositions, etc.

Cette classification ne prend en compte que les mots simples, et non les mots composés et expressions figées tels que *à temps*. Dans l'immense majorité des cas, les mots simples qui composent ces formes complexes appartiennent à l'une de ces 5 catégories.

Les catégories (1) d'une part, (2) et (3) d'autre part, ont des schémas de variation différents. Ces observations et cette classification ont été utilisés pour construire les dictionnaires et lexiques usuels, et repris pour construire des dictionnaires électroniques. Ainsi, on constitue des classes d'équivalence de mots en les regroupant par familles de

² Dictionnaire Electronique du LADL pour les mots Simples.

variations :

- (1) en regroupant les formes simples conjuguées ou non d'un même verbe,
- (2) en regroupant le singulier et le pluriel d'un même nom,
- (3) en regroupant le masculin singulier, le masculin pluriel, le féminin singulier et le féminin pluriel d'un adjectif ou d'un participe.

Traditionnellement, on prend comme représentant des classes d'équivalence, c'est-à-dire comme forme canonique des entrées :

- (1) le verbe à l'infinitif,
- (2) le nom au singulier, et éventuellement au masculin singulier si le genre est variable,
- (3) l'adjectif ou le participe au masculin singulier,

et on considère que les autres membres de la classe s'obtiennent à partir du représentant par flexion (par exemple par conjugaison). Afin de décrire les variations de forme des mots, les classes morphologiques suivantes ont été recensées (B. Courtois, 1984) :

- 99 conjugaisons représentées chacune par un verbe modèle, avec pour chaque combinaison temps-personne-nombre, un couple d'un radical et d'une terminaison.
- 80 classes représentatives des variations des noms et adjectifs. Chacune de ces classes est associée à une suite de quatre terminaisons qui correspondent aux quatre formes possibles du mot.

Les classes sont numérotées. A chaque verbe, nom et adjectif du DELAS est associé le numéro de la ou des classes auxquelles il se rattache. Ce code indique ainsi toutes les variations possibles du mot. Cela permet d'engendrer automatiquement toutes les formes fléchies du français, ensemble qui, avec les mots invariables, constitue le dictionnaire DELAF³ (B. Courtois, 1984).

Le DELAS contient uniquement des mots simples, écrits en minuscules et sans aucun séparateur. Les mots composés figurent dans d'autres fichiers, que le séparateur soit

- un trait d'union : *sang-froid*,
- un espace : *vin blanc*,
- ou une apostrophe : *aujourd'hui*.

Sont également exclues du DELAS les expressions comportant d'autres séparateurs : *c.à.d.*, *id.*, *et/ou*, *par A + B...* ou des lettres majuscules, qu'il s'agisse de noms propres : *Pau*, *Danois*, ou de sigles : *PTT*, *COBOL*, *pH*. Sont inclus, parfois à titre provisoire, des mots tels que *électroacoustique*, qui devra faire l'objet d'une décomposition. Sont également inclus, par convention, des mots simples qui n'apparaissent que dans des mots composés. Par exemple, le mot *escampette* n'apparaît que dans le verbe composé *prendre la poudre d'escampette*. On peut ainsi trouver dans le DELAS les mots *escampette*, *lurette*, *caha*, *cahin*, *aujourd*, *hui*, *habeas*, ... Ces décisions peuvent se justifier dans le cadre des procédures qui reconnaîtront les mots composés d'un texte grâce à des consultations de dictionnaires électroniques (M. Silberztein, 1987). Dans son état actuel,

³ Dictionnaire Electronique du LADL pour les formes Fléchies.

l'ensemble comporte environ 64.000 mots. A titre de comparaison, le *Petit Larousse Illustré 1988* contient 52.500 articles, dont des mots composés, le *Petit Robert 1977* en compte environ 60.000, dont des mots composés, et le *Larousse Lexis 76.000*, dont des mots composés et des mots techniques.

Parmi les différents lexiques orthographiques constitués à partir du DELAS, citons :

- DELAS-L. Il s'agit de la liste des mots simples, sans codage morphologique. Cette liste de référence a l'intérêt de fournir une bonne couverture du français sur support électronique.

- DELAS-C. Il s'agit du DELAS subdivisé en cinq classes qui correspondent aux cinq parties du discours : verbes, noms, adjectifs, adverbes, et catégories résiduelles. Ce dictionnaire est donc comparable aux listes de E. Mater (1966).

- DELAS-I. Ce dictionnaire est constitué des mots du DELAS rangés en ordre alphabétique inverse, c'est-à-dire en partant de la droite des mots. Ce lexique a l'avantage de présenter les mots regroupés d'après leurs terminaisons.

- DELAS-IC. Ce sont les cinq catégories du DELAS-C avec les mots rangés en ordre alphabétique inverse.

- DELAF. Ce dictionnaire regroupe toutes les formes fléchies des mots du DELAS. Les verbes conjugués, ainsi que les noms et les adjectifs mis au féminin et au pluriel, s'ajoutent aux formes canoniques dont ils sont issus, et aux mots invariables. Le DELAF compte actuellement plus de 500.000 entrées. L'une des utilisations de ce dictionnaire est de fournir des informations sur les mots qui constituent un texte, ceci en vue de l'analyser. Outre la partie du discours, une information élémentaire est donc associée à chaque forme. Ainsi, pour un verbe, le DELAF donne le temps, la personne et le nombre : *serait* est donné comme le futur du verbe *être* à la troisième personne du singulier ; pour un nom ou pour un adjectif, le DELAF donne le genre et le nombre : *chevaux* est donné comme le pluriel du nom masculin *cheval*. Dans le cas de formes ambiguës, comme *cause* qui est soit un nom féminin soit une forme du verbe *causer*, le DELAF présente tour à tour les différentes hypothèses possibles et donne pour chacune d'elles les informations grammaticales requises.

- DELAF-L. Cette liste est identique à la précédente, mais les informations morphologiques en sont absentes. Il s'agit donc de la liste des mots simples fléchis du français, indispensable à des opérations telles que la vérification orthographique automatique (M. Silberstein, 1987).

- Tous ces dictionnaires ont été construits au fur et à mesure des besoins, et d'autres sont à l'étude, notamment en vue de recherches sur la dérivation en français.

Les mots composés, avec ou sans trait d'union, figurent dans des dictionnaires séparés. Les plus nombreux étant les noms composés, il sont répartis en plusieurs classes (Gaston Gross, 1986) dont l'union porte le nom de DELAC⁴ et qui rassemblent actuellement près de 100.000 articles. Les mots composés *vin blanc*, *corps de métier*, *date limite*, *tire-bouchon*, *tête-à-queue*, appartiennent ainsi à autant de classes. D'autres classes ont été définies pour les expressions figées de la forme N_0 être *Prép X*, comme *être en avance*, (L. Danlos, 1981), pour les adverbes figés, comme *à temps* (Maurice

⁴ Dictionnaire Electronique du LADL pour les noms Composés.

Gross, 1986), pour les adjectifs composés, comme *sujet à caution*, et pour les verbes composés, comme *se rendre compte* (Maurice Gross, 1982). Pour chaque classe, une table donne la liste des éléments et leurs propriétés syntaxiques.

3. Le dictionnaire DELAP⁵

En parallèle avec cette base de données orthographiques, le dictionnaire DELAP donne des informations phonémiques sur les mots. Il sera décrit en détail dans le chapitre III. Les dictionnaires DELAS et DELAP décrivent les mêmes mots : nous l'avons vu, il s'agit des mots simples, ne comportant aucun séparateur, écrits en minuscules sous leur forme canonique. Un système de comparaison permet d'introduire dans l'un les mots ajoutés à l'autre. Actuellement, le DELAP compte donc environ 64.000 mots, comme le DELAS. Dans chaque entrée figure une représentation de la prononciation du mot, exprimée dans un alphabet phonémique. Lorsqu'un mot peut se prononcer de plusieurs façons, comme *exact* dont la finale *-ct* ne se prononce pas toujours, le DELAP donne la liste des prononciations courantes (cf. chapitre III, section 2, p. 51). Chaque entrée est divisée en trois zones :

- une zone orthographique où figure le mot décrit, sous la même forme que dans le DELAS,
- une zone phonémique donnant la prononciation du mot,
- une zone grammaticale et flexionnelle donnant la catégorie grammaticale du mot et ses variations.

Voici un exemple d'entrée du DELAP avec ses trois zones :

discothèque, /diskotek/, .N21

De même que le DELAS décrit les variations orthographiques des mots lors de la conjugaison, de la mise au pluriel et de la mise au féminin, la troisième zone du DELAP offre la description correspondante du point de vue phonémique et non plus orthographique. Les verbes, les noms et les adjectifs ont donc été classés suivant les variations phonétiques qu'ils subissent avec la flexion. Cette classification est en relation complexe avec la classification correspondante établie pour l'orthographe et utilisée dans le DELAS. Les deux classifications et le système de codage utilisé pour les noter sont présentés dans le chapitre IV. De même que des programmes permettent d'engendrer les formes fléchies à partir des formes canoniques, et donc de construire le DELAF à partir du DELAS, les programmes correspondants ont été écrits pour effectuer la même opération sur les formes phonémiques. En utilisant simultanément ces deux algorithmes, on obtient chaque forme fléchie sous ses deux représentations, orthographique et phonémique :

Singulier	<i>cheval</i>	/ʃ"val/
Pluriel	<i>chevaux</i>	/ʃ"vo/

Cela permet de construire, à partir du DELAP, le dictionnaire de formes phonémiques fléchies, DELAP-F. Ce dictionnaire rassemble toutes les formes fléchies, comme le DELAF, mais il donne en plus leur prononciation.

⁵ Dictionnaire Electronique du LADL pour la Phonémique.

Plusieurs autres lexiques ont été élaborés à partir du DELAP :

- Le DELAP-C est constitué par le DELAP subdivisé en cinq classes correspondant aux cinq parties du discours : verbes, noms, adjectifs, adverbes et autres catégories.

- DELAP-FH. Ce dictionnaire donne la prononciation des formes fléchies, comme le DELAP-F, mais les entrées sont rangées par ordre alphabétique des formes phonémiques, et non plus des formes orthographiques. Cette disposition permet de donner la liste des formes qui admettent une prononciation donnée. Il s'agit donc d'un dictionnaire d'homonymes.

- Le DELAP-FL est la liste des formes fléchies, exprimées en transcription phonémique, mais sans indication de leur orthographe. Ce dictionnaire est un outil théorique pour étudier la structure phonémique des mots, les syllabes et les combinaisons de phonèmes.

- Le DELAP-GC est un outil de recherche plus spécialisé que le précédent, conçu pour étudier systématiquement les groupes de consonnes phonétiques, comme [ksk] dans *excuser*. Il a été constitué comme suit. Une analyse des formes phonémiques du DELAP a permis de marquer tous les groupes de consonnes de toutes les entrées. A chacun de ces groupes a été consacrée une entrée du DELAP-GC. Le dictionnaire obtenu a été trié suivant l'ordre alphabétique des groupes de consonnes. Un traitement ultérieur a permis de calculer, pour chaque groupe de consonnes phonétiques, le nombre d'entrées dans lesquelles ce groupe apparaît.

- Cette liste de dictionnaires n'est pas limitative : de nouveaux éléments, construits de la même façon, rendraient possible, notamment, l'étude systématique de la dérivation à partir de données concrètes.

4. Le phonétiseur de mots connus

Un algorithme écrit en langage C, et donné dans l'annexe 1, calcule la prononciation des 64.000 mots avec un taux d'erreur virtuellement nul, ce qui permet d'engendrer le dictionnaire DELAP. Nous désignons cet algorithme sous le nom de phonétiseur de mots connus. Son originalité réside en deux points :

- Il engendre automatiquement le DELAP à partir d'un dictionnaire qui comporte des données formalisées sur certaines irrégularités des mots. Il s'oppose ainsi aux phonétiseurs par règles traditionnels, dont les entrées sont purement orthographiques et dont la fiabilité est nécessairement limitée.

- Il est adapté à une application pratique précise : la mise au point et la mise à jour du dictionnaire DELAP et d'un second phonétiseur, évoqué dans la section suivante, p. 41.

Dans cette section, nous exposons le principe de l'algorithme.

Pour constituer les algorithmes de phonétisation et le DELAP, nous avons mené une étude approfondie de la correspondance entre l'orthographe et la prononciation du français. Chacun sait que l'orthographe d'un mot reflète approximativement la façon dont il se prononce, mais que cette correspondance est loin d'être parfaite. Nous avons évoqué dans le chapitre I, section 1, p. 9, les "règles locales" du type gé → [Ze] qui permettent de transcrire phonétiquement une lettre ou une séquence de quelques

lettres, en tenant éventuellement compte de son contexte droit ou gauche. Nous avons vu des exemples de ces règles locales. Cependant, la prononciation d'un mot ne se déduit pas à coup sûr de l'examen de son orthographe. Ainsi, le *s* final du nom *cas* ne se prononce pas, mais celui du nom *atlas* se prononce. Cette différence n'est pas explicitement indiquée dans l'orthographe. C'est ce qui rend difficile la phonétisation automatique, puisque cette tâche consiste à calculer la prononciation d'un mot avec pour unique point de départ la séquence de lettres qui correspond à l'orthographe du mot. Plus précisément, la phonétisation automatique de mots se heurte à des difficultés chaque fois que, dans un mot donné, une lettre ou un groupe de lettres est ambigu quant à la façon dont il peut se prononcer. C'est le cas du *s* final des mots *abus*, *amas*, *atlas*, *buis*, *cas*, *mas*, *sas*, *vasistas*...

Pour mieux cerner cette difficulté, nous avons recensé les cas où l'orthographe n'est pas suffisamment explicite pour que la prononciation s'en déduise sans ambiguïté. Nous avons effectué ce recensement sur les 64.000 mots du dictionnaire DELAS. Afin que les résultats de ce travail soient utilisables informatiquement, les mots concernés ont été explicitement marqués dans des fichiers contenant les 64.000 mots. Les marques qui ont été introduites dans ces fichiers signalent les mots dont la prononciation ne se déduit pas automatiquement de l'orthographe. Elles signalent également dans quelle partie du mot se situe l'ambiguïté, car chaque difficulté concerne une lettre ou un groupe de lettres dans un mot. Ainsi, les problèmes posés par le *s* final de *cas*, *atlas*, etc., concernent spécifiquement la fin de ces mots. Enfin, les marques introduites sont différentes suivant la nature de la difficulté qu'elles signalent. Par exemple, des marques distinctes ont été introduites dans le mot *blues*, dont le *s* final se prononce [z], et dans *atlas* et *vasistas*, dont le *s* final se prononce [s]. Les fichiers ainsi constitués ont reçu le nom de DELASP car ils sont issus du DELAS et ont servi à constituer le DELAP. Le DELASP est donc un dictionnaire qui contient des informations explicites sur la correspondance entre l'orthographe et la prononciation. Plus exactement, il contient des informations sur les irrégularités de cette correspondance : lorsque la prononciation d'un mot se déduit de son orthographe sans difficulté par les "règles locales" les plus générales, l'entrée du DELASP ne comporte aucune information particulière et reproduit exactement l'entrée du DELAS. C'est le cas de la plupart des mots : *car*, *deviner*, *compact*, *phénomène*... Les autres mots comportent une ou plusieurs irrégularités dans la correspondance entre l'orthographe et la prononciation. Nous désignons ces irrégularités sous le nom de "particularités de lecture". Lorsqu'un mot comporte une "particularité de lecture", celle-ci concerne une lettre ou un groupe de lettres du mot. Nous avons donné aux "particularités de lecture" une représentation informatique explicite qui reflète ce caractère local : ainsi, nous considérons comme une "particularité de lecture" le *s* final prononcé du nom *atlas*, et nous représentons cette information par une marque accolée au *s* final. Le nom *atlas* apparaît donc dans le DELASP sous la forme *atlas(20.)*, alors que le nom *cas*, qui se prononce suivant les règles les plus générales, apparaît sans aucune marque. De même, dans le nom *aiguille*, le fait que le *u* se prononce constitue une irrégularité que nous représentons par la marque (39.) placée immédiatement devant *u* : *aig(39.)uille* ; les mots *anguille* et *guérir*, dans lesquels *u* ne se prononce pas, sont considérés comme réguliers.

Ainsi, les mots apparaissent dans le dictionnaire DELASP sous une forme plus informative que dans les textes : à partir de cette représentation enrichie en informations, la forme phonétique du mot est calculable sans difficulté par un algorithme. Nous avons donc réalisé en langage C un algorithme qui calcule la prononciation des mots du DELASP. Cet algorithme tire parti à la fois de l'orthographe du mot et des "particularités de lecture". Il dispose donc de toutes les données explicites utiles à ce calcul, c'est pourquoi il phonétise les 64.000 mots avec un taux d'erreur virtuellement nul. Le résultat du calcul est une représentation phonémique du mot donné en entrée ; il s'agit de la représentation phonémique qui figure dans le DELAP et qui est décrite dans le chapitre III, section 7, p. 59. Le "phonétiseur de mots connus" doit son nom au fait qu'il est conçu pour phonétiser non pas les mots tels qu'on les trouve dans les textes, mais des mots exprimés dans la représentation enrichie du DELASP, donc, en particulier, des mots connus.

L'ensemble formé par le phonétiseur de mots connus et le DELASP constitue un modèle de la correspondance entre l'orthographe et la prononciation en français. C'est un outil de recherche qui comporte des informations explicites sur cette correspondance. Cet ensemble de données est donc comparable aux résultats des autres études sur ce sujet, notamment celles de N. Catach (1984) et de J. Véronis (1986). Toutefois, notre étude diffère par le nombre de mots auxquels elle s'applique. Elle couvre une partie représentative du vocabulaire du français courant. Son autre originalité est que les "particularités de lecture" ont été représentées explicitement dans les fichiers, sous une forme utilisable par des algorithmes, et que le phonétiseur de mots connus utilise cette formalisation. Un des résultats de ce travail est que nous avons pu établir une classification des irrégularités constatées dans la correspondance entre l'orthographe et la prononciation. En effet, toutes les "particularités de lecture" des mots traités forment un ensemble de quelques milliers d'éléments que nous avons rangés en 200 classes numérotées de 1 à 200. Par exemple, la classe n° 20 regroupe tous les *s* finaux prononcés, comme ceux de *atlas*, *vasistas* et *bus*, sauf s'il s'agit de deux *s*, comme dans *schuss*. La classe n° 39 regroupe tous les *u* qui se prononcent bien qu'ils soient précédés de *g* et suivis d'une voyelle sans tréma, comme dans *aiguille* et *linguiste*. Chaque classe est consacrée à un cas difficile de la correspondance entre l'orthographe et la prononciation. Beaucoup de ces cas difficiles sont limités à quelques mots d'origine française : le *l* muet d'une trentaine de mots comme *fusil* ou *marsault* ; le *m* d'une vingtaine de mots comme *nom* ou *parfum*. D'autres sont limités à quelques mots d'origine étrangère, comme le *i* d'une trentaine de mots comme *silentbloc* ou *iceberg*.

Le tableau des pp. 34 à 38 est un catalogue des 200 classes de "particularités de lecture". Chaque ligne est consacrée à une de ces classes. Chaque classe est repérée par un numéro de 1 à 200 suivi d'un point et éventuellement d'un ou plusieurs caractères. Ces codes figurent dans la colonne de gauche du tableau. Pour chaque classe, la deuxième colonne donne un exemple de mot. Ce mot est donné sous forme orthographique, mais avec la "particularité de lecture" représentée par son code. (En général, lorsque le mot présente plusieurs "particularités de lecture", seule est représentée celle dont il s'agit dans la ligne considérée.) Dans la troisième colonne, le mot est donné dans la représentation phonémique calculée par l'algorithme. Cette représentation phonémique caractérise la prononciation du mot moyennant un

**Classification des irrégularités
constatées dans la correspondance entre l'orthographe et la prononciation**

Code	Exemple	Représentation phonémique de l'exemple	Module	Fonction
1.o	minim(1.o)um	minimom	r1a	um_mpb ()
2.0	d(2.0)ump	d0mp	r1a	um_mpb ()
3.	gra(3.)s	gras*	r1a	code_s ()
4.em	(4.em)m	em	r1a	um_mpb ()
5.	te(5.)mpo	tempo	r1a	mp_mb ()
6.s	chrestoma(6.s)thie	krestomasi	r1a	t ()
7.s'	coccy(7.s')x	koksis	r1a	code_x ()
8.-	afflu(8.-)x	afly	r1a	code_x ()
9.gz	e(9.gz)xact	egzakt	r1a	code_x ()
10.s	dou(10.s)x	dus*	r1a	code_x ()
11.	(11.)xérès	xeres	r1a	code_x ()
12.z	jalou(12.z)x	Zaluz*	r1a	code_x ()
13.z	di(13.z)xième	diziem	r1a	code_x ()
14.	(14.)hache	haS	r1a	code_h ()
15.	aéro(15.)sol	aerosol	r1a	code_s ()
16.	dy(15.)s(16.)harmonie	disarmoni	r1a	code_s ()
17.-	confin(17.-)s	konfin*	r1a	code_s ()
18.z'	blue(18.z')s	bluz	r1a	code_s ()
19.S	fa(19.S)scisme	faSism	r1a	code_s ()
20.'	bus(20.')s	bys	r1a	s_code ()
21.	dés(21.)habiller	dezabije	r1a	s_code ()
22.	(22.)chaos	kao	r1b	ch ()
23.t	ma(23.t)cho	matSo	r1b	ch ()
24.-	almana(24.-)ch	almana ¹	r1b	ch ()
25.x	(25.x)chamsin	xamsin	r1b	ch ()
26.t	capri(26.t)ccio	kapritSio	r1b	code_c ()
27.-	escro(27.-)c	eskro	r1b	code_c ()
28.s	(28.s)caecal	sekal	r1b	code_c ()
29.	don(29.)c	donk	r1b	code_c ()
30.g	se(30.g)cond	s"gond*	r1b	code_c ()
31.	éq(31.)uateur	ekuat0r	r1b	q ()
32.	éq(32.)uilateral	ekyilateral	r1b	q ()
33.u	q(33.u)uetsche	kuetS	r1b	q ()
34.	(34.)gnou	gnu	r1b	code_g ()
35.d	(35.d)gin	dZin	r1b	code_g ()
36.Z	appo(36.Z)ggiature	apoZiatyr	r1b	code_g ()
37.-	étan(37.-)g	etan*	r1b	code_g ()
38.	jag(38.)uar	Zaguar	r1b	g_code ()

¹ L'autre prononciation de ce mot, /almanak/, est également représentée dans le DELAP.

Code	Exemple	Représentation phonémique de l'exemple	Module	Fonction
39.	ling(39.)uiste	lingyist	r1b	g_code ()
40.y	ling(40.y)ual	lingyal	r1b	g_code ()
41.	grog(41.)gy	grogi	r1b	g_code ()
42.	gage(42.)ure	gaZyr	r1b	g ()
43.d	(43.d)jeep	dZip	r1b	j ()
44.	f(44.)jord	fjord	r1b	j ()
45.x	(45.x)jota	xota	r1b	j ()
46.	persu(46.)ader	persyade	r4b	u ()
47.'	sonn(47.')erie	son'ri	r23	code_e ()
48."	r(48.")essort	r'sor	r23	code_e ()
49.i	b(49.i)ehaviorisme	bi aviorism ²	r23	code_e ()
50.a	ard(50.a)emment	ardaman*	r23	code_e ()
51.0	freez(51.0)er	friz0r	r23	code_e ()
52.-	ri(52.-)esling	risliG	r23	code_e ()
53.an-	(53.an-)enivrer	an-ivre	r23	code_e ()
54.e	(54.e)oesophage	ezofaZ	r23	code_oe ()
55.0	l(55.0)oess	l0s	r23	code_oe ()
56.ua	m(56.ua)oelle	mual	r23	code_oe ()
57.ue	b(57.ue)oette	buet	r23	code_oe ()
58.0:	f(58.0:)oehn	f0:n	r23	code_oe ()
59.	ques(59.)tion	kestion*	r1a	t ()
60.a	(60.a)ennui	annyi	r23	nn ()
61.an	(61.an)emmerder	anmerde	r23	mm ()
62.en	(62.en)immettable	enmetabl	r23	mm ()
63.te	(63.te)t	te	r1a	t ()
64.o:	h(64.o:)all	ho:l	r23	ll ()
65.	vi(65.)lle	vil	r23	ll ()
66.li	cordi(66.li)llère	kordilier	r23	ll ()
67.j	canco(67.j)illotte	kankojot	r23	ill ()
68.d	pi(68.d)zza	pidza	r23	zz ()
69.u	z(69.u)oom	zum	r23	oo ()
70.-	alc(70.-)ool	alkol	r23	oo ()
71.j	a(71.j)ie	aj	r4a	i ()
72.	séquo(72.)ia	sekoia	r4a	i ()
73.-	o(73.-)ignon	oNon*	r4a	i ()
74."	f(74.")aisan	f'zan*	r4a	ai ()
75.ej	s(75.ej)aietter	sejete	r4a	ai ()
76.aj	l(76.aj)eitmotiv	lajtmotiv ³	r4a	ei ()
77.ii	pa(77.ii)ys	pei	r4a	code_y ()
78.j	ka(78.j)yak	kajak	r4a	code_y ()

² Ce mot se prononce également de plusieurs autres façons.

³ Ce mot se prononce également de plusieurs autres façons.

Code	Exemple	Représentation phonémique de l'exemple	Module	Fonction
79.i	fu(79.i)yant	fyiant*	r4a	code_y ()
80.f	alco(80.f)yle	alkoil	r4a	code_y ()
81.aj	b(81.aj)ye	baj	r4a	code_y ()
82.:	meu(82.:)te	m0:t	r4b	u ()
83.aw	(83.aw)out	awt ⁴	r4b	u ()
84.	p(84.)utsch	putS	r4b	code_u ()
85.w	ca(85.w)udillo	kawdijo	r4b	code_u ()
86.o	acup(86.o)uncture	akypontkyr	r4b	code_u ()
87.e	h(87.e)umble	enbl	r4b	code_u ()
88.0	homesp(88.0)un	homesp0n	r4b	code_u ()
89.0	bl(89.0)uff	bl0f	r4b	code_u ()
90.ju	barbec(90.ju)ue	barb"lju	r4b	code_u ()
91.-	b(91.-)uilding	bildiG	r4b	code_u ()
92.-	dreadno(92.-)ught	drednot	r4b	code_u ()
93.G	ri(93.G)ng	riG	r4b	en ()
94.-	vive(94.-)nt	viv	r4b	en ()
95.hen'-	(95.hen'-)nième	hen-iem	r4b	en ()
96.	coré(96.)en	koreen*	r4b	en ()
97.se	(97.se)c	se	r1b	code_c ()
98.Ze	(98.Ze)g	Ze	r1b	code_g ()
99.aS	(99.aS)h	aS	r1a	code_h ()
100.	bi(100.)en(101.--)heureux	bien-0r0z*	r4b	en_code ()
101.--	bi(100.)en(101.--)heureux	bien-0r0z*	r4b	en_code ()
102.'-	moy(103.)en(102.'-)âgeux	muajen aZ0z*	r4b	en_code ()
103.	moy(103.)en(102.'-)âgeux	muajen aZ0z*	r4b	en_code ()
104.Zi	(104.Zi)j	Zi	r1b	j ()
105.ku	(105.ku)q	ky	r1b	q ()
106.'	pollen(106.')	polen	r4b	finale ()
107.-	anspec(107.-)t	anspek	r4b	finale ()
108.t'	vel(108.t')d	velt	r4b	finale ()
109.	che(109.)z	Sez*	r4b	z final ()
110.ts'	kronprin(110.ts')z	kron'prints	r4b	z final ()
111.s'	ruol(111.s')z	ryols	r4b	z final ()
112.	c(112.)edex	sedeks	trm	cinquieme ()
113.:	aro(113.:)me	aro:m	trm	cinquieme ()
114.h	(114.h)onze	honz	trm	cinquieme ()
115.'	lamen(115.')to	lamen'to	trm	cinquieme ()
116.+	emploi interne au programme		r4b	u ()
			trm	cinquieme ()
117.es'	(117.es')s	es	r1a	s ()

⁴ Il s'agit ici de la prononciation [awt] en une syllabe. Le DELAP donne deux autres prononciations courantes de ce mot, plus francisées : [aut], en deux syllabes, et [ut].

Code	Exemple	Représentation phonémique de l'exemple	Module	Fonction
118.iks'	(118.iks')x	iks	r1a	code_x ()
119.0	(119.0)e	0	r23	code_e ()
120.	(120.)i	i	r4a	i ()
121.igrek	(121.igrek)y	igrek	r4a	code_y ()
122.en'	(122.en')n	en	r4b	en ()
123.y	(123.y)u	y	r4b	code_u ()
124.	(124.)a	a	r4a	code_a ()
125.be	(125.be)b	be	trm	cinquieme ()
126.de	(126.de)d	de	trm	cinquieme ()
127.ef	(127.ef)f	ef	trm	cinquieme ()
128.ka	(128.ka)k	ka	trm	cinquieme ()
129.el	(129.el)l	el	r1b	l ()
130.	(130.)o	o	trm	cinquieme ()
131.pe	(131.pe)p	pe	r4a	p ()
132.er'	(132.er')r	er	trm	cinquieme ()
133.ve	(133.ve)v	ve	trm	cinquieme ()
134.dblv	(134.dblv)w	dubl ve	r4a	w ()
135.zed'	(135.zed')z	zed	r4a	z ()
136.-	cu(136.-)l	ky	r1b	l ()
137.0l	pick(137.0l)les	pik0ls	r1b	l ()
138.-	(138.-)août	ut	r4a	code_a ()
139.e	c(139.e)ake	kek	r4a	code_a ()
140.ej	(140.ej)ale	ejl	r4a	code_a ()
141.:	co(141.:)at	ko:t	r4a	code_a ()
142.w	cia(142.w)o	tSaw	r4a	a ()
143.o:	w(143.o:)alkman	wo:kman	r4a	code_a ()
144.o	w(144.o)alkman	wokman	r4a	code_a ()
145.z	al(145.z)sacien	alzasien*	r1a	code_s ()
146.-	des(146.-)quels	dez* kel	r1a	s_code ()
147.v	aq(147.v)uavit	akvavit	r1b	q ()
148.0	angstr(148.0)öm	angstr0m	r4a	o_trema ()
149./	anti(149./)aerien	anti aerien*	r4a	i ()
150.-	plom(150.-)b	plon*	r1a	mp_mb ()
151.*	léger(151.*)	leZer*	r4b	finale ()
152.-	as(152.-)thme	asm	r1a	t ()
153.-	auto(153.-)mne	o:ton	r1a	um_mpb ()
154.-	ba(154.-)ptême	batem	r4a	p ()
155.v	(155.v)wagon	vagon*	r4a	w ()
156.u	cl(156.u)own	klun	r4a	w ()
157.-	slo(157.-)w	slo	r4a	w ()
158.i	j(158.i)ean	dZin	r23	code_e ()
159.-	cin(159.-)q	sen*	r1b	q ()

Code	Exemple	Représentation phonémique de l'exemple	Module	Fonction
160.aj	(160.aj)iceberg	ajsberg	r4a	i ()
161.a	br(161.a)ownien	brawnien*	r4a	w ()
162.iz	b(162.iz)usiness	biznes	r1a	s ()
163.d	scher(163.d)zo	skerdzo	r4a	z ()
164.ts	can(164.ts)zone	cantsone	r4a	z ()
165.s	kon(165.s)zern	konsern	r4a	z ()
166.-	cle(166.-)f	kle	trm	cinquieme ()
167.el	cockt(167.el)ail	koktel	r4a	il ()
168.n	co(168.n)mte	kont	r1a	um_mpb ()
169.ts	(169.ts)czimbalum	tsenbalom	r1b	code_c ()
170.-	débi(170.-)tmètre	debimetr	r1a	t ()
171.0	fash(171.0)ion	faS0n	r4a	i ()
172.o	footb(172.o)all	futbol	r23	ll ()
173.-	pie(173.-)dmont	piemont*	trm	cinquieme ()
174.j	genti(174.j)lhomme	Zantijom	r1b	l ()
175.i	(175.i)hyène	ien	r1a	code_h ()
176.er'/	hyp(176.er')/eracidité	iper asidite	r23	code_e ()
177.ju	n(177.ju)ewton	njuton	r23	code_e ()
178.x	(178.x)khamsin	xamsin	r1a	h ()
179.o	kh(179.o)ôl	kol	r4a	o_circonfl. ()
180.-	(180.-)knickers	nik0rs	trm	cinquieme ()
181.oj	kr(181.oj)eutzer	krojts0r	r23	code_e ()
182.0	patchw(182.0)ork	patSu0rk	trm	cinquieme ()
183.w	mia(183.w)ou	mjaw ⁵	r4b	u ()
184.0	m(184.0)onsieu(185.-)r	m0si0	trm	cinquieme ()
185.-	m(184.0)onsieu(185.-)r	m0si0	trm	cinquieme ()
186.i	piran(186.i)ha	pirania	r1a	code_h ()
187.-	puse(187.-)yisme	pyzeism	r4a	code_y ()
188.u	tennisw(188.u)oman	tenisuuman	trm	cinquieme ()
189.an	sert(189.an)ão	sertan*	r4a	a_circonflexe
190.juw	st(190.juw)eward	stjuward	r23	code_e ()
191.juv	intervi(191.juv)ewer	intervjuve	r23	code_e ()
192.u	st(192.u)ôpa	stupa	r4b	code_u ()
193.s	(193.s)thriller	srl0r	r1a	t ()
194.	-	-	-	-
195.-	fa(195.-)on	fan*	r4a	a ()
196.e	(196.e)innavigable	ennavigabl	r23	nn ()
197.+	li(197.+)er	li+e	r4a	i ()
198.	pro(198.)eutectique	pro0tektik	r4b	u ()
199.e	tj(199.e)äle	tjel	r4a	a_trema ()

⁵ Il s'agit ici de la prononciation [mjaw] en une syllabe. Le DELAP donne une autre prononciation courante de ce mot : [mjaʊ], en deux syllabes.

ensemble de conventions qui sont exposées au chapitre III, section 7, p. 59. Lorsqu'un exemple peut se prononcer de plusieurs façons, nous ne mentionnons qu'une variante dans le tableau. Ainsi, la ligne 9 ne mentionne que la transcription /egzakt/ de *exact*, bien que la transcription /egza/ figure également dans le DELAP. Il en est de même pour plusieurs autres lignes du tableau. Chaque classe de "particularités de lecture" est traitée dans une des fonctions du phonétiseur de mots connus. Le tableau donne le nom de cette fonction, ainsi que le module dont elle fait partie. Six des modules du phonétiseur comportent des fonctions utilisant les "particularités de lecture": r1a.c, r1b.c, r23.c, r4a.c, r4b.c et trm.c.

L'index alphabétique de la p. 40 permet de retrouver par leurs numéros les "particularités de lecture" qui concernent une lettre ou une séquence orthographique données.

D'autres formalisations et classifications que celle-ci seraient également valables, mais le phonétiseur ne peut être modifié indépendamment du DELASP: ces deux éléments reposent sur un même formalisme de représentation des "particularités de lecture". Ils forment un seul système et ne sont utilisables que concurremment.

L'algorithme de phonétisation de mots connus est structuré en 5 passages. Le premier passage opère sur le mot à phonétiser. Il produit une forme intermédiaire dans laquelle seule une partie des lettres ont été transcrites en phonèmes. Les autres passages opèrent sur des formes mixtes du même type. Le dernier passage produit la forme phonémique. La phonétisation se fait donc par l'intermédiaire de formes mixtes, comme dans le phonétiseur par règles développé à l'ENSERG par P. Goyer, D. Degryse et B. Guérin (1979).

Lors de chacun des 5 passages, la forme à transcrire est examinée caractère par caractère, et des fonctions appropriées sont appelées. Chacune de ces fonctions correspond à l'un des 5 passages et à un caractère ou une séquence de caractères particulière: par exemple, lors du 3^e passage, qui est consacré aux lettres doubles, la fonction *mm* est appelée à chaque occurrence de la séquence *mm*. Ces fonctions sont au nombre de 55 environ et effectuent la transcription en tenant compte du contexte droit et gauche et des "particularités de lecture" représentées. Lorsqu'une "particularité de lecture" a été prise en compte par une de ces fonctions et ne servira plus par la suite, elle est effacée. Les dernières "particularités de lecture" sont effacées lors du dernier passage, et la forme phonémique produite par le phonétiseur n'en comporte plus. En revanche, la forme orthographique d'origine et les 4 formes intermédiaires peuvent en comporter: ces formes sont donc représentées non pas par de simples chaînes de caractères, mais par une structure de données appropriée, qui associe à certains caractères des identificateurs de "particularités de lecture". Considérons par exemple le nom *aérosol*, dont le *s* présente une "particularité de lecture" rangée dans la classe n° 15. La structure de données qui permet de le représenter comporte:

- une chaîne de caractères dans laquelle figure *aérosol*;
- une série de pointeurs qui correspondent aux différentes lettres du mot.

a	124, 138, 139	le	137
ä	199	lh	174
ai	74, 75	m	4, 5, 153, 168
aïl	167	n	106, 122
al	137, 143, 144	nc	29
all	64, 172	ng	93
ao	142	nh	186
ão	189	o	113, 130, 182
aon	195	ô	179
au	85	ö	148
ay	77, 78	oa	141
b	125, 150	oe	54, 55, 56, 57
c	26, 27, 29, 30, 97	oé, oê, oè, oë	56, 57
cae	28	oeu	198
ch	22, 23, 24, 25	oill	67
cz	169	on	184
d	108, 126, 173	oo	69, 70
e	47, 48, 49, 52, 112, 119	ou	83, 183
ea	158	ought	92
ei	76	ow	156, 157, 161
emm	50, 61	oy	80
en	53, 96, 100, 101, 102, 103, 115	p	131, 154
enn	60	q	105, 159
ent	94	qu	31, 332, 33, 147
er	51, 151, 176	r	106, 132, 185
es	146	s	3, 15, 17, 18, 19, 20, 117, 145
eu	82, 181	sh	16, 21
ew	177, 190, 191	t	63, 106, 107, 170
f	127, 166	th	152, 193
g	34, 35, 36, 37, 98	thi	6
geu	42	ti	59
gg	36, 41	u	46, 84, 89, 90, 95, 116, 123
gn	34	ù, û	192
gu	38, 39, 40	ui	91
h	14, 99	um	1, 2, 87
hy	175	un	86, 87, 88
i	71, 72, 73, 120, 149, 160, 197	usi	162
ill	65, 66	v	133
imm	62	w	133, 134, 155
inn	196	wo	188
ion	171	x	7, 8, 9, 10, 11, 12, 13, 118
j	43, 44, 45, 104	y	78, 79, 81, 121, 187
k	128	z	106, 109, 110, 111, 135, 163, 164, 165
kh	178	zz	68
kn	180		
l	129, 136		

Tous pointent vers une valeur vide, sauf celui qui correspond à la lettre *s* du mot : celui-là pointe vers l'identificateur de la classe n° 15.

5. Le phonétiseur de mots inconnus

Le phonétiseur de mots connus et le DELASP sont adaptés à un rôle précis. D'une part, ils ont permis la construction du dictionnaire DELAP, et ils servent à sa maintenance et à sa mise à jour. D'autre part, ils ont été indispensables à la réalisation d'un second phonétiseur, que nous désignons sous le nom de "phonétiseur de mots inconnus", car il opère sur des mots nus, sans indication de leurs "particularités de lecture". Ce deuxième phonétiseur est caractérisé par les points suivants :

- (1) Il opère sur des formes purement orthographiques, telles qu'on les trouve dans les textes, c'est-à-dire sans indication de leurs "particularités de lecture", d'où son nom.
- (2) Lorsque l'orthographe d'un mot n'est pas tout à fait explicite sur la façon dont il se prononce, plusieurs transcriptions distinctes sont calculées pour prendre en compte les différentes hypothèses envisageables. Ainsi, *tas* est phonétisé [ta] et [tas] ; de même, *atlas* est phonétisé [atla] et [atlas].
- (3) Il est adapté à une application précise : l'aide à la correction orthographique par phonétisation. Cette application est présentée dans le chapitre VI, section 1, p. 141.

Pour calculer la transcription phonétique d'un mot à partir de son orthographe, on a besoin d'en connaître les "particularités de lecture" éventuelles. Le phonétiseur de mots inconnus doit donc les introduire dans les formes qu'on lui soumet. Considérons par exemple le cas du verbe *compter* : rien dans cette forme n'indique que le *p* est muet ; d'ailleurs, dans *dompter*, le *p* peut se prononcer. Pour obtenir la prononciation sans *p* de *compter*, le phonétiseur doit mettre en place l'hypothèse que ce mot présente cette "particularité de lecture" à cet endroit. Ainsi, à chaque lettre du mot, l'examen de quelques caractères avant et après permet de décider quelles "particularités de lecture" sont possibles. Lorsque l'orthographe du mot n'est pas complètement explicite, plusieurs solutions sont envisagées. La phonétisation peut ensuite se faire en tenant compte des "particularités de lecture" introduites lors de l'étape précédente. Le phonétiseur de mots inconnus est donc modularisé en deux composants :

- Le composant de séparation des hypothèses introduit des "particularités de lecture" à certaines positions du mot, et fournit éventuellement plusieurs hypothèses qui diffèrent par les "particularités de lecture" attribuées au mot.
- Le deuxième composant exploite ces informations et donne la forme phonétique qui correspond à chacune des solutions retenues. Ce deuxième composant n'est autre que le phonétiseur de mots connus évoqué dans la section précédente.

Dans le dictionnaire DELASP, les mots sont classés par ordre alphabétique. Nous avons constitué à partir du DELASP une liste des "particularités de lecture". Cette

liste, appelée DELAP-PL⁶, est structurée en 200 sous-listes qui correspondent aux 200 classes. Chaque sous-liste contient les mots qui comportent une "particularité de lecture" de la classe correspondante. Ainsi, dans la sous-liste n° 20 du DELAP-PL figurent par ordre alphabétique les mots français terminés par un *s* final prononcé, à l'exception de ceux qui sont terminés par deux *s*. En effet, cette définition correspond à la classe n° 20. Chacune des sous-listes regroupe donc les exceptions à une règle générale de phonétisation : un *s* final simple, en règle générale, ne se prononce pas, et la sous-liste n° 20 réunit les exceptions à cette règle. Les mots sont donnés sous deux formes : la forme orthographique, avec les "particularités de lecture" représentées par leur code, est suivie par la forme phonémique calculée par l'algorithme. Nous donnons dans l'annexe 2 la sous-liste n° 20 du DELAP-PL.

Le DELAP-PL est donc un catalogue d'exceptions classées. Les 64.000 mots du DELAS ont été pris en compte pour sa construction. Il en résulte que le DELAP-PL est aussi complet que possible, et ce de deux points de vue :

- il décrit tous les types d'irrégularités constatés dans la correspondance entre l'orthographe et la prononciation, y compris dans les mots d'origine étrangère ;
- pour chacun de ces types d'irrégularités, il donne une liste d'exemples aussi complète que possible.

Ces listes d'exceptions sont importantes pour la réalisation d'algorithmes de phonétisation fiables et robustes.

Avec le DELAP-PL s'achève cette présentation du système de dictionnaires DELA. Nous n'avons guère évoqué les travaux du LADL sur le lexique-grammaire du français. Il s'agit d'une description syntaxique formelle du français sous forme de tables de propriétés et en relation avec le système DELA. Plusieurs applications informatiques utilisent toutes ces données linguistiques pour effectuer diverses opérations telles que l'analyse syntaxique (M. Salkoff, 1973), la génération de textes (L. Danlos, 1986), la vérification orthographique (M. Silberstein, 1987), la correction orthographique (E. Laporte ; Le Kien V., 1987) et la production de textes phonétiques (L. Danlos, F. Emerard, E. Laporte, 1986).

6. L'exploitation d'un système de données linguistiques à des fins de recherche

Pourquoi élaborer des dictionnaires électroniques aussi étendus ? Un des intérêts de cette accumulation de données linguistiques est qu'elle fournit un outil de travail indispensable pour réaliser des programmes robustes. Pour concevoir et construire des algorithmes fiables, nous avons besoin d'un recensement exhaustif des situations difficiles. Lorsqu'il s'agit d'algorithmes de phonétisation, ces situations difficiles correspondent simplement à des exemples de mots. Une des nécessités préalables est donc d'avoir à sa disposition les exemples et contre-exemples pertinents, c'est-à-dire les mots concernés. La recherche d'exemples conditionne ainsi les résultats du travail, puisque les programmes sont construits à partir de ces exemples.

⁶ PL = Particularités de Lecture.

Sans bases de données linguistiques, ou sans moyens d'interroger ces bases de données, un concepteur d'algorithmes sur des mots ne peut compter que sur sa perspicacité, sur son imagination et sur les dictionnaires imprimés pour recenser des exemples. Grâce à l'ordre alphabétique, les dictionnaires usuels permettent une vision exhaustive de certains phénomènes, comme le *h* aspiré en français. À l'aide d'un dictionnaire usuel, on peut facilement obtenir la liste des mots qui présentent un *h* aspiré : en effet, presque tous s'écrivent avec un *h* initial, et donc figurent à la lettre *h* d'un dictionnaire. Mais cela tient uniquement au fait qu'il s'agit d'une propriété liée à l'initiale des mots, ce qui n'est pas le cas de toutes les propriétés à considérer. Les dictionnaires usuels sont de peu d'utilité, par exemple, pour collecter des exemples qui illustrent la prononciation de *l* final, car les mots qui se terminent par *l* sont éparpillés dans l'ordre alphabétique.

Un dictionnaire électronique permet de collecter systématiquement les exemples relatifs à un problème donné, et ce par des traitements automatiques. Un tri alphabétique inverse, c'est-à-dire en commençant par la droite des mots, permet de classer les mots par leur terminaison, et d'avoir sous les yeux, par exemple, tous les mots terminés par *l*. L'extraction de chaînes de caractères est un autre moyen simple d'obtenir automatiquement une liste d'exemples potentiels. Ce procédé sert à extraire d'un dictionnaire électronique toutes les entrées qui ont en commun une propriété. Il est applicable lorsque cette propriété correspond à une chaîne de caractères présente dans l'entrée lexicale. Ainsi, en extrayant du DELAS toutes les lignes qui comportent la séquence *ill*, on obtient une liste d'exemples qui permet de recenser les façons dont cette séquence peut se prononcer : *fil*le, *vil*le, *mail*le, *cordill*ère... Il faudrait plusieurs jours d'un travail fastidieux pour établir à la main la liste des mots en *-ill-* avec un dictionnaire ; encore serait-elle probablement très incomplète.

Plusieurs systèmes d'exploitation et éditeurs programmables offrent des commandes et des fonctionnalités intéressantes pour interroger un dictionnaire. Nous avons utilisé dans ce but l'éditeur EDT de Digital Equipment Corporation, ainsi que les systèmes suivants : RSX 11 M sur PDP 11-34, VMS sur Vax 730, et MS-DOS sur micro-ordinateur IBM. Un tri alphabétique inverse, l'extraction ou le dénombrement des entrées qui comportent une chaîne de caractères, sont des opérations simples avec ces outils. Toutefois, à plusieurs reprises, nous avons eu à utiliser des méthodes plus complexes pour interroger le DELAP :

- pour extraire non plus toutes les occurrences d'une chaîne de caractères, mais toutes les occurrences de plusieurs chaînes de caractères (afin d'obtenir la liste, par exemple, des entrées qui comportent l'une des séquences *ail*, *eil*, *oil*, *uil*) ;
- pour reconnaître les zones qui constituent une entrée, afin de considérer non plus toutes les occurrences d'une chaîne de caractères, mais seulement celles qui figurent dans une zone particulière (l'étude sur la synérèse et la diérèse présentée dans le chapitre V a demandé, par exemple, d'extraire du DELAP les entrées qui comportent la séquence /ua/ dans leur zone phonémique) ;
- pour compter les occurrences d'une chaîne de caractères sans en extraire la liste.

Lorsque ces objectifs sont complexes et se combinent entre eux, les commandes de l'éditeur programmable et des systèmes d'exploitation restent un outil efficace et

maniable, mais on doit les organiser en programmes qui peuvent prendre la forme de fichiers de commandes en *batch*, sans l'intermédiaire d'un langage de programmation. Pour les opérations qui demandent une analyse plus complexe des chaînes de caractères, par exemple une classification des groupes de consonnes, nous adjoignons à ces programmes des procédures en langage C. Enfin, pour les tris élaborés tels que les tris sur plusieurs zones à la fois, nous utilisons le système de tabulation Lexsyn, conçu par Ph. Vasseux en Fortran au LADL. L'ensemble forme un langage d'interrogation disparate, mais puissant, facile à programmer, et efficace sur Vax 730. Dans le chapitre III, nous illustrons par un exemple le fonctionnement de ce langage : le DELAP a fourni les matériaux pour l'étude des groupes de consonnes phonétiques en français, et plus précisément pour le dénombrement, le recensement et la classification des groupes de consonnes initiaux, intérieurs et finaux. Ces matériaux ont été extraits automatiquement du DELAP avec les techniques que nous venons d'évoquer.

Il va de soi que la qualité des listes d'exemples obtenues dépend directement de la qualité des dictionnaires utilisés, et notamment du nombre d'erreurs qu'ils comportent, dans les mots comme dans les informations sur les mots.

7. L'élaboration du DELAS et du DELAP

L'élaboration du système DELA s'est faite par une évolution progressive : chaque modification ou extension systématique a été précédée d'une étape d'interrogation destinée à préciser l'état des dictionnaires sur les points considérés, et à tirer parti des informations déjà présentes dans le système.

Examinons d'abord comment a été élaboré le dictionnaire DELAS. Une première liste de mots a d'abord été saisie sur support informatique. Comme dans le DELAS actuel, les mots étaient accompagnés d'informations grammaticales et morphologiques. Toutefois, cette liste n'était pas conçue comme un dictionnaire électronique, et ces informations s'apparentaient plutôt à celles que l'on trouve dans un dictionnaire usuel. Considérons par exemple la possibilité de mise au féminin des noms humains. Certains noms humains peuvent se mettre au féminin avec un changement de forme : *champion*, *championne*, ou sans changement de forme : *un artiste*, *une artiste*. D'autres peuvent référer à une femme tout en restant au masculin :

Anne est (un assassin, un bon médecin)

Ces deux exemples, d'ailleurs, ne peuvent pas se mettre au féminin, du moins pas sans un important changement de sens :

**Anne est (une assassin(e), une bonne médecin(e))*

D'autres noms humains montrent une combinaison de ces comportements, comme *avocat* : on désigne couramment une femme avocat aussi bien par *un avocat* que par *une avocate*. Ces informations ne figurent pas systématiquement dans les dictionnaires usuels, surtout lorsque l'usage est hésitant. Elles ne se déduisent pas non plus de la forme des mots. Ainsi, nous avons vu que *assassin* et *médecin* ne peuvent pas se mettre au féminin, même lorsqu'ils réfèrent à une femme. Mais les mots *citadin*, *concubin*, *gamin* et *rouquin*, qui ont la même finale, ne peuvent normalement référer à une femme

que sous la forme *citadine, concubine, gamine, rouquine*. De même, *professeur*, qui ne se met guère au féminin en français standard, s'oppose à *chanteur* qui a pour féminin *chanteuse*. De nombreuses difficultés de cet ordre expliquent que l'élaboration du DELAS à partir de cette liste ait signifié un substantiel apport d'information.

Cette élaboration s'est faite - et se continue - par des moyens variés :

- Des mots manquants ont été ajoutés, jusqu'à obtenir une importante couverture du français courant. Ces adjonctions se poursuivent.⁷ Des listes substantielles de verbes, de noms et d'adjectifs ont été dressées par J. Dubois et intégrées au DELAS.
- Des erreurs isolées ont été corrigées : fautes d'orthographe et erreurs de frappe, comme *bailler* pour *bâiller* ; erreurs dans les informations grammaticales, comme *si* marqué préposition.
- Des erreurs systématiques ont été constatées et corrigées. La mise au féminin des noms humains, évoquée ci-dessus, en est un exemple : tous les noms ont été reconsidérés pour rendre plus systématique le codage de cette information. Le DELAS lui-même a fourni des matériaux pour cette étude (B. Courtois, 1984 ; D. Tremblay, 1986).
- Des éléments nouveaux ont été élaborés à partir de données existantes : les programmes de flexion, la prononciation des mots.

En ce qui concerne cette évolution, qui se poursuit depuis quelques années, soulignons une constante qui s'est toujours vérifiée. Qu'il s'agisse de détecter des erreurs ou des lacunes, de corriger des erreurs, ou d'introduire des éléments nouveaux, on ne peut réaliser une opération systématique sur le dictionnaire qu'après avoir effectué une analyse systématique à l'aide du dictionnaire lui-même. Prenons un exemple. La mise au point des programmes de flexion et la constitution du dictionnaire DELAF ont permis d'élaborer un système de vérification orthographique, capable de détecter dans des textes les mots absents du DELAF (M. Silberstein, 1987). Grâce à ce système, appliqué à des textes de grande taille, de nombreuses adjonctions ont pu être faites, après vérification manuelle que les mots détectés existaient et étaient correctement orthographiés. La partie automatique de ce travail (extraire des textes les mots absents du DELAF) ne pouvait être réalisée qu'automatiquement, en raison de la taille du dictionnaire (500.000 mots) et des textes. La partie manuelle (sélectionner les mots corrects) ne pouvait être que manuelle, puisque cela revenait à introduire en machine des informations nouvelles. Loin d'être isolé, cet exemple est typique des traitements semi-automatiques propres à la gestion des dictionnaires électroniques. La plupart du temps, c'est un traitement de ce genre qui permet le meilleur compromis entre la rapidité d'exécution et la qualité du résultat : après une extraction automatique des données pertinentes, on peut étudier et traiter à la main les entrées lexicales obtenues, puis éventuellement leur faire subir un nouveau traitement automatique pour coder ou exploiter les informations introduites manuellement.

⁷ Il nous est arrivé de découvrir dans le dictionnaire des mots surnuméraires, et de les supprimer, mais ceci ne s'est produit que rarement. Par exemple, une entrée était consacrée à un adjectif *phonoque*, peut-être à la suite d'une erreur de saisie sur l'adjectif *phonique*, qui figurait pourtant sous sa forme correcte.

Cet exemple concernait l'extension du système. Les mêmes méthodes sont efficaces pour détecter et corriger des erreurs. De nombreuses erreurs isolées ont été détectées à la suite d'extractions automatiques. En travaillant sur des listes de mots invariables, extraites du DELAS par catégories, nous avons remarqué que si était donné comme préposition et corrigé cette erreur. Les n -grammes, c'est-à-dire les séquences de n lettres pour $n = 1$ à 5, ont été recensés et pour chaque n -gramme, la liste des mots dans lesquels il figure a été extraite. L'observation de ces listes a permis de détecter des mots dans lesquels une erreur de saisie avait amené un trigramme aberrant, comme *ill* dans *bailler* mis pour *bâiller*. De nombreuses corrections ont également pu être faites grâce au DELAS-I, dans lequel les mots sont rangés par ordre alphabétique inverse, donc regroupés par terminaisons : en effet, certaines erreurs sautent aux yeux en raison de ce classement. Le DELAS-IC, qui présente regroupés les noms en *-iste*, les noms en *-eur*, les noms en *-ier*, etc., a facilité l'étude des noms humains et le codage systématique de leur mise au féminin.

Pour introduire des informations phonétiques et construire le DELAP, nous avons suivi la même ligne directrice : tirer parti des informations déjà formalisées et saisies.

Dans la section 5 de ce chapitre, p. 41, nous avons évoqué le dictionnaire DELASP, qui comporte des mots sous leur forme usuelle, mais qui donne également, chaque fois que leur orthographe peut prêter à confusion, des informations complémentaires sur la façon dont elle doit être interprétée pour trouver la prononciation. Ainsi, l'entrée du nom *sas* dans le DELASP se présente sous la forme *sas(20.)*, où le code (20.) indique que le *s* final se prononce, car dans le plupart des mots le *s* final est muet : *cas*, *abus*, *amas*... Le DELASP est couplé avec un algorithme de phonétisation de mots connus, ainsi nommé car il calcule la prononciation des mots à partir de leur orthographe et des informations complémentaires données par le DELASP. Ce sont cet algorithme et ce dictionnaire qui ont servi à la construction du dictionnaire phonémique DELAP. C'est pourquoi nous allons revenir sur le système constitué par le phonétiseur de mots connus et le DELASP, et en évoquer la construction.

Les deux éléments du système ont été construits en même temps, car ils sont étroitement dépendants. Comme ils constituent un système complexe qui prend en compte 64.000 mots, cette construction a été progressive. De plus, comme nous disposions au départ d'importantes listes de mots, dont le DELAS, nous les avons exploitées par des traitements automatiques et semi-automatiques qui ont facilité la construction. Le point de départ était une liste d'environ 32.000 mots sous forme orthographique, issue du dictionnaire de L. Warnant (1964) mis sur support électronique. Lors des premiers essais de construction d'un phonétiseur, il est apparu qu'une solution commode est de transcrire les chaînes orthographiques en chaînes phonémiques en cinq passes environ, c'est-à-dire par l'intermédiaire de chaînes mixtes dans lesquelles seulement une partie des lettres ont été examinées et transformées en phonèmes. Après ces essais, nous avons aussi quelques conclusions sur l'ordre dans lequel les différentes transformations peuvent prendre place. Voyons-en deux exemples. Le calcul de la prononciation des *e* doit intervenir avant la simplification des consonnes

doubles, comme il ressort d'exemples tels que *seller* et *geler*. Si la consonne double de *seller* était simplifiée lors d'une étape antérieure au calcul de la prononciation du *e*, ce dernier calcul deviendrait impossible, car on ne pourrait plus distinguer les cas de *seller* et de *geler*. Par ailleurs, il est prudent de ne transcrire les voyelles accentuées que lors d'une des dernières étapes du traitement : lors des étapes qui ont lieu avant la transcription de ces voyelles, la présence des accents peut constituer une information utile pour calculer la prononciation des lettres voisines. Par exemple, l'examen des mots *gaélique* et *praesidium* montre que le calcul de la prononciation de *ae* doit intervenir avant la transcription de *é* en /e/.

Pour construire l'algorithme, nous avons découpé le problème de la phonétisation en sous-problèmes que nous avons considéré séparément : la prononciation de *-ch-*, celle de *-e-*, celle de *-ill-*, celle de *-en-*, etc. Chacun de ces sous-problèmes a donné lieu à l'écriture d'une procédure en langage C. Ces procédures sont celles qui effectuent les transformations de chaînes de caractères. Dans l'état actuel de l'algorithme, elles sont au nombre de 53 et chacune est appelée lors d'une des cinq passes. Pour la construction de chacune de ces procédures de transcription, le problème se posait ainsi : recenser les difficultés, exceptions et ambiguïtés qui font obstacle à la transcription automatique ; signaler chacune de ces "particularités de lecture" par une marque spécifique dans le DELASP ; exploiter cette information dans une procédure qui effectue la transcription. Par exemple, pour la transcription de la séquence *ch*, il fallait chercher dans quels mots cette séquence se prononce [k] et introduire cette information en machine. Il s'agit bien d'une ambiguïté de l'orthographe. En effet, cette information ne se déduit pas de la forme du mot par des règles locales claires, comme le montrent les paires suivantes, qui opposent chacune un mot où *ch* se prononce [k] et un mot où *ch* a sa valeur phonétique la plus fréquente :

<i>lichen</i>	<i>pichenette</i>
<i>achillée</i>	<i>Achille</i>
<i>bronchoscopie</i>	<i>bronchite</i>

C'est pourquoi les exemples de ce type ont été examinés individuellement et marqués. Pour chacune des procédures de transcription, les exemples douteux ou difficiles ont été passés en revue.

Cela n'impliquait pas d'examiner à chaque fois les 32.000 mots de la liste utilisée, car seule une proportion des mots est concernée par un problème de phonétisation donné. Par exemple, dans le cas de *ch*, les mots qui comportent la séquence *ch* ont été vus individuellement. Toutefois, ce travail n'aurait pas été réalisable sans l'aide de procédures interactives, spécialement écrites dans ce but. Ces procédures sélectionnent automatiquement les mots à examiner et les font défiler à l'écran. Pour chaque mot, l'opérateur peut voir si la transcription se fera correctement suivant les règles les plus générales, ou si une marque doit être introduite dans le DELASP pour signaler une particularité. Il communique cette information au programme en frappant une touche et, suivant le cas, l'entrée du DELASP reste inchangée ou subit le marquage nécessaire. Ce traitement produit donc une nouvelle version du DELASP. Il permet aussi d'examiner systématiquement les cas difficiles et les cas douteux pour la phonétisation, après quoi la procédure de transcription peut être réalisée en connaissance de cause. Ce

processus a été appliqué successivement pour chacun des problèmes de transcription qui ont été rencontrés. Une fois les procédures de transcription correspondantes réalisées, l'algorithme de phonétisation de mots connus et le DELASP formaient un système cohérent apte à phonétiser les 32.000 mots de cette version.

Les développements ultérieurs ont consisté en extensions successives, jusqu'à atteindre les 64.000 mots de la version actuelle. Lors de chaque extension, les nouveaux mots ont été codés à l'aide du système mis au point précédemment. Bien sûr, certains mots nouveaux ont posé des problèmes de phonétisation automatique qui n'avaient pas été rencontrés jusque-là : l'algorithme et le système de codage des "particularités de lecture" ont donc été étendus pour pouvoir prendre en compte ces nouveaux exemples.

Le phonétiseur de mots connus, en accédant au DELASP, génère automatiquement le dictionnaire DELAP en quelques dizaines de minutes. Pour la maintenance et la mise à jour du DELAP, il suffit donc d'effectuer les corrections et les extensions sur le phonétiseur de mots connus et sur le DELASP, puis de générer une nouvelle version du DELAP à chaque modification. C'est cette méthode qui est utilisée.

Chapitre III. Le dictionnaire DELAP¹

Le dictionnaire DELAP est un dictionnaire phonémique et flexionnel du français. Il fait partie du système de dictionnaires électroniques DELA, que nous avons présenté dans le chapitre précédent. Nous avons insisté sur le fait qu'un dictionnaire électronique donne des informations explicites, formalisées et exploitables par des programmes. C'est bien le cas du DELAP : les algorithmes de phonétisation et de flexion opèrent directement sur le dictionnaire, accèdent aux informations qui y figurent et les exploitent. La maintenance et la mise à jour du dictionnaire se font également par des programmes qui comparent le DELAP avec le DELAS, détectent les entrées présentes dans l'un des deux et non dans l'autre, permettent à l'utilisateur d'ajouter des entrées, d'en supprimer, d'en corriger et d'en modifier. D'autres programmes contrôlent l'ordre alphabétique des mots ainsi que la structure des entrées.

1. La structure du dictionnaire et des articles

Nous avons vu dans le chapitre précédent, section 1, que la notion de dictionnaire électronique correspond à une structure de données, particulièrement bien adaptée au traitement informatique des langues naturelles et caractérisée par les points suivants :

- Chaque article comporte un champ qui sert de clé d'entrée au dictionnaire.
- Chaque article comporte, outre la clé d'entrée, des données sur le mot concerné.
- Les éléments modifiables minimaux sont les articles du dictionnaire.

Pour des raisons de commodité, l'orthographe d'un mot étant la seule clé commode pour l'identifier, chaque article du DELAP comporte la forme orthographique du mot concerné, et les entrées sont rangées suivant l'ordre alphabétique des formes orthographiques. La présentation est donc proche de celle des dictionnaires phonétiques usuels.

Le dictionnaire occupe environ 2 Mo, c'est pourquoi il est divisé en 26 fichiers, chacun regroupant les mots qui commencent par une même lettre. Ainsi, les fichiers occupent chacun une quantité de mémoire suffisamment limitée pour qu'on puisse accéder à leur contenu directement par éditeur.

Chaque article occupe une ligne et est structuré en trois zones. Il comporte la forme orthographique du mot, il en donne une transcription phonémique, et il fournit les informations nécessaires à sa flexion : catégories grammaticales, conjugaisons, pluriels, féminins. L'article est donc constitué de trois zones séparées par des virgules :

discothèque, diskotek, N21

La zone orthographique, *discothèque* dans cet exemple, est donnée en minuscules afin que les accents, cédilles et trémas éventuels y apparaissent : la représentation

¹ Dictionnaire Electronique du LADL pour la Phonétique.

DISCOTHEQUE serait moins informative, car, conformément à l'usage actuel en France pour le codage informatique des caractères accentués, l'accent de è n'y figure pas². En particulier, et toujours pour prendre en compte les caractères accentués, l'initiale du mot est sous forme minuscule dans le DELAP, même lorsqu'elle peut apparaître sous forme majuscule dans les textes : *écosais* et non *Ecossais*. Rappelons que les caractères accentués du français ne sont pas représentés dans le code ASCII strict, mais dans un code ASCII étendu, et que chaque constructeur d'ordinateurs utilise un code ASCII étendu différent. Comme le DELAP est actuellement entretenu et utilisé sur un Vax 730 (Digital Equipment Corporation), les caractères accentués sont dans le code ASCII étendu "DEC multinational". Les programmes de flexion et de phonétisation utilisent également ce codage. Il nous arrive de porter le DELAP sur un micro-ordinateur IBM et de l'y utiliser, mais cela pose des problèmes en raison des disparités dans le codage des caractères accentués.

La zone phonémique de notre exemple est /diskotek/. Conformément à l'usage des linguistes, nous citons les représentations phonémiques entre barres obliques dans le texte, mais dans le dictionnaire les trois zones sont délimitées par des virgules, qui gênent moins la lecture. La zone phonémique comporte une représentation de la prononciation du mot sous la forme d'une chaîne de phonèmes. Cette représentation est proche d'une transcription phonétique, mais plus informative et moins redondante : il s'agit d'une représentation abstraite qui permet de traiter simplement les variations phonétiques des mots, et qui tient compte de connaissances linguistiques sur le lexique français. Les sections 4 à 7 de ce chapitre présentent plus en détail les formes phonémiques du DELAP, ainsi que les relations entre ces formes et les formes phonétiques. Pour l'instant, prenons seulement l'exemple du nom *discothèque*, dont nous donnons la transcription phonémique /diskotek/. Dans une transcription phonétique, il aurait fallu tenir compte du fait que le *o* est sujet à des variations phonétiques. Il se prononce plus ou moins ouvert selon les locuteurs et varie couramment entre deux extrêmes : [o], le plus fermé, et [ɔ], le plus ouvert³. Par ailleurs, dans *discothèque*, *è* se prononce obligatoirement ouvert [ɛ]. Pour donner une description phonétique exacte de ce mot sans faire abstraction des variations phonétiques qui l'affectent dans le langage standard, il faudrait donc citer au moins [diskotɛk] et [diskɔɛk], mais pas [diskotek] ni [diskɔtek]. La représentation phonémique /diskotek/ est plus lisible et permettra de retrouver ces deux formes ; elle traitera alors d'une façon plus générale la variation phonétique entre [o] et [ɔ], qui est une variation systématique, c'est-à-dire étendue à un grand nombre de mots dans le lexique.

La zone flexionnelle d'un article du DELAP contient les informations sur les variations de forme du mot, sur ses propriétés grammaticales et sur ses variations phonétiques éventuelles lors de la flexion, c'est-à-dire lors de la conjugaison, de la mise au pluriel, de la mise au féminin. Dans l'exemple *discothèque*, la zone flexionnelle est

² Le codage en usage au Québec distingue les lettres accentuées des autres, aussi bien en majuscules qu'en minuscules.

³ Les représentations phonétiques sont données entre crochets pour les distinguer des représentations phonémiques, qui sont données entre barres obliques.

représentée par le code .N21, qui correspond aux informations suivantes :

- Nom féminin qui peut être au singulier ou au pluriel.
- Le pluriel orthographique s'obtient en ajoutant -s au singulier.
- Le singulier et le pluriel se prononcent de la même façon et reçoivent une même transcription phonémique.

Le chapitre IV décrit le système de flexion orthographique et phonémique du français ainsi que le système de codage utilisé pour expliciter les variations flexionnelles des mots.

2. Les entrées multiples

Certains mots peuvent se prononcer de plusieurs façons, mais sans changement d'orthographe. C'est le cas de l'adjectif au masculin *exact* dont la finale *-ct* peut se prononcer ou être muette. Cette variation ne se retrouve que dans quelques autres mots : elle s'oppose ainsi aux variations de timbre des voyelles du mot *discothèque*, qui sont beaucoup plus systématiques. Cependant, les variations isolées telles que celles de *exact* ont leur importance, car elles sont nombreuses. Dans ce cas, le mot a une seule forme orthographique, mais se voit attribuer plusieurs formes phonémiques : /egza/ et /egzakt/. Pour respecter la structuration de chaque entrée en trois zones, nous avons introduit dans le DELAP autant d'entrées que de formes phonémiques distinctes :

1 = *exact*, egza, .A32

2 = *exact*, egzakt, .A32

Le mot a alors une seule entrée dans le DELAS mais plusieurs dans le DELAP. Ces différentes entrées sont alors liées par une relation d'équivalence à l'intérieur du DELAP en ce sens qu'elles correspondent à une même orthographe. Toutefois, il faut distinguer deux cas possibles pour ces variantes de prononciation.

Dans le cas de *exact*, les deux prononciations correspondent au même adjectif : il s'agit bien de variantes phonétiques d'un même mot. L'équivalence entre les variantes est à plusieurs niveaux. Non seulement elles ont la même orthographe, mais elles appartiennent à la même catégorie grammaticale et elles ont la même flexion orthographique. Elles ont les mêmes emplois syntactico-sémantiques, c'est-à-dire que toute phrase qui comporte le mot *exact* admet ses deux prononciations :

Ce renseignement est exact

[egza]

[egzakt]

Les deux variantes sont interchangeables avec, tout au plus, des différences stylistiques qu'il serait difficile de prendre en compte systématiquement.

Ce n'est pas la seule situation possible. Les deux prononciations de la forme orthographique *poster* correspondent à deux mots différents, un verbe et un nom⁴ :

Guy va <i>poster</i> une lettre	[poste]	*[poster]
Guy a affiché un <i>poster</i> sur le mur	*[poste]	[poster]

La relation entre les deux représentations phonémiques correspondantes, /poste/ et /poster/, consiste seulement en une confusion orthographique : les deux mots appartiennent à des catégories grammaticales distinctes ; ils n'ont pas la même flexion orthographique ni les mêmes emplois syntactico-sémantiques. Ainsi, nous avons donné ci-dessus deux phrases pour illustrer un emploi syntactico-sémantique pour chacun des deux mots *poster*.

Les cas de *exact* et de *poster* sont bien différents, et nous avons voulu rendre cette distinction explicite dans le DELAP. Cela supposait que nous disposions d'un ou plusieurs critères pour décider, pour chaque mot qui admet plusieurs prononciations, s'il s'apparente au type *exact* ou au type *poster*. Le meilleur critère utilisable, d'un point de vue théorique comme du point de vue des applications informatiques, est celui des emplois syntactico-sémantiques. La différence essentielle entre *exact* et *poster* est que les deux prononciations de *exact* s'emploient toujours dans les mêmes phrases, alors que les deux prononciations de *poster* s'emploient dans des phrases distinctes. Ce critère pourra être utilisé lorsque les emplois syntactico-sémantiques des mots seront répertoriés dans le cadre d'un lexique-grammaire (Maurice Gross, 1981) et lorsque la correspondance entre ce lexique-grammaire et le DELAP sera assurée.

Pour distinguer les cas de *exact* et *poster*, nous utilisons pour l'instant un autre critère, qui se trouve être une bonne approximation du précédent : la flexion orthographique. Orthographiquement, le mot *exact* a une mise au féminin et une mise au pluriel bien déterminées, indépendamment de ses deux prononciations. Au contraire, les deux mots *poster* ont chacun leur catégorie grammaticale : l'un se conjugue, l'autre est un nom masculin qui peut se mettre au pluriel. Nous avons marqué cette différence dans le DELAP de la façon suivante : l'adjectif *exact* a deux entrées numérotées 1 et 2 ; le mot *poster* a également deux entrées numérotées 1 et 2, mais le séparateur entre le numéro et la forme orthographique est = dans le cas de *exact* et # dans le cas de *poster*.

1=*exact*,egza,,A32

2=*exact*,egzakt,,A32

1#*poster*,poste,,V3

2#*poster*,poster,,N1

Le séparateur qui suit le numéro est = si les entrées ont la même catégorie grammaticale et la même flexion orthographique, et # dans le cas contraire. Nous avons codé dans le DELAP 774 entrées multiples, dont 720 du type *exact*, 45 du type *poster*, et 9 mixtes. Ces dernières servent au cas où une orthographe correspond à au moins trois

⁴ L'étoile symbolise un jugement d'acceptabilité sur la transcription phonétique qu'elle précède. Une étoile placée devant une transcription phonétique indique que la prononciation transcrite est interdite, c'est-à-dire qu'elle n'est pas employée, du moins dans le français standard. Parfois, il est difficile d'émettre un jugement d'acceptabilité incontestable : ces cas douteux sont signalés par un point d'interrogation, également placé devant la transcription ou devant l'étoile.

prononciations, dont au moins deux relèvent d'une même flexion orthographique, et au moins une autre relève d'une autre flexion orthographique. Toutes les entrées multiples sont données dans l'annexe 3. Une proportion importante d'entre elles comporte des mots d'origine étrangère dont la prononciation française n'est pas fixée.

3. Les mots du DELAP

Le DELAP décrit les mêmes mots que le DELAS, qui a été présenté dans le chapitre précédent. Rappelons que le DELAS comporte des mots simples, c'est-à-dire écrits sans séparateur. Les mots qui s'écrivent avec un trait d'union, comme *sang-froid*, avec un espace, comme *vin blanc*, ou avec une apostrophe, comme *aujourd'hui*, ne figurent donc pas dans le DELAS. En revanche, les mots simples qui n'apparaissent que dans des mots composés sont systématiquement introduits dans le DELAS : c'est le cas de *aujourd* et *hui* pour *aujourd'hui*, de *encontre* pour *à l'encontre de*, de *priori* pour *a priori*, etc. Le DELAS ne comporte pas de noms propres. Compte tenu de ces règles, le DELAS assure une couverture lexicale du français standard aussi complète que possible, excepté dans les domaines techniques spécialisés.

Les mêmes mots figurent dans le DELAS et dans le DELAP. Cette correspondance entre les deux dictionnaires est maintenue lors de leur évolution. Pour cela, un système de programmes de comparaison a été écrit en langage C. Ces programmes contrôlent le respect de l'ordre alphabétique dans les deux dictionnaires. Ils détectent les mots qui figurent dans l'un des deux et non dans l'autre. Quant aux mots qui figurent bien dans les deux dictionnaires, les programmes contrôlent que les informations sur la flexion orthographique concordent. Ils reconnaissent également les articles multiples du DELAP : ils contrôlent leur numérotation et leur correspondance avec les articles du DELAS. Toutes les situations anormales détectées donnent lieu à des messages qui sont envoyés soit à l'écran soit sur un fichier. Les programmes Complap et Complap permettent également à l'opérateur de corriger les lacunes et les erreurs d'une manière interactive, ce qui produit une nouvelle version du DELAP ou du DELAS, suivant le cas. Par exemple, lorsque le programme Complap ne trouve pas un mot du DELAS à son ordre alphabétique dans le DELAP, l'opérateur peut choisir dans un menu l'une des options suivantes :

- Introduire le mot dans le DELAP avec sa prononciation. En ce cas, l'opérateur est invité à coder la prononciation au clavier.
- Introduire le mot dans le DELAP, mais sans sa prononciation, car elle sera calculée automatiquement. En effet, le phonétiseur de mots connus évoqué dans le chapitre II, section 4, p.31, calcule la prononciation des mots lorsque l'orthographe est suffisamment explicite pour que cela soit possible sans ambiguïté.
- Ne pas introduire le mot. En effet, la présence du mot dans le DELAS peut résulter d'une erreur ; le mot peut aussi être mal placé dans l'ordre alphabétique ; enfin, si l'opérateur ne connaît pas le mot, ne sait pas le prononcer, et a besoin de faire une recherche spécifique avant de l'introduire dans le DELAP, il ne l'introduit pas : lors de la comparaison suivante, si le mot a été maintenu dans le DELAS, l'opérateur est à nouveau invité à l'introduire dans le DELAP.

L'opérateur a également le choix entre plusieurs solutions lorsqu'un mot figure bien dans les deux dictionnaires, mais que les informations sur la flexion orthographique ne concordent pas. Le programme Complap lui permet alors soit de maintenir les informations qui figurent dans le DELAP, soit de les remplacer par celles du DELAS ou par une troisième version si les deux premières lui semblent erronées.

Le programme Complap permet d'effectuer les mêmes opérations, mais en produisant une nouvelle version du DELAS et non du DELAP.

Ce système de comparaison a été utilisé intensivement lors de l'élaboration progressive du DELAP. Il a également contribué à l'extension et à la correction du DELAS.

En 1985, le DELAS comportait près de 50.000 mots et le DELAP 32.000, car ils avaient été construits indépendamment. Le système de comparaison a été réalisé et la première confrontation a eu lieu, ce qui a produit deux listes : d'une part, quelques centaines de mots qui figuraient dans le DELAP et non dans le DELAS, et d'autre part plusieurs milliers de mots qui figuraient dans le DELAS et non dans le DELAP. Ces listes ont été examinées manuellement. Des erreurs ont été corrigées et des mots régionaux ont été mis à part, car dans la première version du DELAP la liste des mots était issue du dictionnaire de L. Warnant, qui comporte des mots spécifiques au français de Wallonie et du nord de la France. Les listes ainsi modifiées ont été incorporées au DELAS et au DELAP à l'aide du système de comparaison, ce qui a rendu effective la correspondance entre les deux dictionnaires, avec un même ensemble de 50.000 entrées environ. Cette correspondance a été constamment entretenue depuis. Les erreurs détectées et corrigées lors de la maintenance ou de l'exploitation d'un des dictionnaires ont donc également été corrigées dans l'autre. En 1987, les dictionnaires ont été étendus à 64.000 entrées. Dans une première étape, l'extension a été réalisée entièrement, mais seulement sur le DELAS. Dans une seconde étape, le système de comparaison a permis de dresser automatiquement la liste des 14.000 mots nouveaux. Cette liste a été examinée à la main. Des erreurs qui s'étaient produites lors de l'extension du DELAS ont été détectées et corrigées. Les problèmes liés à la représentation phonémique des mots nouveaux ont été abordés et résolus sans difficulté majeure, car les choix faits pour la représentation phonémique dans la version précédente se sont révélés suffisamment généraux. Les algorithmes de phonétisation ont subi quelques retouches à cette occasion. Les 14.000 mots nouveaux ont alors été incorporés au DELAP à l'aide du système de comparaison, ce qui a achevé de donner au DELAP sa forme actuelle (février 1988).

4. Brève introduction aux représentations phonétiques et phonémiques

L'une des spécificités du DELAP est la représentation phonémique qu'il donne de chaque entrée. Il s'agit d'une chaîne de phonèmes qui caractérise la phonétique du mot, y compris ses variations phonétiques éventuelles.

Une représentation phonétique est constituée d'une chaîne de signes phonétiques qu'on cite entre crochets et qui reflète la prononciation aussi précisément et aussi

directement que possible : [aks] pour *axe*. En particulier, tous les signes phonétiques se prononcent. Par exemple, le nom *fin* s'écrit avec une consonne finale *n*, mais dans la prononciation de ce mot la voyelle n'est pas suivie d'une consonne. Une représentation phonétique de *fin* ne comporte donc pas de consonne finale : [fɛ̃]. Par ailleurs, dans la mesure du possible, on utilise des symboles phonétiques distincts pour représenter des sons distincts. Ainsi, les voyelles des mots *né* et *net* sont différentes : la première est plus fermée que la seconde. On note donc ces voyelles par deux sons différents, [e] et [ɛ] :

<i>né</i>	[ne]	<i>net</i>	[nɛ]
-----------	------	------------	------

Ces représentations phonétiques sont déjà des abstractions, car on y note d'un même signe phonétique des sons qui, en fait, ont des variations en fonction du contexte : ainsi, le [n] est légèrement différent dans *né* et dans *net*. C'est pourquoi plusieurs auteurs français utilisent le terme de "phonème" pour désigner les signes phonétiques, mais nous réservons ce terme à un autre usage.

Jusqu'aux environs des années 1930, on exprimait la prononciation des mots exclusivement à l'aide d'alphabets phonétiques. Puis un besoin nouveau est apparu : connaître quelles combinaisons de sons existent dans une langue donnée, et aussi connaître les variations phonétiques que subissent ces combinaisons. Par exemple, en français, [e] n'est jamais dans la position qu'occupe le [ɛ] de *net* [nɛ], car aucune phrase française ne se termine par la séquence [et]. Dans la conjugaison du verbe *céder*, [e] et [ɛ] sont associés : *céder* se prononce généralement [sede], et (*il*) *cède* [sed]. On peut donc parler d'une variation phonétique entre [e] et [ɛ] dans ce verbe. Pour étudier ces combinaisons et ces variations, on a défini des objets proches de ceux représentés par les signes phonétiques, mais plus abstraits. Ces objets ont été appelés des phonèmes. (Les auteurs qui emploient le terme de "phonème" pour désigner les symboles phonétiques ont recours à celui d'"archiphonème" pour les objets qui sont des abstractions par rapport aux signes phonétiques. Toutefois, nous préférons appeler ces abstractions des phonèmes, car c'est avec ce sens que le terme a été créé.) Par exemple, si on considère que le [e] de *né* [ne] et le [ɛ] de *net* [nɛ] sont des variantes l'un de l'autre, on les représente par un même phonème abstrait /e/, qu'on cite entre barres obliques, et non entre crochets comme les signes phonétiques. On peut alors représenter les mots *né* et *net* par des séquences de phonèmes, également citées entre barres obliques : *né* /ne/, *net* /net/, et de même *céder* /sede/, (*il*) *cède* /sed/. Dans ces représentations phonémiques, la différence entre les voyelles [e] et [ɛ] n'est pas explicite, mais cela n'entraîne pas de perte d'information : en effet, comme aucun mot français ne comporte un [ɛ] dans la position du [e] de *net*, des règles de redondance permettent de retrouver que le /e/ de /net/ représente un [ɛ] et non un [e]. La représentation phonémique est simplifiée par rapport à la représentation phonétique : elle ne comporte pas l'information expresse que la voyelle de *net* est un [ɛ], car cette information est déductible du contexte, c'est-à-dire redondante. Nous reviendrons sur les variations entre [e] et [ɛ] dans la section 7, p. 59, et sur les règles de redondance dans la section 8, p. 68.

Il est à noter que les représentations phonémiques laissent de côté trois types d'informations sur la prononciation :

- les paramètres prosodiques, c'est-à-dire la durée, l'intensité et la fréquence fondamentale ;
- les propriétés acoustiques des phonèmes ;
- l'articulation de la parole.

Ces trois types de données phonologiques concernent elles aussi la prononciation. Elles devraient être intégrées dans une représentation plus complète de la prononciation. Toutefois, de nombreux problèmes de description et de formalisation qui restent à résoudre peuvent s'exprimer en termes de représentations purement phonémiques ; notre travail se situe à l'intérieur de ces limites.

Les linguistes ont conçu les phonèmes abstraits et utilisé des transcriptions phonémiques afin d'explicitier seulement les informations phonétiques pertinentes, à l'exclusion des informations phonétiques qui apparaissent comme redondantes compte tenu de l'ensemble de la langue. Par la suite, notamment avec Z.S. Harris (1951) puis avec la phonologie générative (N. Chomsky et M. Halle, 1968), on a commencé à concevoir des représentations phonémiques plus informatives que les représentations phonétiques, et par exemple à introduire des informations grammaticales dans les représentations phonémiques. Ainsi, Z.S. Harris (1951, p. 225) propose de représenter par deux phonèmes distincts le [f] final du nom anglais *knife* et celui de *fife* parce que *knife* a pour pluriel *knives* [najvz] alors que *fife* a pour pluriel *fifes* [fajfs]. N. Chomsky et M. Halle (1968) introduisent des frontières d'éléments morphologiques dans les chaînes de phonèmes, notamment afin de représenter l'accent en anglais. En français, un traitement analogue à celui proposé par Z.S. Harris a été utilisé pour les adjectifs au masculin tels que *petit*, qui se prononce sans consonne finale, mais dont le féminin *petite* comporte un [t] final : dans certains systèmes phonémiques, ces adjectifs au masculin sont représentés avec la consonne finale. Les chaînes phonémiques de ce genre permettent de traiter d'une façon simple certaines interactions entre grammaire et phonétique. Depuis ces développements théoriques, on peut définir les formes phonémiques des mots comme des représentations plus ou moins abstraites à l'aide desquelles on peut décrire simplement des faits phonétiques observés, éventuellement en relation avec d'autres phénomènes linguistiques, par exemple grammaticaux, syntaxiques ou lexicaux.

Il est à noter que depuis E. Sapir (1933), de nombreux phonologues, notamment en phonologie générative, ont utilisé et utilisent des représentations phonémiques dans un autre but : représenter l'organisation inconsciente du savoir linguistique du sujet parlant. Ces représentations sont alors dites "psychologiquement plausibles", c'est-à-dire susceptibles de ressembler aux connaissances linguistiques implicites des locuteurs de la langue étudiée. Toutefois, il s'agit d'un sujet difficile qui pose des problèmes méthodologiques et que nous n'abordons pas. Nous n'attribuons donc pas de valeur ni de plausibilité psychologique aux représentations phonémiques que nous donnons.

5. L'utilisation de représentations phonémiques dans des applications informatiques

La description phonétique est un investissement qui débouche sur des applications informatiques importantes, notamment la synthèse et la reconnaissance de la parole, ainsi que la correction orthographique de mots. Ces systèmes, dont nous parlerons au chapitre VI, manipulent des informations phonétiques représentées sous une forme ou sous une autre.

Dans de petits systèmes, peu importe que ces informations soient manipulées sous forme de représentations phonétiques ou phonémiques, car la complexité et le volume de ces informations sont de toute façon réduits. Toutefois, pour une utilisation industrielle, un système informatique doit être suffisamment robuste et performant. Or les systèmes de synthèse et de reconnaissance de la parole, ainsi que ceux de correction orthographique, ne pourront atteindre cette robustesse et ces performances que si un volume important de données formalisées est réuni et mis en jeu. Etant donné la complexité de ces données, il est essentiel qu'elles soient représentées sous une forme naturelle et qui facilite les traitements. Dans le cas de données sur la prononciation des mots et sur ses variations, les représentations phonémiques constituent le formalisme qui répond le mieux à cette contrainte, car elles sont conçues pour représenter le fonctionnement du système linguistique. Si on s'interdisait l'emploi de représentations phonémiques et si on utilisait exclusivement des représentations purement phonétiques à toutes les étapes du traitement, cela entraînerait des complications supplémentaires qui conduiraient soit à renoncer à prendre en compte les variations phonétiques, soit à limiter à sa plus simple expression le vocabulaire reconnu, donc dans les deux cas à réduire considérablement les performances réalisables.

Considérons par exemple les systèmes informatiques dont les entrées sont des textes écrits ou parlés. Pour analyser ou traiter les mots des textes, les programmes doivent inclure des connaissances grammaticales plus ou moins approfondies, par exemple en ce qui concerne la conjugaison des verbes, car les verbes qui apparaissent dans les textes écrits ou parlés peuvent être conjugués. Ces programmes manipulent donc à la fois des informations sur la prononciation des mots et des informations grammaticales. Ces deux sortes de données ont des interactions, et la façon la plus simple de manipuler ces concepts est d'utiliser les représentations phonémiques conçues par les linguistes. Si on cherchait à n'utiliser que des représentations phonétiques, on se heurterait à des difficultés liées notamment aux variations phonétiques des mots. Ainsi, nous avons vu que le radical du verbe *céder* connaît des variations phonétiques lors de la conjugaison : dans *il cède*, on prononce obligatoirement un [ɛ] ouvert, alors que dans l'infinitif *céder*, on prononce généralement un [e] fermé. Il en est de même de nombreux verbes courants. Des programmes qui utiliseraient seulement des représentations phonétiques auraient à tenir compte de ces variations. Comme ces variations affectent le radical du mot et s'ajoutent aux variations normales de la conjugaison, qui, elles, affectent la terminaison, c'est un modèle complexe qu'il faudrait construire et implanter en machine, alors que les systèmes de représentations phonémiques sont des modèles conçus pour traiter ces variations d'une manière simple. De nombreux exemples

pourraient s'ajouter à celui du verbe *céder* : on observe des variations phonétiques analogues dans le radical de verbes comme *coller*, *effeuiller*, *peser*, *geler*, etc.

De plus, à ces variations qui ont un caractère grammatical s'ajoutent des variations libres, dont nous n'avons pas toujours conscience lorsque nous parlons mais qui sont un obstacle sérieux à la reconnaissance automatique de la parole. Prendre en compte ces variations en utilisant uniquement des représentations phonétiques reviendrait à lister purement et simplement toutes les variantes phonétiques de toutes les formes. Un tel système de données poserait des problèmes de maintenance et de mise à jour ; la cohérence et l'exhaustivité des informations seraient difficiles à contrôler. Il est préférable de générer automatiquement les variantes à partir de données formalisées sur les formes de base et sur les variations qu'elles peuvent subir.

Ces deux exemples visaient seulement à montrer que l'emploi de représentations phonémiques permet de réaliser des applications informatiques importantes. Nous verrons plus en détail plusieurs autres situations du même type.

6. Le caractère arbitraire des représentations phonémiques

La phonétique concerne la description des sons. La phonémique prend en compte principalement les sons, mais aussi la langue dans son ensemble, car certaines variations phonétiques, par exemple, sont liées à d'autres structures linguistiques comme la conjugaison.

Cette différence entre phonétique et phonémique a des conséquences importantes. Il est théoriquement possible de décrire la prononciation de toutes les langues à l'aide d'un alphabet phonétique unique, du moment qu'il contient assez de signes. L'alphabet phonétique international est d'ailleurs un standard couramment utilisé par les phonéticiens pour les langues du monde entier. En revanche, un alphabet phonémique n'est valable que pour une langue, car il se réfère au lexique et à la grammaire de cette langue. Même pour les différents dialectes ou les accents régionaux d'une langue, des alphabets phonémiques légèrement différents sont parfois nécessaires.

Par ailleurs, comme la phonétique est une description directe de faits observables, les principaux problèmes que rencontrent les phonéticiens ont trait à l'observation des faits. Mis à part ces difficultés d'observation, les résultats de la description phonétique n'ont guère de raisons de diverger d'un auteur à l'autre. Il n'en est pas de même des systèmes phonémiques abstraits. Ceux-ci reposent sur une analyse des relations entre la prononciation et les autres structures de la langue étudiée : aux problèmes d'observation s'ajoutent donc des problèmes spécifiquement linguistiques, qui sont loin d'être résolus et qui peuvent souvent être formalisés de plusieurs façons. C'est pourquoi le consensus entre les auteurs est plus difficile que dans le domaine purement phonétique. En particulier, les spécialistes proposent des représentations phonémiques des mots, mais bien des solutions différentes sont possibles, sans qu'on dispose de critères applicables et clairs pour choisir parmi des solutions concurrentes. Dans l'état actuel des connaissances, l'élaboration et le choix de représentations phonémiques comporte donc une part d'arbitraire.

7. La lecture des formes phonémiques dans le DELAP

Le dictionnaire DELAP donne des formes phonémiques de 64.000 mots français. Ces formes phonémiques couvrent donc une bonne partie de la langue courante. Cela ne signifie pas que ce système apporte des solutions optimales et cohérentes à l'ensemble des problèmes que pose la description phonétique du français. Cependant, ce dictionnaire phonémique est aussi un outil qui s'est révélé utile pour étudier certains de ces problèmes, et qui peut être utilisé pour éclaircir en connaissance de cause les nombreux points qui restent à résoudre.

Nous allons donner les principales conventions nécessaires pour lire, comprendre et utiliser la représentation phonémique du DELAP. Une des difficultés est que cette représentation ne suit pas les règles de lecture de l'orthographe, ce qui demande au lecteur une certaine attention. Les autres difficultés de lecture des formes phonémiques résultent de points où la transcription phonémique s'écarte de la transcription phonétique. Nous allons passer en revue ces différences.

Les voyelles /e/, /ø/, /o/

En français, on observe des variations complexes entre la voyelle [e] de *pré* et la voyelle [ɛ] de *presse*. Dans certains mots, on prononce généralement [e] ou obligatoirement [ɛ], en fonction de la structure de la syllabe et de sa position dans le mot :

<i>pré</i>	[pre]	?[prɛ]
<i>presse</i>	*[pres]	[pres]
<i>céder</i>	[sede]	?[sɛde]
(il) <i>cède</i>	*[sed]	[sed]

Dans ces exemples, il est justifié de représenter [e] et [ɛ] par un même phonème abstrait, puisque la distribution entre les deux voyelles dépend du contexte. Dans d'autres mots, la prononciation dépend beaucoup du locuteur, et on prononce [e], [ɛ] ou un intermédiaire :

<i>lait</i>	[le]	[lɛ]
<i>après</i>	[apre]	[aprɛ]
<i>mêler</i>	[mele]	[mɛle]
<i>serrer</i>	[sere]	[sɛre]
<i>baisser</i>	[bese]	[bɛse]

Ces variations existent depuis longtemps en français (Ph. Martinon, 1913, p. 56) et mériteraient une étude spécifique étendue à tout le lexique. Pour l'instant, nous négligeons provisoirement la variation entre [e] et [ɛ] dans les exemples tels que *lait*, *après*, *mêler*, *serrer*, *baisser*, et nous les notons /e/ comme dans *pré* :

<i>pré</i>	/pre/	<i>lait</i>	/le/
<i>après</i>	/apre/	<i>mêler</i>	/mele/
<i>céder</i>	/sede/	<i>serrer</i>	/sere/
(il) <i>cède</i>	/sed/	<i>baisser</i>	/bese/

Le dictionnaire DELAP fournit les matériaux pour étudier les variations entre [e] et [ɛ] en prenant en compte l'ensemble du lexique : par des traitements automatiques, il est possible d'obtenir la liste des mots terminés par [e] ou par [ɛ], comme *pré* et *après*.

Les voyelles [e] et [ɛ] sont pour l'instant représentées d'une manière approximative, puisqu'elles sont confondues. Ce traitement imprécis convient bien à l'une des applications informatiques, la correction orthographique par phonétisation, présentée au chapitre VI. En effet, en raison de ce choix, le système de correction peut proposer, par exemple, *procès* à la place de *procés*, *léser* à la place de *laiser*, *béret* à la place de *berret*, car il néglige les différences de prononciation introduites par ces erreurs.

De même que la voyelle [e] s'oppose à la voyelle [ɛ], plus ouverte, de même la voyelle [o] de *pot* a une variante plus ouverte, notée [ɔ] dans l'alphabet phonétique international, et qu'on entend dans *port* [pɔʀ]. On retrouve la même opposition entre le [0] de *feu* et le [œ] de *peur*. L'analogie ne s'arrête pas là. On constate dans de très nombreux mots une hésitation entre [o] et [ɔ] :

<i>connu</i>	[kɔny]	[kɔny]
--------------	--------	--------

Les différents locuteurs utilisent [o], [ɔ] ou un intermédiaire. Il en est de même entre [0] et [œ] :

<i>grenouille</i>	[gr0nuj]	[grœnuj]
<i>peureux</i>	[p0r0]	[pœr0]
<i>dis-le</i>	[dil0]	[dile]

Nous avons fait le même choix pour [o] et [ɔ] que pour [e] et [ɛ] : nous négligeons la différence et nous les représentons par /o/ ; de même, nous représentons généralement [0] et [œ] par /0/ :

<i>connu</i>	/kɔny/
<i>grenouille</i>	/gr0nuj/
<i>peureux</i>	/p0r0z*/
<i>dis-le</i>	/diz* l0/

Toutefois, un cas particulier nous a semblé facile à traiter, celui des mots qui comprennent un [o] dans une position qui est caractéristique d'un [ɔ] obligatoire :

	<i>faute</i>	[fot]	*[fɔt]
à comparer à	<i>sotte</i>	*[sot]	[sɔt],
	<i>motte</i>	*[mot]	[mɔt],
	<i>botte</i>	*[bot]	[bɔt], etc.
	<i>grosse</i>	[gros]	*[grɔs]
à comparer à	<i>bosse</i>	*[bos]	[bɔs],
	<i>gosse</i>	*[gos]	[gɔs], etc.
	<i>dôme</i>	[dom]	*[dɔm]
à comparer à	<i>album</i>	*[albom]	[albɔm],
	<i>comme</i>	*[kom]	[kɔm],
	<i>gomme</i>	*[gom]	[gɔm], etc.

Nous avons marqué ces mots en faisant suivre le /o/ du signe /:/, qui est donc employé comme une marque abstraite de fermeture de /o/ :

<i>faute</i>	/fo:t/
<i>grosse</i>	/gro:s/
<i>dôme</i>	/do:m/

Les mots à [ɔ] obligatoire sont représentés avec un /o/ non suivi de /:/ :

<i>sotte</i>	/sot/
<i>gosse</i>	/gos/
<i>gomme</i>	/gom/

On rencontre aussi des mots qui comportent un [0] dans une position qui est caractéristique d'un [œ] obligatoire :

	<i>feutre</i>	[f0tr]	*[fœtr]
à comparer à	<i>oeuvre</i>	*[0vr]	[œvr],
	<i>peuple</i>	*[p0pl]	[pœpl], etc.
	<i>émeute</i>	[em0t]	?*[emœt]
à comparer à	<i>oeuf</i>	*[0f]	[œf],
	<i>jeune</i>	*[Z0n]	[Zœn], etc.

Dans ce cas, nous avons aussi fait suivre le /0/ de la marque /:/ :

<i>feutre</i>	/f0:tr/
<i>émeute</i>	/em0:t/

Les mots à [œ] obligatoire sont représentés avec un /0/ non suivi de /:/ :

<i>peuple</i>	/p0pl/
<i>jeune</i>	/Z0n/

Le traitement de ce type de variations devra être généralisé dans le futur.

Les voyelles nasales

Dans l'orthographe, les voyelles nasales sont signalées par un *n* : *gant*, *bon*, *bien*. Phonétiquement, ce sont des voyelles qu'on transcrit [ã], [ɔ̃], [ɛ̃] :

<i>gant</i>	[gã],	<i>bon</i>	[bɔ̃],	<i>bien</i>	[bjɛ̃]
-------------	-------	------------	--------	-------------	--------

Pour les formes phonémiques, nous avons opté pour les séquences /an/, /on/, /en/, qui sont proches de l'orthographe :

<i>grandir</i>	/grandir/
<i>bondir</i>	/bondir/
<i>feinte</i>	/fent/

Ces voyelles auraient également pu être représentées par /ã/, /ɔ̃/, /ɛ̃/, ce qui aurait été une solution moins abstraite. Il est difficile de choisir entre ces deux solutions, d'autant que des compromis sont également possibles entre elles. Depuis de nombreuses années, les phonologues ont apporté des arguments en faveur de l'un ou l'autre de ces compromis et solutions, entre autres S.S. Schane (1968), F. Dell (1973) et E.O. Selkirk

(1972). Nous n'avons pas considéré la situation dans toute sa complexité, et nous ne prenons donc pas position dans cette controverse, mais pour élaborer le dictionnaire nous avons fait un choix provisoire et arbitraire en faveur de /an/, /on/, /en/, ce qui permet au moins de rayer de la liste des phonèmes les voyelles nasales [ã], [õ], [ẽ]. Le DELAP ainsi constitué fournit les matériaux pour choisir en connaissance de cause une représentation éventuellement plus simple : en effet, par extraction automatique à partir du DELAP, on pourra obtenir la liste des mots concernés par ce problème de représentation, et les classer en fonction du contexte de la voyelle nasale ou d'autres facteurs phonétiques ou grammaticaux.

Quoi qu'il en soit, le choix provisoire en faveur de /an/, /en/, /on/ a des conséquences immédiates : les voyelles nasales [ã], [õ], [ẽ], ainsi représentées, doivent pouvoir être distinguées des séquences phonétiques [an], [on], [en] qu'on entend dans *panne* [an], *reine* [ɛn], *tonne* [ɔn]. En effet, si ces séquences sont également représentées /an/, /en/ et /on/, cela introduit une confusion entre [ã] et [an], entre [õ] et [on], entre [ẽ] et [ɛn]. Pour éviter ces confusions, nous avons adopté deux solutions différentes suivant que la nasale est située en fin de mot ou non.

- Cas où la nasale n'est pas située en fin de mot. Il s'agit de différencier le [ã] et le [an] de *ânerie* [anri], qui sont tous les deux suivis d'une consonne. Pour cela, nous avons marqué d'une apostrophe /' / les séquences telles que [an]. Ainsi, dans ces mots, [ã] est représenté /an/, et [an] est représenté /an' / :

<i>Henri</i>	[ãri]	/anri/
<i>ânerie</i>	[anri]	/an'ri/

Cette marque /' / correspond généralement à un *e* muet dans l'orthographe :

<i>ânerie</i>	/an'ri/	[anri]	?*[anOri]	?*[anœri]
(il) <i>tannera</i>	/tan'ra/	[tanra]	?*[tanOra]	?*[tanœra]
<i>canneler</i>	/kan'le/	[kanle]	?[kanOle]	?[kanœle]
<i>anneeler</i>	/an'le/	[anle]	?[anOle]	?[anœle]

Dans quelques mots d'origine étrangère, la marque /' / n'a aucun correspondant orthographique :

<i>hanbalisme</i>	[anbalism]	/han'balism/
-------------------	------------	--------------

La même convention permet de différencier le [ẽ] de *teinter* et le [ɛn] de *ancienneté*. Le premier est représenté /en/, le deuxième /en' / :

<i>teinter</i>	[tẽte]	/tente/
<i>ancienneté</i>	[asjɛnte]	/ansien'te/
<i>stencil</i>	[stɛnsil]	/sten'sil/

Enfin, nous avons représenté par /on/ le [õ] de *bonté* et par /on' / le [ɔn] de *bonneteau* :

<i>bonté</i>	[bõte]	/bonte/
<i>bonneteau</i>	[bɔnto]	/bon'to/

- Cas où la nasale est située en fin de mot. Le problème est le même, mais il s'agit de différencier le [ã] de *temps* et le [an] de *panne*. Ce problème de représentation est en

relation avec les effacements de consonnes finales, évoqués dans le chapitre V, section 2, p. 100. C'est pour cette raison que les nasales finales ont été représentées différemment des autres nasales. En partie en raison des conclusions sur les effacements de consonnes finales (cf. chapitre V), nous avons adopté pour [ã] final la séquence de phonèmes /an*/ , et pour [an] final la séquence /an/ :

<i>temps</i>	[tã]	/tan*/
<i>panne</i>	[pan]	/pan/

La marque abstraite /*/, employée en fin de mot, signale donc l'effacement du /n/ et la nasalité de la voyelle /a/. De même, elle permet de différencier le [ĕ] de *gain* et le [ɛ.n] de *gaine*, qui, nous l'avons vu, est représenté /en/ :

<i>gain</i>	[gĕ]	/gen*/
<i>gaine</i>	[gɛ.n]	/gen/

ainsi que le [õ] de *ton* et le [ɔ.n] de *tonne* :

<i>ton</i>	[tõ]	/ton*/
<i>tonne</i>	[tɔ.n]	/ton/

Nous avons étendu ce système de représentation et incorporé dans certains articles du DELAP des séquences finales /in*/ et /yn*/, qui s'interprètent [ĕ] et parfois aussi [œ], bien que cette dernière voyelle nasale puisse toujours être remplacée par [ĕ] en français standard :

<i>fin</i>	/fin*/	[fĕ]	
<i>commun</i>	/komyn*/	[komĕ]	[komœ]

L'utilisation de formes en /in*/ et en /yn*/ simplifie la mise au féminin de certains adjectifs comme *fin* et *commun*. Dans ces mots, la marque /*/ ajoutée aux séquences /in/ et /yn/ signale que cette séquence se prononce [ĕ] et non [in] ou [yn].

Consonnes finales muettes

Pour des raisons qui sont exposées dans le chapitre V, section 2, p. 100, et notamment pour représenter simplement la mise au féminin de nombreux noms et adjectifs, nous avons introduit des consonnes finales muettes dans les représentations phonémiques de certains mots comme *plat*. Pour les distinguer des consonnes finales prononcées, comme celle de *patte*, nous avons fait suivre les consonnes finales muettes de la marque /*/ dont il vient d'être question à propos des voyelles nasales finales :

<i>plat</i>	[pla]	/plat*/
<i>patte</i>	[pat]	/pat/

La marque abstraite /*/ signale donc ici, comme dans le cas des voyelles nasales finales, la présence d'une consonne finale abstraite, qui a été introduite dans les représentations pour simplifier la prise en compte du féminin et d'autres variations grammaticales, mais qui n'est pas prononcée. De nombreux mots sont représentés avec une consonne finale

muette, qui est souvent /t*/ , /n*/ , /z*/ ou /r*/ :

<i>plat</i>	/plat*/	<i>patte</i>	/pat/
<i>bon</i>	/bon*/	<i>tonne</i>	/ton/
<i>creux</i>	/kr0z*/	<i>dose</i>	/doz/
<i>léger</i>	/leZer*/	<i>père</i>	/per/

Dans le cas où la consonne est /n*/ , la voyelle qui précède est prononcée sous la forme d'une voyelle nasale. Pour plus de précisions, voyez le chapitre V et les listes de mots données dans l'annexe 6.

Certains mots comme *tacot* ont été représentés avec une consonne finale muette pour des raisons de pure commodité informatique : /takot*/ , mais pourraient aussi être transcrits sans aucune consonne finale : /tako/. Leur représentation gagnerait ainsi en simplicité.

E dits muets

On ne sait pas comment représenter au mieux les *e* "muets" dans un système phonémique. Plusieurs solutions ont été proposées, mais toutes sont controversées. Pourtant, un examen systématique et exhaustif des mots qui comportent un *e* muet devrait permettre d'observer et de recenser l'ensemble des faits à prendre en compte, et c'est à partir d'une telle description qu'une représentation pourra être choisie.

Pour l'instant, seules quelques conclusions se dégagent. Phonétiquement, lorsque le *e* orthographique intérieur sans accent se prononce, comme dans *premier* ou dans *une fenêtre*, on ne sait pas le distinguer du [0] de *peuplier* ni du [œ] de *meurtrier* d'une façon reproductible. Par ailleurs, dans certains mots comme *premier*, il se prononce obligatoirement ; dans d'autres, comme *semer*, il est facultatif ; dans les autres, comme *ancienneté*, il semble ne jamais se prononcer en français standard. Quant au *e* dit "muet" final, il ne se prononce effectivement pas : ainsi, *mer* et *mère* sont homonymes. C'est pourquoi, dans les mots qui comportent un *e* muet final orthographique, comme *mère*, nous n'avons pas introduit de phonème correspondant :

<i>mère</i>	/mer/
<i>mer</i>	/mer/

Pour les "e muets" situés à l'intérieur des mots, nous les avons généralement représentés par un phonème spécial. Le symbole traditionnellement utilisé est le schwa /ə/, mais il n'est pas représenté dans le code ASCII, c'est pourquoi nous lui avons substitué le guillemet /"/ :

<i>projeter</i>	/proZ"te/
<i>dégeler</i>	/deZ"le/
<i>redire</i>	/r"dir/

Toutefois, nous avons donné une autre représentation à certains *e* dits muets. Certains

mots comportent un *e* muet intérieur, et, dans la syllabe qui précède, l'une des voyelles ouvertes [ɛ], [œ], [ɔ] :

<i>saignement</i>	[sɛnj0ma]	?*[sɛnj0ma]
<i>peuplement</i>	[pœpl0ma]	?*[p0p10ma]
<i>(il) cognera</i>	[kɔnj0ra]	*[konj0ra]

Dans ces mots, [ɛ], [œ], [ɔ] semblent obligatoires : on peut difficilement les remplacer par les voyelles plus fermées [e], [o], [ɔ]. Ces [ɛ], [œ], [ɔ] obligatoires sont remarquables : en-dehors de ce cas, on ne trouve de [ɛ], [œ], [ɔ] obligatoires en français qu'en syllabe finale, devant certains groupes de consonnes particuliers comme [rt], et dans quelques mots isolés. En particulier, si on remplace le *e* muet par une autre voyelle dans les exemples ci-dessus, la prononciation [ɛ], [œ], [ɔ] de la voyelle précédente n'est plus obligatoire :

<i>saigner</i>	[sɛnje]	[senje]
<i>peupler</i>	[pœple]	[p0ple]
<i>cogner</i>	[kɔnje]	[konje]

Le fait qu'on observe des [ɛ], [œ], [ɔ] obligatoires dans *saignement*, *peuplement*, *il cognera* est en relation avec la présence du *e* muet intérieur dans la syllabe qui suit. Dans ce cas, nous avons transcrit ce *e* muet par un autre symbole que /"/, pour indiquer la présence d'un [ɛ], [œ] ou [ɔ] obligatoire. Nous avons utilisé pour cela l'apostrophe /'/' :

<i>saignement</i>	/sɛN'mant*'/
<i>peuplement</i>	/p0pl'mant*'/
<i>(il) cognera</i>	/kɔN'ra/

Lorsque la voyelle de la syllabe précédente n'est ni /e/, ni /o/, ni /ɔ/, nous avons choisi arbitrairement entre /"/ et /'/' pour représenter le *e* muet :

<i>pelouse</i>	/p"luz/
<i>mouvement</i>	/muv'mant*'/

Il est à noter que le phonème /'/' est utilisé dans deux situations différentes :

- dans la situation que nous venons d'évoquer, c'est-à-dire en présence d'un [ɛ], [œ], [ɔ] obligatoire (*saignement*, *peuplement*, *il cognera*) ;
- dans la situation évoquée plus haut à propos des voyelles nasales, à savoir lorsque l'une des séquences [an], [ɛn], [ɔn] est suivie d'une consonne (*ânerie*, *ancienneté*, *bonneteau*) et doit être différenciée des voyelles nasales [ā], [Ț], [Ȯ] (*Henri*, *teinter*, *bonté*).

Dans certains mots, ces deux situations sont réunies à la fois. C'est le cas de *ancienneté* et de *bonneteau*. Dans ces mots, on a d'une part un [ɛ] ou un [ɔ] obligatoire, et d'autre part une séquence [ɛn] ou [ɔn] suivie d'une consonne. Ces mots sont représentés /ansien'te/ et /bon'to/. Le phonème abstrait /'/' remplit alors deux fonctions : d'une part, il représente l'ouverture obligatoire du /e/ et du /o/ précédent ; d'autre part, il indique que la séquence phonémique /en/ ou /on/ représente non pas une voyelle nasale, comme dans *teinter* ou *bonté*, mais une séquence *Voy*-[n]. Ce double usage ne crée donc pas d'ambiguïté.

Consonnes doubles

Les consonnes prononcées doubles sont plus rares que les consonnes doubles de l'orthographe. Dans la plupart des mots qui comportent une consonne double orthographique, on prononce obligatoirement une consonne simple :

appel [apɛl] *[appɛl]

Dans d'autres mots, on constate une variation phonétique entre deux variantes, l'une avec consonne simple, l'autre avec une consonne double, cette dernière étant souvent d'un niveau plus recherché :

allusion [alyzjɔ̃] [allyzjɔ̃]
immédiat [imedja] [immedja]

Dans quelques mots, on prononce obligatoirement une consonne double :

transsaharien [trāsaaɾjɛ] *[trāsaarjɛ]
(il) flairera [flɛɾra] *[flɛɾra]

Le type *appel*, prononcé avec une consonne simple obligatoire, et le type *transsaharien*, prononcé avec une consonne double obligatoire, n'ont pas posé de problèmes de représentation. Le premier a été représenté avec une consonne simple et l'autre avec deux consonnes, éventuellement séparées par le phonème /' / :

transsaharien /transsaarien* /
(il) flairera /flɛr'ra /

Le type *allusion*, qui présente une variation phonétique, demanderait une étude particulière : il sera utile de connaître la liste des mots concernés. Nous n'avons pas abordé ce travail. Pour l'instant, le DELAP confond donc les types *appel* et *allusion* en les représentant tous les deux avec des consonnes simples :

appel /apɛl /
allusion /alyzion* /
immédiat /imediāt* /

Les semi-consonnes [j], [w], [ɥ]

Les sections 3 à 6 du chapitre V sont consacrées à ces semi-consonnes. Nous y renvoyons le lecteur. Nous ne donnons ici que quelques indications qui aideront à la lecture des formes phonémiques. Les semi-consonnes apparaissent dans de nombreux mots comme *pied* [pje], *pois* [pwa] et *puis* [pyi]. Dans ces mots, nous les avons représentées par les voyelles correspondantes /i/ de *pire*, /u/ de *pour* et /y/ de *pur* :

pied /pie /
pois /pua /
puis /pyi /

Ainsi, nous n'avons utilisé le phonème /w/ que dans quelques onomatopées et quelques mots d'origine étrangère comme *out* /awt / ; le phonème /ɥ/ n'a été utilisé pour aucun

mot. En revanche, le phonème /j/ a été fréquemment utilisé après voyelle :

<i>paille</i>	/paj/
<i>bail</i>	/baj/

Nous avons introduit la marque abstraite /+ / dans certains mots comme *jouable*, pour les différencier de mots comme *joignable* :

<i>jouable</i>	/Zu + abl/
<i>joignable</i>	/ZuaNabl/

Les raisons théoriques de ces choix sont exposées en détail dans le chapitre V.

Le phonème /N/

La séquence *gn* de l'orthographe correspond généralement à la séquence phonétique [nj] :

<i>gagner</i>	[ganje]
<i>saignement</i>	[senj0ma]

Toutefois, dans certains contextes, *gn* a une prononciation légèrement différente de [nj], et que les phonéticiens transcrivent par [ñ]. Comme le symbole ñ n'est pas représenté dans le code ASCII strict, nous l'avons remplacé par *N* :

<i>un signe de la main</i>	[ĚsiNd0lamĚ]
<i>il gagnera</i>	[iganj0ra] [igaNra]
<i>baignoire</i>	[beNwar]

Dans les représentations phonémiques de ces mots, on transcrit traditionnellement [nj] et [N] de la même façon, généralement par le phonème /N/. Dans le DELAP, les mots cités ci-dessus sont représentés conformément à cet usage :

<i>gagner</i>	/gaNe/
<i>saignement</i>	/seN'mant* /
<i>signe</i>	/siN/
<i>(il) gagnera</i>	/gaN'ra/
<i>baignoire</i>	/beNuar/

Le *h* aspiré

Des applications informatiques importantes, comme la production automatique de textes phonétiques, nécessitent un traitement des interactions phonétiques entre les mots, et notamment des liaisons en français. Le *h* aspiré est un des paramètres à prendre en compte en ce qui concerne les liaisons. Certains mots orthographiés avec un *h* initial interdisent la liaison avec le mot qui précède, d'autres non :

<i>un hérissément</i>	[Ěrisma]	?*[Ěnerisma]
<i>un héritage</i>	*[ĚritaZ]	[ĚneritaZ]

On dit des premiers qu'ils ont un *h* aspiré. Quelques mots orthographiés sans *h* interdisent également la liaison avec le mot qui précède :

<i>les onze autres</i>	[leʒotr]	*[lezotr]
<i>un ululement</i>	[Ëlylma]	?*[Ënylma]

Traditionnellement, dans les représentations phonémiques, on attribue aux mots à *h* aspiré un phonème initial qui est souvent /h/ :

<i>hérissément</i>	/heris'mant*/
<i>héritage</i>	/eritaZ/
<i>onze</i>	/honz/

Nous avons employé cette convention dans la représentation phonémique du DELAP. G. Tep (1979) emploie la convention inverse. Il attribue un phonème initial /x/ aux mots qui admettent la liaison à gauche :

<i>hérissément</i>	/eris'mant*/
<i>héritage</i>	/xeritaZ/
<i>onze</i>	/onz/

Les deux conventions sont équivalentes en ce sens que l'on peut passer automatiquement de l'une à l'autre. La convention traditionnelle que nous avons reprise est plus économique, car parmi les mots dont l'initiale phonétique est une voyelle, ceux qui admettent la liaison à gauche sont beaucoup plus nombreux que ceux qui l'interdisent.

8. Le calcul des transcriptions phonétiques à partir des transcriptions phonémiques

Par définition, les transcriptions phonémiques sont des représentations abstraites des transcriptions phonétiques. Elles sont simplifiées, mais moins redondantes et plus informatives. Il est donc possible de générer automatiquement celles-ci à partir de celles-là. Un algorithme a donc été réalisé pour analyser les transcriptions phonémiques du DELAP et pour calculer les transcriptions phonétiques correspondantes (E. Laporte, 1984). Cet algorithme est indispensable pour les applications informatiques, aussi bien pour la synthèse et pour la reconnaissance de la parole que pour la correction orthographique. En effet, les représentations phonémiques utilisées lors des traitements ne correspondent pas directement aux prononciations observables, mais servent à les représenter et à les retrouver automatiquement.

Cet algorithme, appliqué au DELAP, produit des transcriptions phonétiques des 64.000 entrées. Il tient compte des conventions qui sont utilisées dans les représentations phonémiques et que nous avons énumérées. Cet ensemble de conventions et cet algorithme ne peuvent donc pas être modifiés indépendamment l'un de l'autre. En particulier, si la représentation phonémique prend en compte des variations phonétiques plus nombreuses qu'elle ne le fait dans son état actuel, l'algorithme pourra produire toutes les variantes des mots concernés.

L'algorithme est structuré en 6 passes qui sont appliquées une seule fois et dans un ordre déterminé. Il produit des transcriptions phonétiques de mots isolés. Il pourra être étendu à certaines des variations phonétiques qui ont lieu systématiquement à la jonction de deux mots, par exemple dans la séquence *quatre lits*, qui peut se prononcer [katli], [katrɔli] ou [katrɔli].

9. Les groupes de consonnes phonétiques

Le DELAP a fourni des matériaux pour l'étude des groupes de consonnes phonétiques en français. On peut distinguer trois sortes de groupes de consonnes :

- les groupes de consonnes initiaux, comme [kr] dans *cri* ;
- les groupes de consonnes intérieurs ou médians, comme [rkl] dans *encercler* ;
- les groupes de consonnes finaux, comme [kst] dans *mixte*.

Les dissemblances entre l'orthographe et la prononciation font qu'un même groupe de consonnes phonétiques s'écrit souvent de différentes façons. C'est le cas de [rk] qui correspond à trois groupes de consonnes orthographiques différents dans *marque*, *mark* et *arc* : *rq*, *rk* et *rc*. L'inverse se produit aussi : plusieurs groupes de consonnes phonétiquement différents s'écrivent par un même groupe de consonnes orthographiques. Ainsi, [rs] et [rk] s'écrivent *rc* dans *farceur* et *arcade*. C'est donc sur des listes de formes phonétiques ou phonémiques, et non de formes orthographiques, que les groupes de consonnes peuvent être recensés.

Les exemples précédents sont des groupes de deux ou trois consonnes phonétiques. On peut leur ajouter des exemples de consonnes simples :

- Dans *tu* et dans *phare*, une consonne initiale [t] ou [f] est suivie d'une voyelle.
- Dans *ami* et dans *appel*, les consonnes [m] et [p] sont intervocaliques, c'est-à-dire situées entre deux voyelles du même mot. Comme dans les exemples précédents, ces consonnes constituent la limite entre deux syllabes.
- Dans *ail* et dans *balle*, les consonnes finales [j] et [l] sont précédées d'une voyelle.

Par convention, nous considérons également ces consonnes simples comme des "groupes de consonnes" qui ne comportent qu'une consonne.

Des recensements de groupes de consonnes phonétiques ont été faits à partir du *Dictionnaire du français fondamental* (G. Gougenheim, 1964) dans le cadre d'une étude sur la syllabe en vue de la reconnaissance de la parole (P. Juban, 1974). Des données statistiques sur la structure des syllabes ont ainsi été réunies. Il serait utile de disposer de données de ce type sur un dictionnaire plus étendu.

C'est pourquoi nous avons réalisé un premier programme pour recenser les groupes de consonnes phonétiques intérieurs, comme [ksk] dans *excuser*. Un tel groupe de consonnes constitue la limite entre deux syllabes. Nous avons considéré comme un cas particulier de groupes de consonnes, outre le cas des consonnes simples, celui des hiatus, c'est-à-dire de la succession entre deux voyelles, comme dans *boa*. En effet, les hiatus constituent également des limites entre deux syllabes. Ils ont été inclus dans le

traitement au titre de groupes de 0 consonnes. Pour chaque groupe de consonnes, le programme recense les mots qui comportent ce groupe. Les listes obtenues à partir de la version 3 du dictionnaire DELAP (50.000 entrées) comportent plus de 100.000 éléments. Ces éléments ont été classés dans un ordre alphabétique qui tient compte prioritairement des phonèmes qui précèdent le groupe. Nous donnons dans l'annexe 4 deux pages qui correspondent aux groupes de consonnes [ksf], [ksk], [kskl], [kskr], [ksm] et [ksp]. Rappelons que les chaînes recensées sont des groupes de consonnes phonétiques, comme [ksk], et non des groupes de consonnes orthographiques comme xc. Les entrées sont présentées sur 4 colonnes. La première colonne comporte la forme orthographique de chaque mot. Les trois autres colonnes constituent sa forme phonémique, découpée en trois parties dont la deuxième est le groupe de consonnes. Les voyelles nasales non finales, qui sont représentées /an/, /en/, /on/, /in/ et /yn/ dans le DELAP, sont considérées comme les autres voyelles : ainsi, malgré le /n/ qui apparaît dans l'entrée *grandir* /grandir/, ce mot fournit un exemple de la consonne simple /d/, et non du groupe de deux consonnes /nd/. La semi-consonne [j] de *amitié* [amitje], qui est représentée par un /i/ dans la forme phonémique /amitie/, est traitée comme une voyelle : ce mot fournit donc des exemples de la consonne simple /t/ et d'un hiatus entre /i/ et /e/. Cette convention permet d'obtenir la liste des mots représentés avec une séquence /i/-Voy : *papier, déplier, relier*, etc. Au vu de cette liste, nous avons pu classer ces mots suivant leurs variations phonétiques. Les conclusions de cette étude sont présentées en détail dans le chapitre V.

Le programme de recensement des groupes de consonnes phonétiques constitue un exemple des méthodes d'interrogation automatique du DELAP évoquées dans le chapitre II, section 6, p. 42. Il comporte sept modules. Le module principal est un fichier de commandes VMS. Il appelle les six autres, qui sont des fichiers de commandes EDT. Les 26 parties du DELAP, qui correspondent aux 26 lettres de l'alphabet, sont traitées les unes après les autres.

Pour chacune des 26 parties, la première opération consiste à séparer les trois zones du DELAP. Trois fichiers temporaires sont ainsi créés : *tempor.or* qui contient la zone orthographique de chaque entrée, *tempor.ph* qui contient la zone phonémique, et *tempor.fl* qui contient la zone flexionnelle. Cette séparation permet de ne rechercher les groupes de consonnes que dans la zone phonémique. Elle est effectuée par trois fichiers de commandes EDT, qui créent chacun un des trois fichiers temporaires. Parmi ceux-ci, le fichier *tempor.ph* contient les formes phonémiques des mots. Le module qui crée ce fichier effectue également le marquage des groupes de consonnes intérieurs, y compris, comme nous l'avons vu, les consonnes intervocaliques (*appel*) et les hiatus (*boa*). Pour effectuer ce marquage, sachant que les voyelles sont moins nombreuses que les consonnes, nous avons tiré parti du fait que tout groupe de consonnes intérieur, par définition, est limité à gauche et à droite par des voyelles. Des commandes de l'éditeur insèrent un caractère spécial, \$, de part et d'autre de chaque voyelle. Pour l'entrée *approbation*, ce marquage effectué sur /aprobasion*/ donne le résultat /\$a\$pr\$oS\$b\$a\$ss\$i\$sson\$*/. Le premier et le dernier caractère \$ de chaque entrée est ensuite effacé :

/a\$pr\$oS\$b\$a\$ss\$i\$sson\$*/

Dans la chaîne de caractères obtenue, les caractères \$ isolent les groupes de consonnes /pr/, /b/ et /s/. Le hiatus entre /i/ et /on*/ est signalé par la séquence \$\$\$. Le même traitement, sur l'entrée *sonir*, donne le résultat /so\$rt\$ir/. A la suite de ce marquage, les trois zones de chaque entrée peuvent être réassemblées à l'aide de commandes VMS. Le fichier produit est appelé *tempor.3m* ; les fichiers *tempor.or*, *tempor.ph* et *tempor.fl* sont détruits.

La seconde étape, effectuée par deux modules spécifiques, produit une liste dans laquelle chaque ligne est consacrée à un groupe de consonnes rencontré dans un mot particulier. Pour cela, les entrées qui ne comportent pas de groupes de consonnes intérieurs, comme *col*, sont effacées ; celles qui en comportent un seul, comme *panir*, restent inchangées ; celles qui en comportent plusieurs, comme *approbation*, sont multipliées, et chacune des lignes constituées correspond à un des groupes de consonnes intérieurs du mot :

```
/a$pr$obasion*/
/apro$b$asion*/
/aproba$$s$ion*/
/aprobasi$$on*/
```

L'algorithme suivant, qui est adapté d'un algorithme conçu par Ch. Leclère au LADL, permet d'arriver à ce résultat à partir du fichier *tempor.3m*, dans lequel les groupes de consonnes intérieurs sont marqués par des caractères \$:

1. Dans chaque ligne du fichier *tempor.3m*, remplacer les deux premiers caractères \$, s'ils existent, par des caractères % : /a%pr%o\$b\$a\$\$s\$i\$\$on*/.
2. Extraire dans un fichier *tempor.m1* les lignes qui comportent des caractères %.
3. Effacer les caractères % du fichier *tempor.3m* : /apro\$b\$a\$\$s\$i\$\$on*/.
4. Itérer douze fois les étapes 1, 2 et 3 en créant les fichiers *tempor.m2...*, *tempor.m12*. Par exemple, la ligne /apro%b%a\$a\$\$s\$i\$\$on*/ figure dans le fichier *tempor.m2*. Aucun mot du DELAP ne comporte plus de douze groupes de consonnes intérieurs.
5. Fusionner les fichiers *tempor.m1...*, *tempor.m12* en un fichier *tempor.m*.
6. Dans le fichier *tempor.m*, effacer les caractères \$.

Cet algorithme a été implanté sous forme de deux fichiers de commandes EDT et quelques commandes VMS. Le résultat du calcul est un fichier dans lequel chaque ligne est consacrée à un groupe de consonnes rencontré dans un mot particulier. Ce groupe est marqué par deux caractères % qui indiquent les limites gauche et droite du groupe :

```
/apro%b%asion*/
```

La troisième étape est réalisée par le système de tabulation Lexsyn conçu par Ph. Vasseux au LADL. Les 100.000 lignes issues des 26 fichiers sont classées dans un

ordre alphabétique qui tient compte :

- en premier lieu, du groupe de consonnes marqué dans la ligne ;
- en second lieu, du début de la forme phonémique, c'est-à-dire /apro/ dans l'exemple ci-dessus ;
- en troisième lieu, de la fin de la forme phonémique.

Les lignes sont ensuite formatées de manière à constituer les quatre colonnes. Le fichier obtenu est alors prêt à imprimer.

La liste construite suivant ce processus est un recensement des groupes de consonnes phonétiques intérieurs. Elle présente l'avantage de comporter, pour chaque groupe de consonnes, un grand nombre d'exemples, ce qui permet de prendre en compte les cas exceptionnels ou difficiles, et aussi de détecter dans le DELAP des erreurs isolées ou systématiques. Son inconvénient majeur est qu'elle est peu maniable : sur papier, elle occupe plus de 2.000 pages. Par ailleurs, cette liste ne permet de calculer automatiquement ni le nombre de groupes de consonnes intérieurs distincts, ni le nombre d'exemples d'un groupe de consonnes donné.

Nous avons donc complété cette liste par des dénombrements. Par un programme écrit en C, nous avons recensé :

- 85 groupes de consonnes initiaux distincts, dont 54 de 2 consonnes,
- 550 groupes de consonnes intérieurs, dont 273 de 2 consonnes,
- et 138 groupes de consonnes finaux, dont 83 de 2 consonnes.

Ce recensement inclut les consonnes simples, considérées comme des groupes réduits à une seule consonne, mais pas les hiatus. Pour chaque groupe de consonnes, les exemples correspondants dans le DELAP ont été dénombrés séparément. Le résultat de ces comptages est un tableau qui donne les effectifs des différents groupes de consonnes recensés. Les groupes de deux consonnes, qui sont les plus nombreux, sont illustrés par quelques exemples chacun : 4 par groupe au maximum. Le tableau et les exemples constituent l'annexe 5.

Etant donné la complexité de ces comptages, ils ont été faits par un programme écrit en C. Lors d'une première passe, les groupes de consonnes sont analysés et marqués. Lors d'une deuxième passe, ils sont dénombrés et classés suivant la séquence de consonnes qui constituent le groupe et suivant sa position : initiale, intérieure ou finale.

Ces comptages et ces exemples permettent de détecter des erreurs. Par exemple, le groupe intérieur [tf] est donné avec un exemple unique, l'adjectif *motphochronologique*, qui résulte d'une erreur de saisie. Ils mettent en évidence les limites du phonétiseur : celui-ci ne tient pas compte des assimilations à l'intérieur des groupes de consonnes, par exemple dans *absolu* [apsoly] qui est actuellement transcrit /absoly/. En effet, c'est au vu des listes d'exemples tels que *absolu* qu'un traitement adéquat pourra être adopté. Pour l'instant, 79 occurrences du groupe /bs/ intérieur ont été dénombrées dans le DELAP. L'examen de ces exemples dans la liste des groupes de consonnes intérieurs montre que dans tous ces mots, on prononce [ps] et non [bs]. A

l'aide de ces listes d'exemples substantielles, une étude systématique sur les autres assimilations consonantiques pourra également être menée.

Chapitre IV. La flexion : conjugaisons, pluriels, féminins

Pour achever d'analyser la structure d'un article du DELAP, il reste à décrire le contenu de la troisième zone, qui comporte les informations grammaticales et flexionnelles sur les mots. Ces informations sont directement issues de la zone flexionnelle du DELAS, qui a été évoquée dans le chapitre II, section 2, p. 27. Rappelons que les mots peuvent être classés en cinq catégories grammaticales, dites aussi parties du discours :

- (1) les verbes,
- (2) les noms,
- (3) les adjectifs et participes,
- (4) les adverbes, invariables,
- (5) les parties résiduelles : conjonctions, prépositions, etc.

Les catégories (4) et (5) sont traitées comme invariables. Dans le DELAP comme dans le DELAS, les mots correspondants reçoivent une description grammaticale simple :

- .ADVE pour les adverbes,
- .PREP pour les prépositions,
- .CONS pour les conjonctions de subordination,
- .CONC pour les conjonctions de coordination,
- .DETE pour les déterminants,
- .PRON pour les pronoms,
- .QUAN pour les quantificateurs.

Les catégories (1), (2) et (3) rassemblent les mots variables : les verbes se conjuguent, les noms sont variables en nombre et parfois en genre, les adjectifs sont variables en genre et en nombre.

1. Les noms et les adjectifs

Dans le DELAS, la zone flexionnelle spécifie les variations orthographiques des noms et des adjectifs lors de la mise au pluriel et de la mise au féminin. Environ 80 schémas de variation distincts ont été recensés et numérotés de 0 à 80. La zone flexionnelle d'un nom comporte un code de la forme .N*n*, où *n* est le numéro du schéma de variation :

col,.N1
vitrier,.N42

De même, la zone flexionnelle d'un adjectif comporte un code de la forme .A*n* :

léger,.A42
seul,.A32

A titre d'exemple, le schéma de variation 42, qu'on rencontre dans de nombreux noms et

adjectifs, correspond aux propriétés suivantes :

- Le mot admet les deux genres, masculin et féminin. Au singulier, le masculin et le féminin sont en *-er* et en *-ère*, et on passe de l'un à l'autre en intervertissant ces finales.
- Le mot admet les deux nombres, singulier et pluriel. Au masculin comme au féminin, le pluriel s'obtient en ajoutant *s* au singulier.

Des programmes réalisés par B. Courtois effectuent la mise au pluriel et la mise au féminin des noms et des adjectifs, compte tenu de ces schémas de variation.

Pour plus de commodité dans le contrôle de la correspondance entre le DELAS et le DELAP, les informations flexionnelles du DELAS ont été reproduites intégralement dans le DELAP. Toutefois, les variations flexionnelles d'un mot sont différentes suivant qu'on les observe sur l'orthographe, sur des transcriptions phonétiques ou sur des transcriptions phonémiques. Ainsi, la mise au pluriel orthographique et la mise au pluriel phonétique ne concordent généralement pas. La plupart des noms et des adjectifs font leur pluriel orthographique en *s*, mais le pluriel se prononce comme le singulier, donc la mise au pluriel phonétique de ces mots se fait sans changement. On constate aussi des disparités entre la mise au féminin orthographique et la mise au féminin phonétique. Considérons les noms et les adjectifs des classes N42 et A42 : ils ont la même mise au féminin orthographique. Quant à la mise au féminin phonétique, elle se fait généralement, dans ces mots, en remplaçant la finale [e] par la finale [ɛr] :

<i>léger</i>	[leZe]	<i>légère</i>	[leZɛr]
--------------	--------	---------------	---------

Cependant, dans les adjectifs *amer*, *cher* et *fier*, qui ont les mêmes variations orthographiques que *léger*, le masculin et le féminin se prononcent de la même façon :

<i>amer</i>	[amɛr]	<i>amère</i>	[amɛr]
-------------	--------	--------------	--------

Pour permettre la flexion automatique des noms et des adjectifs sous leur forme phonémique, nous avons donc dû construire un système parallèle au système de flexion orthographique. Les noms et les adjectifs ont été classés suivant leurs schémas de variation phonémiques. La classification obtenue est en relation complexe avec la classification correspondante établie pour l'orthographe et utilisée dans le DELAS. De même que les schémas de variation orthographiques, les schémas de variation phonémiques ont été recensés, numérotés et incorporés au DELAP.

Cette numérotation et cette correspondance sont donnés dans le tableau ci-après. Chaque ligne correspond à un cas différent, illustré par un exemple donné dans la colonne la plus à droite. Chacun de ces cas est défini par un schéma de variation orthographique et par un schéma de variation phonémique. Les deux colonnes numériques de gauche permettent de repérer ces schémas de variation par leurs numéros. La première colonne donne le schéma de variation orthographique par son numéro, compris entre 0 et 80. La deuxième colonne numérique donne le schéma de variation phonémique, également par un numéro compris entre 0 et 80. La zone grammaticale et flexionnelle du DELAP prend en compte les deux classifications et permet de retrouver les deux schémas de variation : le premier est donné par son

numéro de 0 à 80, et le deuxième est donné ensuite entre parenthèses, sauf lorsqu'il se déduit du premier sans ambiguïté. Il est à noter que pour la plupart des mots, le schéma de variation phonémique se déduit sans ambiguïté du schéma de variation orthographique, et les informations du DELAS sont donc reproduites telles quelles dans le DELAP.

Les colonnes du milieu donnent les finales des différentes formes du mot, exprimées dans la même représentation phonémique que le DELAP. Le signe . indique une forme à laquelle ne vient s'ajouter aucune terminaison. Le signe - indique une forme inexistante, par exemple le pluriel du nom *bercail*. Le signe & indique que les formes fléchies constituent des entrées séparées du DELAP, solution qui a été adoptée dans des cas exceptionnels.

Noms et adjectifs masculins

Schéma de variation orth.	phon.	Terminaison		Exemples
		sing.	plur.	
0S, 0P	0S, 0P	&	&	ledit, lesdits ; ail, aux ; lieudit, liexdits ¹
1	1	.	.	col, cols
1	2	.	.	oeuf, oeufs ; boeuf, boeufs
2	1	.	.	as, as
2	2	.	.	os, os
2S	1S	.	-	bercail
2P	1P	-	.	abats
3	1	.	.	pieu, pieux
4	4	al	o	cheval, chevaux
4	17	el	0	ciel, cieux
5	5	aj	o	travail, travaux
6	6	.	.	amour(s) ; délice(s) ; orgue(s) ²
7	7	ys	i	naevus, naevi
8	8	om	a	quantum, quanta
9	9	om	zom	bonhomme, bonshommes
9	12	jom	zom	gentilhomme, gentilshommes
10	10	an	en	barman, barmen
11	11	.	z	lobby, lobbies
11	1	.	.	lobby, lobbies
12	1	.	.	lunch, lunches
13	13	o	i	tempo, tempi
14	14	.	im	kibboutz, kibboutzim
15	15	.	m	sefardi, sefardim
15	16	j	im	goï, goïm

¹ Dans cette classe, chaque forme fléchie est donnée dans un article du DELAP et du DELAS.

² Mots qui sont masculins au singulier, mais des deux genres au pluriel.

Noms et adjectifs féminins

Schéma de variation		Terminaison		Exemples
orth.	phon.	sing.	plur.	
20S, 20P	20S, 20P	&	&	ladite, lesdites ; madame, mesdames ³
21	21	.	.	vie, vies
22	21	.	.	fois, fois
22S	21S	.	.	basoche
22P	21P	.	.	fiançailles
23	21	.	.	eau, eaux
24	24	.	z	lady, ladies
24	21	.	.	lady, ladies
25	25	an	en	tenniswoman, tenniswomen

Noms et adjectifs masculins et féminins dont le pluriel est phonétiquement semblable au singulier

Schéma de variation		Terminaison		Exemples
orth.	phon.	masc.	fém.	
30M, 30F	30M, 30F	&	&	docteur, doctoresse ⁴
31	31	.	.	artiste, artiste
32	31	.	.	ami, amie
32	32	*	.	lent, lente
32	43	.	kt	suspect, suspecte
32	44	*	kt	distinct, distincte
32P	31P	.	.	quadruplés, quadruplées
33	32	*	.	andalou, andalouse
34	32	*	.	sot, sotté
34	31	.	.	net, nette
66	33	0 (ej)	ej	vieux (vieil), vieille
72	34	o	el	jumeau(x), jumelle(s)
35	35	r	z	diseur, diseuse
36	36	0r	ris	acteur, actrice
37	37	0r	"res	demandeur, demanderesse
38	38	f	v	juif, juive
38	32	*	.	serf, serve
39	39	.	es	maire, mairesse
39	32	*	.	bêta, bêtas
40	31	.	.	nul, nulle
40	32	*	.	gentil, gentille
73	40	o (el)	el	beau(x) (bel), belle(s)
41	32	*	.	bon, bonne
42	32	*	.	léger, légère
42	31	*	.	fier, fière
43	32	*	.	secret, secrète

³ Dans cette classe, chaque forme fléchie est donnée dans un article du DELAP et du DELAS.

⁴ Dans cette classe, chaque forme fléchie est donnée dans un article du DELAP et du DELAS.

44	38	f	v	bref, brève
45	45	k	S	sec, sèche
46	32	*	.	franc, franque
47	32	*	.	franc, franche
48	48	k	Ses	duc, duchesse
49	32	*	.	long, longue
50	31	.	.	aigu, aiguë
51	51	.	esk	maure, mauresque
52	52	n*	N	malin, maligne
53	53	u	ol	fou (fol), folle
54	35	r	z	challenger, challengeuse
54	54	er	Oz	challenger, challengeuse
55	55	.	in	tsar, tsarine
55	56	Or	"rin	speaker, speakerine
56	37	Or	"res	quaker, quakeresse
56	39	.	es	clown, clownesse
57	57	o	a	flamenco, flamenca
60M, 60F	30M, 30F	&	&	tiers, tierce ⁵
61	32	*	.	danois, danoise
61	31	.	.	ours, ourse
62	32	*	.	bas, basse
62	31	.	.	métis, métisse
63	32	*	.	creux, creuse
64	32	*	.	faux, fausse
65	32	*	.	doux, douce
66	33	0 (ej)	ej	vieux (vieil), vieille
67	32	*	.	profès, professe
70M, 70F	30M, 30F	&	&	hébreu(x), hébraïque(s) ⁶
72	34	o	el	jumeau(x), jumelle(s)
73	40	o (el)	el	beau(x) (bel), belle(s)
80	31	.	.	orange, orange

Noms et adjectifs masculins et féminins dont le pluriel est phonétiquement différent du singulier

Schéma de variation orth.	phon.	Terminaison				Exemples
		ms	fs	mp	fp	
71	71	0l	0l	0	0l	aïeul, aïeule(s), aïeux
75	75	t*	t	s	t	tout, toute(s), tous
76	76	al	al	o	al	animal, animale(s), animaux
78	78	om	a	a	a	minimum, minima, minima
80P	31P	-	-	.	.	vingt, vingt

⁵ Dans cette classe, chaque forme fléchie est donnée dans un article du DELAP et du DELAS.

⁶ Dans cette classe, chaque forme fléchie est donnée dans un article du DELAP et du DELAS.

Des programmes écrits en langage C effectuent la flexion automatique des noms et des adjectifs, compte tenu de cette classification flexionnelle. Ils effectuent parallèlement la flexion orthographique et la flexion phonémique, et donnent chaque forme sous ses deux représentations.

2. La conjugaison des verbes

Dans le cas d'un verbe, la zone flexionnelle du DELAS et du DELAP spécifie comment le verbe se conjugue. Cette zone comporte un code de la forme .V*n*, où *n* est un numéro de conjugaison compris entre 0 et 99 :

monter, monte., V3

faire, fer., V80

Les verbes défectifs, c'est-à-dire ceux dont certaines formes n'existent pas, sont signalés par la lettre D placée après le numéro de conjugaison :

déchoir, deSuar., V53D

Les verbes dont le participe passé est invariable dans tous leurs emplois sont signalés par la lettre U placée après le numéro de conjugaison. Cette information est issue du lexique-grammaire des verbes simples (Maurice Gross, 1975 ; J.P. Boons, A. Guillet et Ch. Leclère, 1976) :

suffire, syfir., V89U

Toutes ces informations flexionnelles concernent la conjugaison orthographique des verbes. Elles sont exploitées par les programmes de conjugaison automatique réalisés par B. Courtois et Y. Nicolas. Pour des raisons de commodité, elles ont été reproduites intégralement dans le DELAP. Toutefois, de même que pour les noms et les adjectifs, la conjugaison est différente suivant le mode de représentation : orthographique, phonémique ou phonétique. C'est pourquoi, pour permettre la conjugaison automatique des verbes dans leur représentation phonémique, nous avons élaboré un système de conjugaison phonémique, parallèle au système de conjugaison orthographique et construit sur le même modèle.

La conjugaison des verbes, de même que la flexion en général, est liée à l'interaction de deux variables : d'une part, le verbe employé, choisi parmi l'ensemble des verbes français, et d'autre part les combinaisons de traits flexionnels (temps, personne, genre, nombre). La première variable, qui est une donnée lexicale, est indépendante de la deuxième. Ainsi, décrire la conjugaison des verbes consiste à expliciter leurs variations de forme en fonction (1) du verbe employé, et (2) des traits flexionnels. Etant donné un verbe, par exemple *prendre*, et une combinaison de traits flexionnels telle que la suivante :

futur - 3^e personne - singulier

on peut lui faire correspondre une forme conjuguée, ici *il prendra*. Dans la suite de ce chapitre, nous abordons les problèmes liés à la représentation de la conjugaison sous la forme d'un système formel. Nous ne nous intéressons ici qu'aux formes simples, et non aux formes composées telles que *il a pris*, qui sont des séquences de formes verbales simples existant aussi isolément. Les formes composées sont étudiées par Maurice Gross

(1968). Par ailleurs, en français courant, la notion traditionnelle de mode n'est plus guère séparable de celle de temps, ni sur des bases morphologiques, ni sur des bases syntaxiques⁷. A la suite de Maurice Gross (1968, p. 10), nous confondons donc sous le terme de "temps" ces deux notions traditionnelles et nous distinguons onze temps simples⁸ pour lesquels nous utilisons les abréviations suivantes :

<i>Vpr</i>	présent	<i>Vps</i>	passé simple
<i>Vimpft</i>	imparfait	<i>Vimptf</i>	impératif
<i>Vfut</i>	futur	<i>Vinf</i>	infinitif
<i>Vcond</i>	conditionnel	<i>V-ant</i>	participe présent
<i>Vsubj</i>	subjonctif	<i>Vpp</i>	participe passé
<i>Vimpft-subj</i>	imparfait du subjonctif		

Ces abréviations ont l'avantage d'être relativement transparentes, en particulier pour noter les combinaisons temps-personne-nombre : ainsi, *Vfut1,2,3sg* signifie "futur, 1^o, 2^o et 3^o personnes du singulier".

3. Les notions de radical et de terminaison

Les notions de radical et de terminaison, prises dans un sens strict, se réfèrent à un système idéal où chaque forme conjuguée serait constituée d'un radical qui ne dépendrait que du verbe employé (il s'agirait donc d'une donnée lexicale) et d'une terminaison qui ne dépendrait que des traits flexionnels de temps, personne, genre et nombre. L'ensemble de la conjugaison en français n'est pas un système aussi régulier. On peut cependant définir des notions algébriques de radical et de terminaison. C'est ce que nous allons faire. Pour simplifier, nous nous intéresserons pour l'instant à des formes orthographiques et non à leur prononciation. Pour chaque verbe, on peut définir son "radical invariable" comme la plus longue séquence initiale qui ne dépend ni du temps, ni de la personne, ni du genre, ni du nombre, par exemple *pr-* pour le verbe *prendre*. La séquence *pren-* serait trop longue par rapport à cette définition, comme le montre le rapprochement entre l'infinitif *prendre* et le participe passé *pris*. De même, pour chaque combinaison possible de traits flexionnels, on peut définir comme "terminaison caractéristique" la plus longue séquence terminale commune à tous les verbes, par exemple *-ra* pour le *Vfut3sg*, car cette terminaison se rencontre aussi bien dans *il prendra* que dans *il aura* et au futur de tous les verbes. Si l'on fait ce découpage pour quelques formes verbales, on constate que :

- Le radical invariable, tel que nous l'avons défini, peut être nul. C'est le cas du verbe *aller*, comme le montrent les formes *aller* et *il ira*, qui n'ont aucun élément initial commun ;

⁷ Ainsi, à la suite de la disparition progressive de l'imparfait du subjonctif, le subjonctif ne compte plus qu'un seul temps simple, comme le conditionnel et l'impératif ; de plus, l'imparfait, traditionnellement considéré comme un temps, se comporte comme un mode dans certaines subordonnées telles que celle de la phrase suivante :

Si Guy trichait, il gagnerait.

⁸ Dans un souci d'exhaustivité, nous avons pris en compte dans les tableaux et dans les programmes le passé simple et l'imparfait du subjonctif, bien qu'ils aient disparu de la langue parlée standard.

- La terminaison caractéristique peut aussi être nulle. C'est le cas du *Vpp* au masculin singulier : *perdu* et *gardé* n'ont aucun suffixe commun.

Bien sûr, chaque forme conjuguée commence par le radical ainsi défini, et finit par la terminaison ainsi définie, mais on constate que les limites de ces éléments ne correspondent généralement pas. Par exemple, nous avons vu que le radical invariable de *prendre* est *pr-* et que la terminaison caractéristique du *Vfut3sg* est *-ra* : en conséquence, dans le *Vfut3sg* *il prendra*, ces deux éléments sont séparés par un élément *-end-*. Même dans le cas d'une conjugaison régulière comme celle de *garder*, le *e* de *il gardera* ne fait partie ni du radical invariable du verbe (cf. *il gardait*), ni de la terminaison caractéristique du *Vfut3sg* (cf. *il prendra*). Les définitions algébriques précédentes permettent donc, en fait, de découper chaque forme verbale en trois parties. Appelons "élément intermédiaire" le tronçon obtenu en soustrayant du mot d'une part le radical invariable et d'autre part la terminaison caractéristique :

-end- dans *il pr-end-ra*, *-e-* dans *il gard-e-ra*

Dès lors, par définition :

- le radical invariable ne dépend que du verbe employé ;
- la terminaison caractéristique ne dépend que des traits flexionnels ;
- l'élément intermédiaire dépend des deux variables.

Dans aucun verbe on n'observe de chevauchement entre le radical et la terminaison tels que nous les avons définis, ce qui se produirait par exemple si le futur *il courra* s'écrivait avec un seul *r*. C'est cela qui rend possible le découpage en trois parties. L'élément intermédiaire peut toutefois être nul : c'est le cas dans les formes *il perdra* et *il gardait*.

La conclusion de ces jeux algébriques est que le découpage rigoureux en deux éléments, un radical et une terminaison, n'est pas possible si on prend la conjugaison dans son ensemble, en incluant tous les verbes. Nous n'avons pris que des exemples orthographiques, mais il en est de même phonétiquement et phonémiquement. La conjugaison des verbes est relative à un mode de représentation : orthographique, phonémique ou phonétique. Par exemple, *filer* et *figer* n'ont pas la même conjugaison orthographique. Leurs radicaux invariables sont *fil-* et *fig-* ; à certains temps, l'élément qui s'y ajoute est le même pour les deux verbes : *fil-é*, *fig-é* ; à d'autres temps, cet élément est différent pour les deux verbes : *il fil-ait*, *il fig-eait*. Toutefois, ces éléments sont toujours phonétiquement identiques, donc les deux verbes ont la même conjugaison phonétique. Un autre exemple est celui des verbes *croire* et *pouvoir*. A cause du *e* muet final de l'infinitif *croire*, ces deux verbes n'ont pas la même conjugaison orthographique, mais à part cette différence ils se conjuguent de la même façon. Comme le *e* muet ne se prononce ni plus ni moins dans *croire* que dans *pouvoir*, nous avons pu les ranger dans la même conjugaison phonémique.

Pour décrire les conjugaisons, nous utilisons un découpage en trois parties, comme ci-dessus, mais qui ne correspond pas exactement à celui que nous venons de définir. Notre découpage n'est pas défini d'une manière aussi absolue, et il est parfois un peu arbitraire, mais il permet d'explicitier d'une façon simple toutes les conjugaisons. Nous allons voir sur quelles bases nous effectuons ce découpage.

4. Le radical

Le premier élément est le radical invariable tel que nous l'avons défini plus haut : il ne dépend que du verbe et reste inchangé dans toute la conjugaison. Ainsi, le radical de *ouvrir* est *ouv-*, comme le montre le rapprochement entre *ouvrir* et *ouvert*.

Cette notion de radical invariable fournit un principe de classification pour répartir les verbes en classes de conjugaison. En effet, considérons un verbe qui se conjugue de la même façon que *ouvrir* : *offrir*. Le radical invariable de *offrir* est *off-*. Pour vérifier que les deux verbes se conjuguent sur le même modèle, il suffit de constater que les éléments qui s'ajoutent au radical invariable, c'est-à-dire *-rir* dans *ouvrir* et *-ert* dans *ouvert*, sont les mêmes pour les deux verbes, et ce à tous les temps et à toutes les personnes. Plus généralement, on peut dire que deux verbes se conjuguent sur le même modèle s'ils diffèrent seulement dans leur radical invariable et non dans le reste du mot. Cette propriété est une définition des classes de conjugaison ou conjugaisons. En fait, ce principe de répartition des verbes en conjugaisons convient aux conjugaisons phonémiques, mais il n'est pas adapté aux conjugaisons orthographiques, car il multiplie trop les classes. En effet, en raison des variations entre *é* et *è* observées dans certains verbes en *-er* (*céder*, il *cède*; *lécher*, il *lèche*; *paramétrer*, il *paramètre*; *téter*, il *tète*; *disséquer*, il *dissèque*...), le radical invariable, tel que nous l'avons défini plus haut, se réduirait dans ces verbes à *c-éder*, *l-écher*, *param-étrer*, *t-éter*, *diss-équér*. Si on utilisait la définition ci-dessus pour classer ces verbes selon leur conjugaison orthographique, on devrait les ranger dans autant de classes distinctes, ce qui serait peu satisfaisant. Lorsqu'on passe à une représentation phonémique, ce problème disparaît et on obtient moins de classes. En effet, les variations entre [e] et [] dans ces mots sont conditionnées par le contexte, et on peut y représenter [e] et [] par le phonème /e/. Le radical invariable de ces verbes devient alors /sed/, /leS/, /parametr/, /tet/, /disek/, et ils se conjuguent tous de la même façon. Effectivement, nous avons étudié la conjugaison phonémique des verbes en utilisant la notion de radical invariable, nous avons classé les verbes suivant la règle ci-dessus, et les résultats obtenus sont satisfaisants ; en particulier, cette classification compte environ 65 conjugaisons.

Le tableau suivant donne la correspondance entre les conjugaisons orthographiques et les conjugaisons phonémiques. Les conjugaisons phonémiques sont listées à gauche, et les conjugaisons orthographiques à droite. Chaque classe est spécifiée par un numéro et un verbe type. Les deux colonnes de nombres sont les numéros de conjugaisons. La concordance est simple, sauf lorsqu'une seule conjugaison phonémique, par exemple celle de *chanter*, correspond à plusieurs conjugaisons orthographiques (*chanter*, *placer*, *manger*, *semer*, *gérer*...). Dans ce cas, la classe phonémique est citée une seule fois à gauche, et les classes orthographiques qu'elle regroupe sont détaillées à droite.

Cette liste de conjugaisons ne tient pas compte du fait que certaines formes n'existent pas, comme les formes du verbe *falloir* à la 1^{re} et à la 2^{de} personnes.

Correspondance entre les conjugaisons

Conjugaisons phonémiques

1	<i>avoir</i>
2	<i>être</i>
3	<i>chanter</i>
13	<i>noyer</i>
15	<i>envoyer</i>
16	<i>aller</i>
18	<i>finir</i>
19	<i>haïr</i>
20	<i>acquérir</i>
21	<i>assaillir</i>
23	<i>couvrir</i>
24	<i>cueillir</i>
25	<i>fleurir (florissant)</i>
27	<i>fuir</i>
28	<i>sentir</i>
30	<i>vêtir</i>
31	<i>courir</i>
32	<i>mourir</i>
33	<i>tenir</i>
34	<i>ouïr</i>
36	<i>gésir</i>
37	<i>chauvir</i>
40	<i>asseoir</i>
41	<i>devoir</i>
42	<i>mouvoir</i>
44	<i>pouvoir</i>
45	<i>prévoir</i>

Conjugaisons orthographiques

1	<i>avoir</i>
2	<i>être</i>
3	<i>chanter</i>
4	<i>placer</i>
5	<i>manger</i>
6	<i>semer</i>
7	<i>gérer</i>
8	<i>jeter</i>
9	<i>appeler</i>
10	<i>assiéger</i>
11	<i>dépecer</i>
12	<i>rapiecer</i>
13	<i>noyer</i>
15	<i>envoyer</i>
16	<i>aller</i>
18	<i>finir</i>
19	<i>haïr</i>
20	<i>acquérir</i>
21	<i>assaillir</i>
23	<i>couvrir</i>
24	<i>cueillir</i>
25	<i>fleurir (florissant)</i>
27	<i>fuir</i>
86	<i>conclure</i>
95	<i>rire</i>
28	<i>sentir</i>
22	<i>bouillir</i>
26	<i>dormir</i>
29	<i>servir</i>
30	<i>vêtir</i>
31	<i>courir</i>
32	<i>mourir</i>
33	<i>tenir</i>
34	<i>ouïr</i>
36	<i>gésir</i>
37	<i>chauvir</i>
40	<i>asseoir</i>
41	<i>devoir</i>
46	<i>recevoir</i>
42	<i>mouvoir</i>
59	<i>promouvoir</i>
44	<i>pouvoir</i>
45	<i>prévoir</i>

47	<i>savoir</i>	47	<i>savoir</i>
48	<i>surseoir</i>	48	<i>surseoir</i>
49	<i>valoir</i>	49	<i>valoir</i>
		55	<i>falloir</i>
50	<i>voir</i>	50	<i>voir</i>
51	<i>vouloir</i>	51	<i>vouloir</i>
53	<i>choir</i>	53	<i>choir</i>
56	<i>pleuvoir</i>	56	<i>pleuvoir</i>
57	<i>seoir</i>	57	<i>seoir</i>
58	<i>prévaloir</i>	58	<i>prévaloir</i>
61	<i>absoudre</i>	61	<i>absoudre</i>
62	<i>coudre</i>	62	<i>coudre</i>
63	<i>moudre</i>	63	<i>moudre</i>
64	<i>peindre</i>	64	<i>peindre</i>
65	<i>résoudre</i>	65	<i>résoudre</i>
66	<i>prendre</i>	66	<i>prendre</i>
67	<i>perdre</i>	67	<i>perdre</i>
		60	<i>vaincre</i>
		68	<i>battre</i>
		76	<i>foutre</i>
		77	<i>rompre</i>
69	<i>connaître</i>	69	<i>connaître</i>
70	<i>croître</i>	70	<i>croître</i>
		73	<i>décroître</i>
71	<i>mettre</i>	71	<i>mettre</i>
72	<i>naître</i>	72	<i>naître</i>
74	<i>suivre</i>	74	<i>suivre</i>
75	<i>vivre</i>	75	<i>vivre</i>
78	<i>fiche</i>	78	<i>fiche</i>
79	<i>joindre</i>	64(79)	<i>joindre</i>
80	<i>faire</i>	80	<i>faire</i>
81	<i>traire</i>	81	<i>traire</i>
82	<i>plaire</i>	82	<i>plaire</i>
		88	<i>taire</i>
83	<i>boire</i>	83	<i>boire</i>
84	<i>croire</i>	84	<i>croire</i>
		43	<i>pourvoir</i>
87	<i>inclure</i>	87	<i>inclure</i>
90	<i>confire</i>	90	<i>confire</i>
		89	<i>suffire</i>
91	<i>cuire</i>	91	<i>cuire</i>
		97	<i>nuire</i>
92	<i>écrire</i>	92	<i>écrire</i>
93	<i>dire</i>	93	<i>dire</i>
94	<i>lire</i>	94	<i>lire</i>
96	<i>maudire</i>	96	<i>maudire</i>

98	<i>circoncire</i>	98	<i>circoncire</i>
		85	<i>clore</i>

5. L'élément intermédiaire et la terminaison

Le radical invariable une fois isolé, il reste à expliciter les variations du reste du mot, qui est *-endre* dans *prendre* et *-ir* dans *courir*. Comme nous l'avons vu, cette partie droite du mot dépend bien sûr des traits flexionnels de temps, personne, genre et nombre, mais aussi du verbe employé : *courir* fait son *Vpp* en *-u* et *dormir* le fait en *-i*. Toutefois, les variations en fonction du verbe sont de toutes façons limitées. En effet, des verbes qui se conjuguent de la même façon, comme *ouvrir* et *offrir*, diffèrent seulement dans leur radical invariable et non dans le reste du mot, par définition. Le reste du mot dépend donc de la classe du verbe, plutôt que du verbe lui-même : à l'intérieur d'une classe de conjugaison, cette partie de mot est indépendante du verbe considéré. Ainsi, les variations de la partie droite du mot sont les mêmes pour *ouv-rir* que pour *couv-rir* ou *off-rir*.

Pour expliciter ces variations, il est commode de découper la partie droite en deux éléments : un élément intermédiaire et une terminaison. Cependant, nous n'avons pas repris la notion de terminaison caractéristique évoquée ci-dessus. Cette notion a l'avantage d'avoir une définition rigoureuse. Elle est intéressante à certains temps comme l'imparfait, car tous les imparfaits sont en *-ait*. Mais à d'autres temps, les terminaisons sont différentes suivant les conjugaisons, et n'ont souvent aucun élément commun : ainsi, on trouve des *Vpp* en *-é, -i, -u*, etc. et des *Vinf* en *-er, -ir, -oir*, etc. Les "terminaisons caractéristiques" de ces temps sont donc vides, alors que les éléments *-é, -i, -u, -er, -ir, -oir*, fonctionnent comme des terminaisons : ils sont peu nombreux et commutent avec d'autres terminaisons. C'est pourquoi nous avons préféré d'autres critères de découpage des formes verbales. Pour voir sur quelles bases on peut effectuer ce découpage, prenons d'abord un exemple simple : celui de la classe *chanter*.

Le verbe *chanter* est régulier. A son radical /Sant/ s'ajoutent des suffixes répandus dans de nombreux verbes. Nous donnons à ces suffixes le statut de terminaison, ce qui laisse vide l'élément intermédiaire. Voici les valeurs de la terminaison pour les principaux temps de ce verbe :

Vimpft = : il chantait /Sante/. La terminaison /e/ se retrouve dans les *Vimpft* de tous les verbes.

Vinf = : chanter /Sante/, comme dans *noyer, envoyer, aller*.

Vpr1,2,3sg = : il chante /Sant/. Cette terminaison vide se retrouve dans de nombreux *Vpr1,2,3sg* : il court, il sait, il voit, ...

Vpr3pl = : ils chantent /Sant/. La terminaison vide est également répandue au *Vpr3pl* : ils courent, ils veulent, ils voient, ils disent, ...

Vpp = : chanté /Sante/, comme dans *noyé, envoyé, allé, été*.

Au futur et au conditionnel, nous représentons par le phonème /' / le *e* muet qui précède la finale /ra/, /re/, etc. : il chantera /Sant'ra/. Cet *e* muet peut se prononcer : il se prononce d'ailleurs généralement dans *vous aimeriez, il parlera* ou *il gagnera*.

Prononcé ou non, il a la propriété d'ouvrir obligatoirement en [ɛ], [æ], [ɔ] un /e/, /o/, /o/ de la syllabe précédente :

<i>il cèdera</i>	[sɛdra] *[sedra]
<i>il pleurera</i>	[plœrra] *[plœrra]
<i>il donnera</i>	[dɔnra] *[donra]

Avec cette représentation, la partie de mot qui s'ajoute au radical est /'ra/, /'re/, ... Les finales /ra/, /re/, ..., se retrouvent dans tous les futurs : *il chantera, il courra, il voudra, il rendra, il dira, il verra, il prendra*. Etant donné la généralité de ces finales, nous découpons /Sant'ra/ de la façon suivante :

- radical /Sant/
- élément intermédiaire /'/
- terminaison /ra/.

Nous faisons de même au futur et au conditionnel de tous les verbes de la classe.

Etant donné ce découpage des formes verbales de la classe *chanter*, on peut considérer séparément les variations de l'élément intermédiaire et celles de la terminaison.

En ce qui concerne l'élément intermédiaire dans la classe *chanter*, nous lui avons donné deux valeurs : /'/ au futur et au conditionnel, et zéro aux autres formes. L'élément intermédiaire prend donc des valeurs peu nombreuses par rapport au nombre total des formes conjuguées distinctes, qui est 11 sans compter le *Vps* et le *Vimprf-subj*. De plus, les valeurs de l'élément intermédiaire ne sont pas distribuées au hasard parmi les combinaisons temps-personne-nombre. Au contraire, cette distribution reflète la distinction entre les temps verbaux, et aussi la parenté entre le futur et le conditionnel : pour tous les verbes français, le futur et le conditionnel se construisent sur un même radical. Toutes ces remarques sont des propriétés du découpage des formes conjuguées. Elles se résument ainsi : pour un verbe de la classe *chanter*, l'élément intermédiaire varie lors de la conjugaison, mais ces variations sont limitées. Nous verrons que pour chaque conjugaison, on peut trouver un découpage des formes verbales qui vérifie cette propriété.

Quant aux valeurs de la terminaison, nous les avons passées en revue ci-dessus et nous avons remarqué que chacune d'elles est un suffixe répandu dans de nombreux verbes : les terminaisons /e/, /ion*/ et /ie/ de l'imparfait se retrouvent même dans tous les verbes français⁹. D'autres terminaisons sont moins générales, mais sont communes à quelques classes de conjugaisons : ainsi, le /e/ du *Vpp* se retrouve dans les classes *noyer, envoyer, aller, être* ; le /e/ du *Vinf* se retrouve dans *noyer, envoyer* et *aller*. Si on regarde les valeurs de la terminaison dans les autres *Vinf*, on constate qu'elles sont en petit nombre : /ir/ dans *courir*, /uar/ dans *valoir*, /r/ dans *rendre* et zéro dans *fiche*. On peut vérifier qu'il en est de même pour les *Vpp*, pour les *Vpr3pl* et pour toutes les combinaisons temps-personne-nombre. En d'autres termes, étant donné une combinaison temps-personne-nombre, la terminaison correspondante n'est pas la même

⁹ Dans la terminaison /ion*/, la séquence /on*/ représente la voyelle nasale finale [ɔ̃] (cf. chapitre III, section 7, p. 61).

pour toutes les classes, mais ces variations en fonction de la classe sont limitées. On peut faire le découpage de toutes les formes verbales de sorte que cette propriété soit toujours vérifiée.

L'examen des conjugaisons moins régulières que celle de *chanter* confirme ces constatations. Il est possible d'effectuer le découpage entre l'élément intermédiaire et la terminaison de telle sorte que :

- pour un verbe donné, l'élément intermédiaire connaisse seulement des variations limitées lors de la conjugaison ;

- étant donné une combinaison temps-personne-nombre, la terminaison ait seulement des variations limitées en fonction des classes considérées.

Lorsque le découpage satisfait à ces deux propriétés, il permet de décrire séparément les variations de l'élément intermédiaire et celles de la terminaison. On peut alors expliciter d'une façon simple et maniable les variations de forme des verbes lors de la conjugaison. Ces deux propriétés, ou un compromis entre elles, sont donc la raison d'être et l'objectif du découpage. Elles constituent un critère de découpage, quoiqu'il ne s'agisse pas d'un critère rigoureux tel que celui que nous avons présenté plus haut. Dans les pages qui suivent, nous prenons en compte l'ensemble du système de conjugaisons français et nous décrivons successivement les variations de forme de l'élément intermédiaire et de la terminaison, afin de donner des exemples d'un découpage commode et de mettre en évidence les deux propriétés ci-dessus. Le fait qu'un tel découpage soit possible est une propriété du système de conjugaisons français.

6. L'élément intermédiaire

Dans la classe *chanter*, nous avons donné deux valeurs à l'élément intermédiaire : /' / au futur et au conditionnel, et zéro ailleurs. Le cas de la classe *courir* est encore plus simple : l'élément intermédiaire y est vide à toutes les formes conjuguées. En effet, toutes les terminaisons qui se rattachent au radical invariable /kur/ sont des suffixes répandus dans plusieurs conjugaisons, comme le montre l'examen des principaux temps :

<i>Vimpft</i> =: il cour-ait	/kur-e/
<i>Vinf</i> =: cour-ir	/kur-ir/
<i>Vpr1,2,3sg</i> =: il court	/kur/
<i>Vpr3pl</i> =: ils courent	/kur/
<i>Vfut</i> =: il cour-ra	/kur-ra/
<i>Vpp</i> =: cour-u	/kur-y/

Toutefois, dans la plupart des verbes, l'élément intermédiaire prend entre 2 et 4 valeurs qui se distribuent sur la conjugaison. Le verbe *mourir* a pour radical invariable /m/, et l'élément intermédiaire prend les valeurs /ur/, /0r/ et /or/ :

<i>Vimpft</i> =: il m-our-ait	/m-ur-e/
<i>Vinf</i> =: m-our-ir	/m-ur-ir/
<i>Vpr1,2,3sg</i> =: il m-eurt	/m-0r/
<i>Vpr3pl</i> =: ils m-eurent	/m-0r/
<i>Vfut</i> =: il m-our-ra	/m-ur-ra/
<i>Vpp</i> =: m-or-t	/m-or-t*/

Dans chaque conjugaison, les différentes valeurs de l'élément intermédiaire se répartissent sur les formes conjuguées. Cette distribution suit quelques règles générales. Toutes les formes de l'imparfait se construisent toujours avec un même élément intermédiaire, qui est /ur/ pour mourir (*je m-our-ais, tu m-our-ais, il m-our-ait, nous m-our-ions, vous m-our-iez, ils m-our-aient*) et zéro pour chanter (*je chant-ais, tu chant-ais, il chant-ait, nous chant-ions, vous chant-iez, ils chant-aient*). Le futur et le conditionnel partagent toujours un même élément intermédiaire, qui est vide pour courir (*il cour-ra, il cour-ra-it*), /' / pour parler (*il parl-e-ra, il parl-e-ra-it*), /i/ pour dormir (*il dorm-i-ra, il dorm-i-ra-it*), /d/ pour coudre (*il cou-d-ra, il cou-d-ra-it*). Même les verbes *aller, avoir* et *être*, qui sont très irréguliers, suivent ces deux règles sur l'imparfait, le futur et le conditionnel. Pour les autres temps, la distribution des valeurs de l'élément intermédiaire suit des règles un peu moins générales. Ainsi, l'élément intermédiaire est souvent le même au *Vpr3pl*, au *Vsubj1,2,3sg* et au *Vsubj3pl*. C'est le cas pour mourir (*ils m-eurent, qu'il m-eure, qu'ils m-eurent*), pour dormir (*ils dorment, qu'il dorme*), pour prendre (*ils pr-ennent, qu'il pr-enne*), et pour la plupart des conjugaisons. Les verbes *aller, avoir* et *être* ne suivent pas cette règle. Il y a plusieurs autres exceptions : *valoir* (*ils v-alent, mais qu'il v-aille*), *pouvoir* (*ils p-euvent, mais qu'il p-uisse*), *savoir* (*ils s-avent, mais qu'il s-ache*), *vouloir* (*ils v-eulent, mais qu'il v-uille*) et *faire* (*ils f-ont, mais qu'il f-asse*). Toutefois, la distribution de l'élément intermédiaire montre, pour la plupart des verbes, une parenté entre le *Vpr3pl*, le *Vsubj1,2,3sg* et le *Vsubj3pl*.

Plus généralement, les règles de distribution des valeurs de l'élément intermédiaire, lorsqu'on observe toutes les conjugaisons, constituent un critère de parenté entre les formes. Par exemple, les six formes de l'imparfait sont proches. Il n'en est pas de même pour les six formes du subjonctif. Le *Vsubj1,2,3sg* et le *Vsubj3pl* se rattachent généralement au *Vpr3pl* (*qu'il meure, ils meurent*) alors que le *Vsubj1,2pl* se rattache généralement au *Vimpft* (*que vous mouriez, il mourait*).

Ce critère de parenté entre les formes permet d'établir une classification des 54 combinaisons temps-personne-genre-nombre. Une telle classification, si elle est construite en tenant compte de tous les verbes, est indépendante des classes de conjugaison. Elle a nécessairement une part d'arbitraire, puisqu'elle n'est pas fondée sur un critère absolu, mais elle facilite l'explicitation des variations de forme des verbes. Nous avons adopté une classification des 54 combinaisons temps-personne-genre-nombre en 7 séries, qui reflète à peu près les parentés entre formes indépendamment des classes de conjugaison. Voici cette classification.

1) *Vimpft*, *Vpr1,2pl*, *V-ant*, *Vsubj1,2pl*, *Vimpft1,2pl*. Dans la plupart des conjugaisons, toutes ces formes se construisent avec un même élément intermédiaire. C'est le cas de la classe *boire*, qui utilise l'élément intermédiaire /yv/ pour toutes les formes de cette série :

<i>Vimpft</i> =: il b-uv-ait	/b-yv-e/
<i>Vpr1,2pl</i> =: vous b-uv-ez	/b-yv-e/
<i>V-ant</i> =: en b-uv-ant	/b-yv-ant*/
<i>Vsubj1,2pl</i> =: que vous b-uv-iez	/b-yv-ie/
<i>Vimpft1,2pl</i> =: b-uv-ez	/b-yv-e/

Il en est de même de presque toutes les classes : *tenir, valoir, prendre, voir, aller, ...* Les verbes *avoir* et *être* ne suivent pas cette règle, et il y a plusieurs autres exceptions : *pouvoir* (il *p-ouv-ait*, mais *que vous p-uiss-iez*), *savoir* (il *s-av-ait*, mais *en s-ach-ant, que vous s-ach-iez, s-ach-ez*), *vouloir* (il *v-oul-ait*, mais *que vous v-euill-iez, v-euill-ez*), *dire* (il *di-s-ait*, mais *vous di-tes, di-tes*), *faire* (il *f-ais-ait*, mais *vous f-aîtes, que vous f-ass-iez, f-aîtes*), et peut-être *déchoir* et *seoir*.

2) *Vinf*. Cette série ne comprend qu'une seule forme, l'infinitif.

3) *Vpr1,2,3sg, Vimptf2sg*. Ces formes comportent presque toujours un même élément intermédiaire, qui est /0r/ dans la classe *mourir* :

Vpr1,2,3sg = : il *m-eurt* /m-0r/

Vimptf2sg = : *m-eurs* /m-0r/

Les verbes *aller, avoir* et *être* sont des exceptions, ainsi que *savoir* (il *s-ait* mais *s-ache*), *vouloir* (il *v-eut* mais *v-euille*) et *pouvoir* (pour la forme *je p-uïs*).

4) *Vpr3pl, Vsubj1,2,3sg, Vsubj3pl*. Nous avons déjà évoqué la parenté entre ces formes et donné les exceptions.

5) *Vfut, Vcond*. Nous avons vu que ces deux temps se construisent toujours avec un même élément intermédiaire.

6) *Vpp*. Cette série ne comprend que le participe passé.

7) *Vps, Vimptf-subj*. Comme le futur et le conditionnel, ces deux temps se construisent toujours avec un même élément intermédiaire.

Dans chacune de ces 7 séries, le temps cité en premier ci-dessus est considéré arbitrairement comme un temps primitif. Par exemple, nous considérons que le *Vsubj1,2,3sg* se construit sur le modèle du *Vpr3pl*, et non le contraire.

Le découpage en trois parties, effectué sur tous les verbes, permet une description systématique des conjugaisons phonémiques. Cette description prend la forme d'une table morphologique donnée pp. 94-95, dont les lignes représentent des classes de conjugaison, et dont les colonnes correspondent à des temps ou à des combinaisons de traits flexionnels. Les colonnes relevant d'une même série sont présentées côte à côte : il y a donc 7 séries de colonnes pour les 7 séries de temps.

7. La terminaison

La terminaison d'une forme verbale dépend bien sûr des traits flexionnels : dans *mourir*, /ir/ est caractéristique de l'infinitif et commute avec la terminaison /e/ de l'imparfait. Mais nous avons vu que le découpage des formes verbales peut se faire de façon à ce que pour une combinaison temps-personne-nombre donnée, la terminaison ait seulement des variations limitées en fonction des classes de conjugaison. Par exemple, les *Vinf* des différents verbes ont des terminaisons distinctes : /e/ dans *manger*,

/ir/ dans *courir*, /uar/ dans *valoir*, /r/ dans *rendre* et zéro dans *fiche*, mais elles sont en petit nombre.

Aux autres temps que l'infinitif, les variations de la terminaison en fonction de la classe de conjugaison sont aussi peu importantes ou encore moins. Les terminaisons de l'imparfait, /e/, /ion*/ et /ie/, se retrouvent dans toutes les conjugaisons, y compris *aller*, *avoir* et *être*. Les terminaisons du futur et du conditionnel, /ra/, /re/, etc., sont tout aussi générales. La terminaison /ant*/ du participe présent est également commune à toutes les conjugaisons. Au participe passé, on note les terminaisons /e/ de *chanté*, /i/ de *dormi*, /y/ de *couru*, /z*/ de *mis*, /t*/ de *mort*, et la terminaison vide de *fini* (dans ce verbe, le /i/ appartient au radical invariable puisqu'il se retrouve dans toute la conjugaison).

Le *Vpr1,2,3sg* pose un problème de représentation de la terminaison. Phonétiquement, tous les verbes ont une terminaison nulle au *Vpr1,2,3sg* : *il chante* /Sant/, *il court* /kur/, *il faut* /fo/, *il résout* /rezu/, *il voit* /vua/. Cependant, de nombreux verbes perdent une consonne au *Vpr1,2,3sg*, par exemple *rendre*. A l'imparfait de ces verbes, la terminaison /e/ est précédée d'une consonne, qui est /d/ dans le cas de *rendre* :

il rend-ait /rand-e/

Cette consonne se retrouve aux temps construits sur l'imparfait (cf. chapitre V, section 2, p. 100) :

<i>vous rend-ez</i>	/rand-e/
<i>en rend-ant</i>	/rand-ant*/
<i>que vous rend-iez</i>	/rand-ie/
<i>rend-ez</i>	/rand-e/

ainsi qu'au *Vpr3pl* et aux formes qui en dérivent :

<i>ils rendent</i>	/rand/
<i>qu'il rende</i>	/rand/

La consonne disparaît au *Vpr1,2,3sg* et au *Vimptf2sg* :

<i>il rend</i>	[rã] sans [d]
<i>rends</i>	[rã] sans [d]

Voici quelques autres verbes qui présentent la même particularité :

<i>ils dorment</i>	<i>il dort</i>	sans [m]
<i>ils vivent</i>	<i>il vit</i>	sans [v]
<i>ils finissent</i>	<i>il finit</i>	sans [s]
<i>ils connaissent</i>	<i>il connaît</i>	sans [s]
<i>ils disent</i>	<i>il dit</i>	sans [z]

Ces effacements de consonnes finales sont étudiés au chapitre V, section 2, p. 100. Dans les représentations phonémiques, nous introduisons la consonne manquante et nous

signalons son effacement par une marque explicite notée /*/ :

<i>il rend</i>	/rand*/
<i>il dort</i>	/dorm*/
<i>il vit</i>	/viv*/
<i>il finit</i>	/finis*/
<i>il connaît</i>	/kones*/
<i>il dit</i>	/diz*/

Dans ces formes phonémiques, le radical et l'élément intermédiaire peuvent donc avoir la même valeur qu'au *Vpr3pl* et au *Vimpft*, puisque la consonne effaçable y figure. Pour cela, dans le découpage de ces formes, la marque /*/ doit être placée dans la terminaison. Le découpage de *il rend* est alors le suivant :

Radical	/rand/
Élément intermédiaire vide	
Terminaison	/*/

La terminaison du *Vpr1,2,3sg* et du *Vimpft2sg* de ces verbes est donc /*/. Dans les verbes où il ne se produit pas d'effacement de consonne finale au *Vpr1,2,3sg*, la terminaison est vide : *il chante* /Sant/. Avec ce système de représentation, la valeur de la terminaison au *Vpr1,2,3sg* indique simplement si la consonne finale s'efface.

La table de conjugaisons phonémiques des pp. 94-95 donne les valeurs de l'élément intermédiaire et de la terminaison pour toutes les conjugaisons. Ces valeurs constituent des critères formels pour la classification morphologique des verbes. Les classes de conjugaison sont toutes différentes, par définition, mais certaines ont des points communs. On peut ainsi, sur la base de ces critères formels, ranger en 7 séries les 65 classes de conjugaison :

La première série rassemble les classes *chanter*, *noyer*, *envoyer* et *aller*. Elle est caractérisée par le fait que le *Vinf* (*chant-er*) et le *Vpp* (*chant-é*) sont semblables : il sont construits sur le *Vimpft* (*il chant-ait*) et comportent la terminaison /e/.

Deuxième série (*acquérir*, *assaillir*, *offrir*, *cueillir*, *sentir*, *vêtir*, *courir*, *mourir*, *tenir*, *gésir*, *chauvir*) : le *Vinf* (*sent-ir*) est construit sur le *Vimpft* (*il sent-ait*) et comporte la terminaison /ir/.

Troisième série (*avoir*, *devoir*, *mouvoir*, *pouvoir*, *savoir*, *valoir*, *vouloir*, *pleuvoir*, *prévaloir*) : le *Vinf* (*d-ev-oir*) est construit sur le *Vimpft* (*il d-ev-ait*) et comporte la terminaison /uar/ qui est précédée d'une des consonnes /v/ et /l/ ; le *Vpp* est en /y/ (*d-û*).

Quatrième série (*être*, *rendre*, *mettre*, *suivre*, *vivre*, *fiche*) : le *Vinf* (*rend-re*) est construit sur le *Vimpft* (*il rend-ait*) et comporte la terminaison /r/ qui est précédée d'une consonne ; le *Vpr3pl* (*ils rendent*) et le *Vfut* (*il rend-ra*) sont également construits sur le *Vimpft*.

Cinquième série (*finir, haïr, fleurir, faire, plaie, boire, confire, cuire, écrire, dire, lire, maudire, circoncrire*) : le *Vinf* (*cui-re*) comporte la terminaison /r/ qui est précédée d'une voyelle. L'élément intermédiaire du *Vimpft* (*il cui-s-ait*) se termine par l'une des consonnes /z/, /v/, /s/, qui disparaît au *Vinf*, au *Vfut* (*il cui-ra*) et au *Vpp* (*cui-t*).

Sixième série (*fuir, ouïr, asseoir, prévoir, surseoir, voir, déchoir, seoir, traire, croire, inclure*) : le *Vinf* (*cr-oi-re*) comporte la terminaison /r/ qui est précédée d'une voyelle. L'élément intermédiaire du *Vimpft* (*il croyait, /kr-uaj-e/*) se termine par un /j/ ou un /+ / qui disparaît au *Vinf*, au *Vfut* (*il cr-oi-ra*) et au *Vpr123sg* (*il cr-oit*).

Septième série (*absoudre, coudre, moudre, peindre, joindre, résoudre, prendre, connaître, croître, naître*) : le *Vinf* (*pr-end-re*) comporte la terminaison /r/, qui est précédée d'une des consonnes /d/ et /t/. Cette consonne ne se retrouve pas au *Vimpft* (*il pr-en-ait*) : l'élément intermédiaire est différent au *Vinf* et au *Vimpft*. Le *Vfut* (*il pr-end-ra*) est construit sur le *Vinf*.

Numéro de conjugaison	Exemple	Radical invariable de l'exemple	Vimpft Elément intermédiaire	Vpri1pl = Vimpft	Vpri2pl = Vimpft	V-ant = Vimpft	Vsubj.2pl = V-ant	Vimpft1pl = Vimpft	Vimpft2pl = Vimpft	Vinf Elément intermédiaire	Terminaison
3	chanter	Sant	.	+	+	+	+	+	+	+	e
13	noyer	nua	j	+	+	+	+	+	+	+	ee
15	envoyer	env	uaj	+	+	+	+	+	+	+	ee
16	aller	.	al	+	+	+	+	+	+	+	ee
20	acquérir	ak	er	+	+	+	+	+	+	+	ir
21	assaillir	asa	j	+	+	+	+	+	+	+	ir
23	offrir	of	r	+	+	+	+	+	+	+	ir
24	cueillir	ko	j	+	+	+	+	+	+	+	ir
28	sentir	sant	.	+	+	+	+	+	+	+	ir
30	vêtir	vet	.	+	+	+	+	+	+	+	ir
31	courir	kur	.	+	+	+	+	+	+	+	ir
32	mourir	m	ur	+	+	+	+	+	+	+	ir
33	tenir	t	"n	+	+	+	+	+	+	+	ir
36	gésir	Z	iz	+	+	+	+	+	+	ez	ir
37	chauvir	So:v	.	+	+	+	+	+	+	+	ir
1	avoir	.	av	+	+	e	+	+	+	+	uar
41	devoir	d	"v	+	+	+	+	+	+	+	uar
42	mouvoir	m	uv	+	+	+	+	+	+	+	uar
44	pouvoir	p	uv	+	+	+	yis	-	-	+	uar
47	savoir	s	av	+	+	as	+	+	+	+	uar
49	valoir	v	al	+	+	+	+	+	+	+	uar
51	vouloir	v	ul	+	+	+	+	0j	0j	+	uar
56	pleuvoir	pl	ov	+	+	+	+	+	+	+	uar
58	prévaloir	prev	al	+	+	+	+	+	+	+	uar
2	être	.	et	son	et	+	sua	sua	sua	+	r
67	rendre	rand	.	+	+	+	+	+	+	+	r
71	mettre	m	et	+	+	+	+	+	+	+	r
74	suiivre	syiv	.	+	+	+	+	+	+	+	r
75	vivre	v	iv	+	+	+	+	+	+	+	r
78	fiche	fis	.	+	+	+	+	+	+	+	.
18	finir	fini	s	+	+	+	+	+	+	.	r
19	hair	h	ais	+	+	+	+	+	+	al	r
25	fleurir	fl	oris	-	+	-	-	-	-	Ori	r
80	faire	f	"z	+	et	+	as	+	et	e	r
82	plaire	pl	ez	+	+	+	+	+	+	e	r
83	boire	b	yv	+	+	+	+	+	+	ua	r
90	confire	konfi	z	+	+	+	+	+	+	.	r
91	cuire	kyl	z	+	+	+	+	+	+	.	r
92	écrire	ékri	v	+	+	+	+	+	+	.	r
93	dire	di	z	+	t	+	+	+	t	.	r
94	lire	l	iz	+	+	+	+	+	+	.	r
96	maudire	modi	s	+	+	+	+	+	+	.	r
98	circoncire	sirkonsi	z	+	+	+	+	+	+	.	r
27	fuir	fyl	+	+	+	+	+	+	+	.	r
34	ouir	.	uaj	+	+	+	+	+	+	u+i	r
40	asseoir	as	e	+	+	+	+	+	+	ua	r
45	prévoir	prev	uaj	+	+	+	+	+	+	ua	r
48	surseoir	syrs	uaj	+	+	+	+	+	+	ua	r
50	voir	v	uaj	+	+	+	+	+	+	ua	r
53	déchoir	deS	uaj	+	+	+	+	-	-	ua	r
57	seoir	e	e	-	-	+	-	-	-	ua	r
81	traire	tre	j	+	+	+	+	+	+	.	r
84	croire	kr	uaj	+	+	+	+	+	+	ua	r
87	inclure	enkly	+	+	+	+	+	+	+	.	r
61	absoudre	aps	olv	+	+	+	+	+	+	ud	r
62	coudre	ku	z	+	+	+	+	+	+	d	r
63	moudre	mu	l	+	+	+	+	+	+	d	r
64	peindre	pe	N	+	+	+	+	+	+	nd	r
65	résoudre	rez	olv	+	+	+	+	+	+	ud	r
66	prendre	pr	"n	+	+	+	+	+	+	and	r
69	connaître	kon	es	+	+	+	+	+	+	et	r
70	croître	kr	uas	+	+	+	+	+	+	uat	r
72	naître	n	es	+	+	+	+	+	+	et	r
79	joindre	Zu	aN	+	+	+	+	+	+	end	r

$V_{pr2,3sg}$ Elément intermédiaire	Eff. cons. fin.	$V_{impft2sg} = V_{pr2,3sg}$	$V_{pr1sg} = V_{pr2,3sg}$	V_{pr3pl} Elément intermédiaire	$V_{subj1,2,3sg,3pl} = V_{pr3pl}$	V_{fut} Elément intermédiaire	V_{pp} Elément intermédiaire Terminaison	V_{ps} Elément intermédiaire Terminaison				
+	+	+	+	+	+	+	e	+	a	3 chanter		
.	-	-	+	Vpr23sg	+	.	+	+	a	13 noyer		
ua	-	-	+	Vpr23sg	+	e	+	+	a	15 envoyer		
va	-	-	+	von*	aj	i	+	+	a	16 aller		
ier	-	-	+	Vpr23sg	+	impft	i	z*	vpp	20 acquérir		
+	-	-	+	+	+	i	+	i	+	21 assaillir		
+	-	-	+	+	+	ri	er	t*	+	23 offrir		
+	-	-	+	+	+	'	+	i	+	24 cueillir		
+	-	-	+	+	+	i	+	i	+	28 sentir		
+	-	-	+	+	+	i	+	y	+	30 vêtir		
+	-	-	+	+	+	impft	+	y	+	31 courir		
Or	-	-	+	Vpr23sg	+	impft	or	t*	+	32 mourir		
ien	+	+	+	Vpr23sg	+	iend	+	y	en*	33 tenir		
+	+	+	+	+	+	-	-	-	-	36 gésir		
i	-	-	+	+	+	i	+	i	+	37 chauvir		
+	+	e	e	on*	e	o	.	y	vpp	y	1 avoir	
uav	+	+	+	Vpr23sg	+	impft	.	y	vpp	y	41 devoir	
ov	+	+	+	Vpr23sg	+	impft	.	y	vpp	y	42 mouvoir	
ov	-	-	+	yi	Vpr23sg	yis	u	+	y	vpp	y	44 pouvoir
e	-	as	+	+	+	o	.	y	vpp	y	47 savoir	
o	-	+	+	+	+	od	+	y	+	y	49 valoir	
Ol	+	0j	+	Vpr23sg	0j	ud	+	y	+	y	51 vouloir	
ov	+	+	+	+	+	impft	.	y	vpp	y	56 pleuvoir	
o	-	+	+	+	+	od	+	y	+	y	58 prévaloir	
+	+	sua	syi	son*	sua	s"	+	e	f	y	2 être	
+	+	+	+	+	+	impft	+	y	+	i	67 rendre	
+	+	+	+	+	+	impft	i	z*	vpp	.	71 mettre	
+	+	+	+	+	+	impft	+	i	+	i	74 suivre	
+	+	+	+	+	+	impft	ek	y	vpp	y	75 vivre	
+	+	+	+	+	+	impft	+	y	-	-	78 fiche	
+	+	+	+	+	+	vinf	vinf	.	+	.	18 finir	
e	-	-	+	+	+	vinf	vinf	.	vinf	.	19 hair	
-	-	-	-	-	-	-	-	-	-	-	25 fleurir	
e	-	+	+	on*	as	"	vinf	t*	.	i	80 faire	
+	+	+	+	+	+	vinf	.	y	vpp	y	82 plaie	
uav	+	+	+	Vpr23sg	+	vinf	.	y	vpp	y	83 boire	
+	+	+	+	+	+	vinf	vinf	t*	vinf	.	90 confire	
+	+	+	+	+	+	vinf	vinf	t*	+	i	91 cuire	
+	+	+	+	+	+	vinf	vinf	t*	+	i	92 écrire	
+	+	+	+	+	+	vinf	vinf	t*	vinf	.	93 dire	
+	+	+	+	+	+	vinf	.	y	vpp	y	94 lire	
+	+	+	+	+	+	vinf	vinf	t*	vinf	.	96 maudire	
+	+	+	+	+	+	vinf	vinf	z*	vinf	.	98 circoncrire	
Vinf	-	+	+	Vinf	+	Vinf	Vinf	.	Vinf	.	27 fuir	
ua	-	+	+	Vpr23sg	+	o	Vinf	.	Vinf	.	34 ouir	
ie	-	+	+	+	+	ie	i	z*	vpp	.	40 asseoir	
Vinf	-	+	+	Vinf	+	Vinf	.	y	vpp	.	45 prévoir	
Vinf	-	+	+	Vinf	+	Vinf	i	z*	vpp	.	48 surseoir	
Vinf	-	+	+	Vinf	+	e	.	y	vpp	.	50 voir	
Vinf	-	+	+	Vinf	+	e	.	y	vpp	y	53 déchoir	
ie	-	-	-	Vpr23sg	+	ie	i	z*	-	-	57 seoir	
Vinf	-	+	+	Vinf	+	Vinf	vinf	t*	-	-	81 traire	
Vinf	-	+	+	Vinf	+	Vinf	.	y	vpp	y	84 croire	
Vinf	-	+	+	Vinf	+	Vinf	vinf	z*	vinf	.	87 inclure	
u	-	+	+	+	+	Vinf	u	t*	ol	y	61 absoudre	
+	+	+	+	+	+	Vinf	+	y	+	i	62 coudre	
+	+	+	+	+	+	Vinf	+	y	+	y	63 moudre	
n	-	+	+	+	+	Vinf	n	t*	+	i	64 peindre	
u	-	+	+	+	+	Vinf	ol	y	vpp	y	65 résoudre	
an	+	+	+	+	+	Vinf	i	z*	vpp	.	66 prendre	
+	+	+	+	+	+	Vinf	.	y	vpp	y	69 connaître	
+	+	+	+	+	+	Vinf	.	y	vpp	y	70 croître	
+	+	+	+	+	+	Vinf	.	e	ak	i	72 naître	
en	+	+	+	+	+	Vinf	en	t*	+	i	79 joindre	

Chapitre V. Variations phonétiques

Nous parlons de variations phonétiques entre des formes phonétiques apparentées lorsque ces formes présentent une différence de prononciation. Par exemple, certains mots possèdent des variantes phonétiques libres, comme le verbe *lier* qui peut se prononcer en une syllabe ou en deux syllabes :

lier ?*[lie] [lije] [lje]

Certains éléments morphologiques offrent des différences de prononciation suivant le contexte phonémique. C'est le cas de l'élément final *-ien*, qui, suivant les mots, se prononce soit obligatoirement en une syllabe, soit obligatoirement en deux syllabes :

<i>italien</i>	*[italijɛ]	[italjɛ]
<i>ombrien</i>	[ɔbrijɛ]	*[ɔbrjɛ]

1. Différences d'emploi des variantes phonétiques

Ces trois exemples permettent déjà de faire une distinction importante. Les variantes monosyllabique et dissyllabique de l'infinitif *lier* s'emploient dans des conditions à peu près équivalentes, et on peut les considérer comme des variantes libres, c'est-à-dire interchangeables. Au contraire, les variantes monosyllabique et dissyllabique du suffixe *-ien* s'excluent mutuellement : elles sont en distribution complémentaire. On parle alors de variation conditionnée, car les différentes variantes n'ont pas les mêmes conditions d'emploi. Une variation phonétique est rarement complètement libre : en général, les variantes ne sont pas équivalentes.

Différences stylistiques, géographiques et sociales

Parfois, la différence de prononciation est associée à une différence de niveau de langue ou à une nuance de style. Ainsi la conjonction *si*, employée immédiatement devant le pronom *il*, admet une variante d'un niveau légèrement recherché :

Luc sait (si il, s'il) viendra

La phrase figée *il y a*, lorsqu'elle est conjuguée à un temps fini, admet au moins trois prononciations différentes, dont deux d'emploi familier :

On dit qu'il y (a, aurait, eut) du vent [ilj] [ij] [j]

Dans la phrase suivante, la particule préverbale (*Ppv*) *y* peut soit constituer une syllabe distincte, soit former avec le [a] qui suit une syllabe unique :

Jean y a mis un cendrier [ja] [ia]

et ce choix est lui aussi lié à une nuance dans le style de l'élocution. Une différence de prononciation peut aussi être associée à une différence géographique ou sociale : certains parlers français sont caractérisés par des [u] fricatifs uvulaires, d'autres par des

[r] roulés apicaux¹. Le mot *oxygène* peut se prononcer en [gz] en français du Québec, alors qu'il est obligatoirement en [ks] en français standard.

Dans une description formelle, les différences d'emploi dans l'échelle des niveaux de langue ou dans le style, et les différences sociales ou géographiques, peuvent *a priori* être prises en compte de deux manières. On peut séparer des façons de parler présentant une différence de ce type, pour les décrire comme des idiomes distincts. On peut, au contraire, décrire indistinctement les formes relevant de ces différentes façons de parler, en les considérant comme des variantes libres. Bien sûr, cette dernière attitude se justifie particulièrement lorsque les variations d'emploi qu'elle néglige sont fines. C'est cette attitude que nous choisissons dans un premier temps puisque ces façons de parler présentent des différences stylistiques, sociales ou géographiques difficiles à distinguer les unes des autres d'une manière reproductible. Nous écartons seulement les façons de parler vieilles ou marginales du point de vue stylistique ou social, comme la diction poétique traditionnelle. Pour ce qui est des variations géographiques, nous écartons les façons de parler trop spécifiques à des régions particulières². En dehors de ces cas extrêmes, nous considérons comme équivalentes les variantes phonétiques qui ne diffèrent que d'un point de vue stylistique, géographique ou social.

Différences dans le contexte phonémique

Les différences d'emploi entre les variantes phonétiques d'un élément concernent parfois le contexte phonémique amené par la flexion, par la dérivation ou par l'enchaînement des mots dans le discours. Par exemple, les prononciations [pli] et [plij] du radical du verbe *plier* alternent au cours de la conjugaison suivant que le suffixe est nul, commence par une voyelle ou commence par une consonne :

Luc <i>pliait</i> la carte	*[plie] [plije]
Luc (<i>plie, pliera</i>) la carte	[pli] *[plij]

L'élément final *-ien* associé à des noms de lieux ou de personnes se prononce [j] ou [ij] suivant le groupe consonantique qui le précède immédiatement :

<i>italien</i>	*[italijɛ] [italjɛ]
<i>ombrien</i>	[ɔbrijɛ] *[ɔbrjɛ]

L'article *les* se prononce avec ou sans [z] suivant l'initiale vocalique ou consonantique du mot qui le suit :

<i>Les chats dorment</i>	[le] *[lez]
<i>Les enfants dorment</i>	*[le] [lez]

¹ Les variétés de français étudiées ici utilisent le [v] fricatif uvulaire et non le [r] roulé apical, mais pour des raisons typographiques, nous le notons [r].

² Cette délimitation de l'objet décrit ne découle pas d'un jugement de valeur et ne reflète pas une intention normative. Ce sont des conditions qui délimitent *a priori* l'étude de façon explicite mais qui peuvent être modifiées à volonté selon les développements théoriques ou des considérations pratiques.

Différences syntaxiques et prosodiques

Nous avons évoqué la notion de prosodie dans le chapitre 1, section 4, p. 19, et rappelé que la structure syntaxique et la structure prosodique d'une phrase sont liées. Elles interviennent parfois dans les conditions d'emploi des variantes phonétiques d'un mot ou d'un élément donné. Ainsi, on sait que les *Ppv* sont atones sauf quand elles sont placées après le verbe, ce qui ne se produit qu'à l'impératif :

<i>Jean y a mis un cendrier</i>	[ja] [ia]
<i>Mets-y un cendrier</i>	*[zɛ] [ziɛ]

Suivant le cas, la *Ppv* y peut ou ne peut pas se prononcer [j]. Il y a donc une nette corrélation entre la place de la *Ppv* y dans la structure syntaxique de la phrase, sa place dans la structure prosodique, et l'acceptabilité de la prononciation [j].

Différences à plusieurs niveaux

Enfin, ces différents types de conditions d'emploi - 1° stylistiques, géographiques et sociales, 2° phonémiques, 3° syntaxiques et prosodiques - peuvent se combiner pour définir la distribution des variantes phonétiques d'un mot ou d'un élément de mot. Les trois variantes phonétiques du mot *six*, par exemple, alternent en fonction de deux facteurs. Le premier facteur, de nature phonémique, est l'initiale vocalique ou consonantique du mot qui suit :

<i>Luc a six billes</i>	[si] *[siz] *[sis]
<i>Luc a six ans</i>	*[si] [siz] *[sis]

L'autre facteur est syntaxique et prosodique. La prononciation de *six* diffère selon que le nom quantifié suit le mot *six*, ou bien est effacé ou situé dans un autre groupe nominal :

<i>Luc a six billes à la main</i>	[si] *[siz] *[sis]
<i>Luc en a six à la main</i>	*[si] *[siz] [sis]

Suivant le cas, le mot *six* occupe une place nettement différente dans la structure prosodique de la phrase.

Parenté entre variantes phonétiques

Nous ne parlons de variations phonétiques que si les différentes variantes sont apparentées. Dans tous les cas précédents, les variantes phonétiques restent globalement apparentées malgré leurs différences d'emploi. En revanche, la différence entre les mots *piller* et *pied*, par exemple, ne constitue pas une variation phonétique, puisque ces deux mots n'ont rien à voir par ailleurs. Dans d'autres exemples, la notion de parenté entre deux formes phonétiques est moins claire. En particulier, cette parenté peut sembler plus ou moins marquée selon les niveaux d'analyse. Ainsi, le verbe *resurgir* et le nom *résurgence*, malgré une certaine ressemblance morphologique, et donc phonémique, présentent une telle différence d'emploi du point de vue syntactico-sémantique que les phrases suivantes ne sont probablement pas dans une relation de

nominalisation :

La rivière resurgit à 2 km

La rivière a une résurgence à 2 km

et qu'on doit probablement les considérer comme relevant d'éléments lexicaux sans relation entre eux aujourd'hui. Le cas inverse se produit avec les paires *oser* et *audace*, ou *ressusciter* et *résurrection*, qui présentent des correspondances d'emploi remarquables, mais une importante dissemblance phonétique, au point que si on met en relation *Guy a osé crier* et *Guy a eu l'audace de crier*, on peut se poser la question de lier *Guy mange* et *Guy fait un repas*. Ces exemples difficiles mettent en évidence une des limites de la notion de relation entre éléments linguistiques : cette notion repose sur une analyse globale, et elle met en jeu plusieurs niveaux de l'analyse linguistique. En particulier, la description phonémique n'est pas séparable de la description syntactico-sémantique : par exemple, les notions morphologiques de suffixe, de préfixe et de dérivation interviennent dans ces deux domaines.

En raison de la difficulté et de la complexité de l'analyse linguistique, la notion de parenté entre formes phonétiques reste à définir plus clairement. Il s'agit là d'un problème théorique important, car la notion de variation phonétique repose inévitablement sur celle de parenté entre formes phonétiques. Ce problème s'est posé constamment et en termes concrets au cours du travail de description et de formalisation qui fait l'objet de ce chapitre : deux formes phonétiques sont-elles des variantes l'une de l'autre ou sont-elles sans relation formalisable ? C'est en analysant chaque situation difficile, cas par cas, que nous avons donné des solutions locales à ce problème.

En résumé, une relation entre deux variantes phonétiques peut être décrite comme une relation d'équivalence, à ceci près que les variantes ont parfois des conditions d'emploi différentes, voire complémentaires. Dans cette étude, nous prenons en compte les conditions d'emploi phonémiques et syntaxiques, car elles se prêtent bien à l'observation, mais nous ne considérons pas systématiquement les conditions stylistiques, géographiques ou sociales.

En ce qui concerne la distinction entre variations libres et variations conditionnées, nous pouvons déjà faire une remarque générale. Pour représenter une variation libre, il suffit d'un système formel capable de fournir la liste des variantes libres de chaque mot concerné. La représentation d'une variation conditionnée met en jeu un élément supplémentaire : la spécification des conditions d'emploi de chaque variante.

2. Effacement de consonnes finales

On observe dans de nombreux mots, au cours de la flexion et de la dérivation, une alternance entre un radical terminé par une consonne, comme *plate*, et un radical dans lequel cette consonne finale n'apparaît pas : *plat*. Ces variantes ont été appelées forme longue et forme courte du radical. Toutes les deux peuvent se rencontrer sans suffixe, mais les dérivés comportant un suffixe non nul, comme *platement*, ont

généralement la forme longue. En outre, les deux formes alternent au cours de la flexion : c'est le cas pour l'adjectif *plat*. Nous allons préciser l'extension lexicale de cette alternance.

(a) Un premier groupe d'exemples est constitué par des noms, des adjectifs et des participes variables en genre, comme l'adjectif *plat*. Dans ces mots, le masculin est constitué par la forme courte et le féminin par la forme longue. Toutefois, dans le cas des adjectifs susceptibles d'être placés avant le nom, il y a parfois des formes particulières pour l'adjectif au masculin singulier placé avant le nom et devant voyelle : on rencontre alors soit obligatoirement la forme courte, comme dans

un court extrait [kurekstre] *[kurtεkstre]

soit obligatoirement la forme longue, comme dans

un petit arbre *[ptiarbr] [ptitarbr]

soit obligatoirement une troisième forme, comme dans

un grand arbre *[grāarbr] *[grādarbr] [grātarbr]

soit une combinaison de ces possibilités, comme dans

un léger accent ?[leZeaksā] [leZεraksā] [leZeraksā]

(b) Un second groupe d'exemples est celui de verbes tels que *sentir*. Leur radical est sous sa forme longue à l'imparfait et aux temps morphologiquement apparentés :

il sentait ; nous sentons, vous sentez ; sentant ; que nous sentions, que vous sentiez ; sentons, sentez

ainsi qu'au *Vpr3pl* et aux formes qui en dérivent³ :

ils sentent ; qu'il sente

Le radical est sous sa forme courte au *Vpr1,2,3sg* : *il sent*, et au *Vimptf2sg* : *sens*. Les autres formes conjuguées montrent des comportements variables suivant les verbes. Ainsi, à l'infinitif, *sentir* a la consonne [t] mais *dire* n'a pas la consonne [z] de l'imparfait.

(c) Un troisième groupe d'exemples regroupe des noms de genre constant, comme *regard*, et des adverbes, comme *tard*, qui admettent des dérivés suffixés : respectivement *regarder* et *tarder* ou *tardif*. Le mot de base est constitué par la forme courte du radical et les dérivés sont construits sur la forme longue. Ces exemples sont particulièrement difficiles à recenser. En effet, les relations syntaxiques entre mots et mots dérivés sont complexes. Elles mettent en jeu, en fait, des relations transformationnelles entre phrases, telles que les nominalisations. Il est probable, par exemple, que *regard* et *regarder* sont liés par une nominalisation, et que *coup* et *couper*, qui l'ont été, ne le sont plus aujourd'hui, car leur relation s'est perdue avec l'évolution de la langue. Mais la plupart des exemples sont litigieux, car ces relations sont mal connues, et on ne sait pas quels critères formels employer. Dans l'état actuel des connaissances, on doit donc se contenter d'exemples plus sûrs : ceux qui mettent en jeu

³ La liste des abréviations utilisées pour exprimer les formes verbales figure dans le chapitre IV, section 2, p. 81.

la variation en genre et la conjugaison. C'est pourquoi les recensements qui suivent concernent seulement les exemples des groupes (a) et (b).

La classification en trois groupes (a), (b), (c) est fondée sur des considérations grammaticales et syntaxiques. Parmi tous ces exemples, la structure phonémique de la fin du radical permet de distinguer plusieurs situations et d'établir une seconde classification, que nous allons détailler en rappelant à quel groupe, (a) ou (b), se rattachent les exemples.

Dans des exemples tels que *plat*, *léger*, *sentir*, le radical se termine par une voyelle, orale ou nasale, éventuellement suivie de la consonne intermittente. Cette consonne est :

- [t] dans (a) 336 noms, 1372 adjectifs et 110 participes passés variables en genre (*plat*), ainsi que dans les participes présents de tous les verbes,
et (b) dans 43 verbes irréguliers (*mettre*),
- [z] dans (a) 251 noms, 727 adjectifs et 51 participes passés variables en genre (*creux*)
et (b) 62 verbes irréguliers (*dire*),
- [s] dans (a) 6 noms et 14 adjectifs variables en genre (*gros*)
et (b) 439 verbes (*finir*),
- [r] dans (a) 410 noms et 218 adjectifs variables en genre (*léger*),
- [d] dans (a) 37 noms et 43 adjectifs variables en genre (*grand*)
et (b) 41 verbes irréguliers (*prendre*),
- [v] dans (b) 35 verbes irréguliers (*écrire*),
- [k] dans (a) 1 nom et 1 adjectif variables en genre (*franc*)
et (b) 2 verbes irréguliers (*vaincre*),
- [g] dans (a) 3 adjectifs variables en genre (*long*),
- [l] dans (a) 1 adjectif variable en genre (*soûl*)
et (b) 5 verbes irréguliers (*vouloir*),
- [S] dans (a) 1 nom et 2 adjectifs variables en genre (*blanc*),
- [p] dans (b) 4 verbes irréguliers (*rompre*),
- [j] dans (a) 1 nom et 1 adjectif variable en genre (*gentil*)
et (b) 4 verbes irréguliers (*bouillir*).

Des échantillons des listes de mots correspondantes sont donnés dans l'annexe 6, dans le même ordre que ci-dessus. Les effectifs indiqués ont été calculés sur une version du DELAP qui comptait 58.000 mots. Les noms et les adjectifs qui se terminent par une même consonne intermittente sont donnés en deux listes : une liste de noms et une liste d'adjectifs. Les mots qui peuvent être soit nom soit adjectif sont donc listés et comptés deux fois, une fois en tant que nom et une fois en tant qu'adjectif.

Les autres consonnes [h Z f w] ne fournissent aucun exemple. La voyelle précédant la consonne intermittente subit parfois des alternances entre [e] et [ɛ], entre

[o] et [ɔ], ou entre [0] et [œ], également observées ailleurs :

<i>légère, léger</i>	cf. <i>il cède, céder</i>
<i>sotte, sot</i>	cf. <i>il note, noter</i>
<i>ils veulent, il veut</i>	cf. <i>il pleure, pleurer</i>

Toutefois, en français standard, ni l'alternance entre [o] et [ɔ] ni celle entre [0] et [œ] ne se produisent lorsque la consonne intermittente est [s] ou [z] :

<i>gros</i>	[gro] * [grɔ]	<i>creux</i>	[kr0] * [krœ]
<i>grosse</i>	[grɔs] * [grɔs]	<i>creuse</i>	[kr0z] * [kræz]

Dans des exemples tels que *vert* et *perdre*, le radical se termine par une voyelle orale suivie de [r] et éventuellement suivie de la consonne intermittente. Cette consonne est :

- [t] dans (a) 2 noms, 14 adjectifs et 11 participes passés variables en genre (vert),
- [s] dans (a) 1 nom et 11 adjectifs variables en genre (pervers),
- [d] dans (a) 114 noms et 107 adjectifs variables en genre (bavard)
- et (b) 12 verbes irréguliers (perdre),
- [m] dans (b) 5 verbes irréguliers (dormir),
- [v] dans (b) 3 verbes irréguliers (servir).

Dans des exemples tels que *bon* et *plein*, la consonne intermittente est [n]. La forme courte du radical se termine par une voyelle nasale : *bon*. La forme longue se termine par une voyelle orale suivie de [n] : *bonne*. Cette alternance introduit des relations entre certaines voyelles orales et les voyelles nasales :

- [o] et [ɔ] correspondent à [ɔ̃]
dans 63 noms et 51 adjectifs variables en genre (bon) ;
- [e] et [ɛ] correspondent à [ɛ̃]
dans 374 noms et 559 adjectifs variables en genre (plein) ;
- [a] correspond à [ã]
dans 28 noms et 35 adjectifs variables en genre (gitan) ;
- [i] correspond à [ĩ]
dans 73 noms et 138 adjectifs variables en genre (fin) ;
- [y] correspond à [ɥ̃] et à [œ̃]
dans 2 noms et 8 adjectifs variables en genre (commun).

Quelques exemples supplémentaires ressemblent à l'un des trois types ci-dessus, sans y correspondre en tous points :

<i>pouvoir</i>	formes longues : <i>pouvoir, ils peuvent</i> forme courte : <i>il peut</i>
<i>tenir</i>	formes longues : <i>tenir, ils tiennent</i> forme courte : <i>il tient</i>
<i>peindre</i>	forme longue : <i>ils peignent</i> forme courte : <i>il peint</i>
<i>malin</i>	formes longues : <i>maligne, maline</i> forme courte : <i>malin</i>
<i>suspect</i>	forme longue : <i>suspecte</i> forme courte : <i>suspect</i>
<i>oeuf</i>	forme longue : <i>oeuf</i> forme courte : <i>oeufs</i>

Enfin, l'alternance entre [v] et [f] observée par exemple entre *juive* et *juif* pourrait également être rapprochée de celle étudiée ici, car elle a lieu dans les mêmes conditions de flexion et de dérivation. Elle concerne quelque 450 noms et adjectifs variables en genre. Toutefois, nous nous en tiendrons aux exemples réguliers, et donc aux trois types dont nous avons donné des effectifs.

Même ainsi, l'alternance a une extension lexicale importante. Toutefois, dans le reste du lexique, on observe aussi de nombreux mots où elle ne se produit pas, bien que toutes les conditions flexionnelles et phonémiques énumérées ci-dessus soient remplies. Ainsi, il existe de nombreux mots dont le radical est constamment analogue à une forme courte

- du type *plat* - *léger* - *il sent* : par exemple *gai* et *il crée* ;
- du type *vert* - *il perd* : par exemple *fier* et *il court*.

Dans ces exemples, les conditions phonémiques de l'alternance sont donc remplies. Les conditions de flexion et de dérivation le sont aussi, puisqu'il s'agit soit de verbes, soit de noms, adjectifs et participes passés admettant les deux genres. Malgré cela, l'alternance ne se produit pas⁴.

Symétriquement, de nombreux mots ont un radical constamment analogue à une forme longue

- du type *plate* - *légère* - *ils sentent* : par exemple *vide* et *il gratte* ;
- du type *verte* - *ils perdent* : par exemple *corse* et *il avorte* ;
- du type *bonne* - *pleine* : par exemple *diaphane* et *il dîne*.

⁴ Pour ce qui est du type *bon* - *plein*, la situation est différente. Il existe bien des mots comme *flan* dont le radical est constamment analogue à une forme courte de ce type, c'est-à-dire terminée par une voyelle nasale. Mais aucun de ces mots n'admet de dérivés suffixés ; on ne trouve parmi eux aucun verbe ; et seuls une vingtaine d'entre eux sont des noms ou adjectifs des deux genres, comme *marron*. Les conditions de flexion et de dérivation, dans ces mots à voyelle nasale, ne sont donc remplies que dans quelques exemples.

La solution formelle adoptée pour rendre compte de ces alternances, depuis G.L. Trager (1944), est de représenter les deux formes du radical alternant à partir de la forme la plus informative, c'est-à-dire la forme longue. La consonne intermittente figure alors dans les représentations abstraites de l'une comme de l'autre. Pour retrouver la prononciation de la forme courte, on efface cette consonne, et aussi, dans le cas du type *bon - plein*, on remplace la voyelle orale par la voyelle nasale correspondante. Dans ce cas particulier, c'est encore la forme longue qui est plus informative que la forme courte, car la donnée d'une voyelle orale suffit à déterminer sans ambiguïté la voyelle nasale correspondante :

<i>pleine</i> [ɛn]	<i>plein</i> [ɛ̃]
<i>fine</i> [in]	<i>fin</i> [ɛ̃]
<i>commune</i> [yn]	<i>commun</i> [ɛ̃]
<i>bonne</i> [ɔn]	<i>bon</i> [ɔ̃]
<i>gitane</i> [an]	<i>gitan</i> [ã]

La réciproque est fautive : la voyelle nasale [ɛ̃] peut correspondre aux trois voyelles orales [ɛ], [i] et [y]. Il est donc devenu traditionnel, à juste titre, de représenter les deux formes du radical à partir de la forme longue. Pour que la forme longue et la forme courte reçoivent des représentations distinctes, il faut donc donner aussi une représentation phonémique aux conditions de l'alternance. En présence d'un suffixe, comme dans *platement*, on a vu qu'il ne peut s'agir que de la forme longue. Mais en fin de mot, on ne peut pas représenter indistinctement les deux formes par la forme longue, car cela introduirait une confusion entre les deux formes :

<i>plat</i>	/plat/	<i>plate</i>	/plat/
-------------	--------	--------------	--------

Pour éviter cette confusion, on doit introduire une marque abstraite soit dans les formes à consonne prononcée :

(1)	<i>plat</i>	/plat/	<i>plate</i>	/platə/
-----	-------------	--------	--------------	---------

soit dans les formes à consonne effacée :

(2)	<i>plat</i>	/plat*/	<i>plate</i>	/plat/
-----	-------------	---------	--------------	--------

La solution (1) est celle de l'orthographe : la marque attribuée aux formes longues est le *e* muet. Depuis G.L. Trager (1944), les phonologues reprennent traditionnellement cette solution et marquent d'un schwa les formes longues. Nous lui avons préféré la solution (2) et nous avons marqué du symbole /*/ les formes à consonne effacée. F. Dell et M. Plénat (1985) ont indépendamment adopté la même solution, qui a été implantée dans BDLEX. Pour notre part, ce choix vient de la prise en compte de l'ensemble du lexique. En effet, considérons les mots dans lesquels l'alternance n'a pas lieu bien que toutes les conditions soient remplies. Nous avons vu qu'on peut en distinguer deux types, suivant que le radical est constamment analogue à une forme courte, comme dans *gai*, *il crée*, *fier*, *il court*, *marron*, ou à une forme longue, comme dans *vide*, *il gratte*, *corse*, *il avorte*, *diaphane*, *il dîne*. Cette dernière série de mots est pertinente au problème, car lorsqu'on fait figurer la consonne effaçable dans les deux formes du type *plat - plate*, on les rapproche ainsi du type *vide*.

Prenons le cas de la solution (1) : les formes à consonne finale prononcée, comme *plate*, sont marquées d'un schwa qui les distingue des formes courtes correspondantes. Les mots du type *vide*, qui ont eux aussi une consonne finale prononcée, sont analogues aux formes longues telles que *plate*, et on doit également les marquer d'un schwa, sinon rien ne permettrait de les distinguer des formes courtes telles que *plat* :

Conditions de flexion	masculin ou <i>Vpr123sg</i>	féminin ou <i>Vpr3pl</i>
Type <i>plat</i> - <i>plate</i>	/plat/	/platə/
Type <i>vide</i>	/vidə/	/vidə/

On marque donc toutes les formes des deux types, sauf celles à consonne effacée.

Si on adopte la solution (2), ce sont les formes à consonne finale effacée, comme *plat*, qui sont marquées du symbole /*/. La représentation des mots du type *vide* peut rester inchangée :

Conditions de flexion	masculin ou <i>Vpr123sg</i>	féminin ou <i>Vpr3pl</i>
Type <i>plat</i> - <i>plate</i>	/plat*/	/plat/
Type <i>vide</i>	/vid/	/vid/

Les mots du type *vide* sont alors représentés d'un façon analogue aux formes longues telles que *plate*, ce qui reflète leur ressemblance phonétique.

Formellement, les deux solutions sont bien sûr équivalentes, mais la deuxième a des avantages secondaires. En effet, pour la plupart des consonnes, le type *vide* est numériquement beaucoup plus important que le type *plat* - *plate* ; pour la consonne [r], l'effacement ne peut avoir lieu qu'après [ɛ], et le type *mer* est beaucoup plus important que le type *léger* - *légère* ; enfin, pour les consonnes [b Z f w], le type *plat* - *plate* n'existe pas. Au niveau du lexique, les formes courtes telles que *plat* ont donc un caractère exceptionnel par rapport aux formes à consonne finale prononcée, qui regroupent aussi bien celles du type *plate* que celles du type *vide*. Il est donc économique de marquer ces formes exceptionnelles plutôt que les autres. De plus, comme il s'agit de représenter par un système abstrait une alternance qui affecte les mots du type *plat* - *plate*, il est naturel de marquer ces mots plutôt que ceux du type *vide*, qui, eux, ne participent pas à l'alternance. Ces arguments ne sont pas décisifs : ils n'impliquent pas qu'on doive écarter définitivement la solution calquée sur l'orthographe. Ils désignent la solution (2), non pas comme plus informative ou plus puissante, puisqu'elles sont équivalentes, mais comme plus logique et plus économique. C'est celle que nous avons adoptée. Nous donnons donc à chaque forme courte une représentation abstraite constituée de la forme longue suivie de la marque /*/ :

<i>plat</i>	/plat*/
<i>léger</i>	/leZer*/
<i>il perd</i>	/perd*/
<i>il finit</i>	/finis*/
<i>vert</i>	/vert*/
<i>bon</i>	/bon*/
<i>plein</i>	/plen*/

Les représentations des autres formes ne sont pas affectées par ce marquage, qu'il s'agisse des formes longues correspondantes :

<i>plate</i>	/plat/
<i>légère</i>	/leZer/
<i>ils perdent</i>	/perd/
<i>ils finissent</i>	/finis/
<i>verte</i>	/vert/
<i>bonne</i>	/bon/
<i>pleine</i>	/plen/

des mots dont le radical est constamment analogue à une forme longue :

<i>vide</i>	/vid/
<i>il gratte</i>	/grat/
<i>corse</i>	/kors/
<i>il avorte</i>	/avort/
<i>diaphane</i>	/diafan/
<i>il dîne</i>	/din/

ou des mots dont le radical est constamment analogue à une forme courte :

<i>gai</i>	/ge/
<i>il crée</i>	/kre/
<i>fier</i>	/fier/
<i>il court</i>	/kur/
<i>marron</i>	/marö/ ou /maron*/

3. Synérèse et diérèse : généralités

Aux trois voyelles les plus fermées du français, [i u y], que nous appellerons les voyelles fermées, correspondent les trois consonnes [j w ɥ], dites semi-consonnes. On entend par exemple la voyelle [y] dans le mot *nu* [ny] et la semi-consonne [ɥ] dans *nuît* [nɥi]. Chaque semi-consonne a le même timbre que la voyelle fermée correspondante, mais constitue une transition d'un son à un autre, sans passage par un stade stationnaire, alors que la voyelle comprend dans sa durée une partie stable. La suite de ce chapitre est consacrée à l'étude des variations phonétiques entre voyelles fermées et semi-consonnes en français. Ces variations constituent un exemple intéressant car elles sont nombreuses et se produisent dans des situations variées.

En français, les trois voyelles fermées et les trois semi-consonnes obéissent à une distribution complexe. Une semi-consonne est souvent suivie d'une voyelle : [wa] dans *loi*, [ɥi] dans *lui*, [je] dans *pied* sont des séquences courantes. Dans ces exemples, la semi-consonne et la voyelle appartiennent à une seule et même syllabe. Les trois mots cités sont donc monosyllabiques. En phonétique et en termes de versification, on désigne cette situation sous le nom de synérèse, par opposition à une autre possibilité : dans certaines conditions, on rencontre une voyelle fermée dans la position qu'occupe la semi-consonne dans les exemples précédents, c'est-à-dire devant une voyelle. On a par

exemple les séquences phonétiques [ue] dans *clouer*, [yã] dans *truand*, [io] dans *lui aussi*. Les deux voyelles en contact forment alors un hiatus phonétique et appartiennent à deux syllabes distinctes. Pour opposer cette configuration à la précédente, on parle de diérèse. La séquence [ije] de *crier* représente une troisième situation. Les deux voyelles, [i] et [e], appartiennent comme précédemment à deux syllabes distinctes, mais ici la transition entre les deux voyelles prend la forme de la semi-consonne [j], qu'on perçoit d'ailleurs nettement à l'oreille : le [j] de *crier* est aussi net que celui de *briller*. Comme les deux voyelles appartiennent à deux syllabes distinctes, on parle également de diérèse. Le contraste entre les séquences [io] dans *lui aussi* et [ije] dans *crier* montre donc qu'il existe deux types de diérèse en [i], suivant qu'on prononce ou non un [j] entre le [i] et la voyelle qui suit. La différence entre ces deux sortes de diérèse, bien perceptible sur ces deux exemples, est cependant moins claire sur d'autres. Curieusement, on ne retrouve pas une distinction analogue pour les diérèses en [u] et en [y]. En effet, on peut parfois déceler un léger [w] à la transition entre les deux syllabes de *clouer*, ou un léger [ɥ] entre celles de *truand*, mais ce [w] et ce [ɥ] sont peu nets à l'oreille, alors que certaines combinaisons de mots amènent des séquences telles que [uwa], dans *canard ou oie*, et [ɥɥi], dans *la dame du huit*, dans lesquelles la transition entre les deux syllabes est sans aucun doute une semi-consonne. Conformément à la tradition des phonéticiens français, et à la suite de F. Dell et M. Plénat (1985), mais à l'encontre de l'usage en phonologie générative, nous négligeons le [w] éventuel de *clouer* et nous notons ce mot [klue]. De même, nous notons [yã] la diérèse de *truand*.

Les exemples ci-dessus se classent donc en trois types.

- Synérèse : *loi* [lwa], *lui* [lɥi], *piéd* [pje].
- Diérèse avec hiatus phonétique : *clouer* [klue], *truand* [tryã], *lui aussi* [io].
- Diérèse avec [j] : *crier* [krije].

Les différences entre ces trois types sont parfois des nuances difficiles à observer, que ce soit à l'oreille ou sur des spectrogrammes. En effet, la différence entre synérèse et diérèse concerne le nombre de syllabes, or le nombre de syllabes d'un énoncé n'est pas toujours net. De même, l'observation d'une transition en [j] entre [i] et voyelle est parfois douteuse, car tous les intermédiaires existent entre un hiatus tel que [io] dans *lui aussi* et une transition en [j] telle que [ije] dans *crier*. Ainsi, entre les prononciations [lje] et [lje] du verbe *lier*, qui sont toutes les deux employées, tous les intermédiaires semblent également acceptables. On peut alors parler d'une variation phonétique continue entre [lje] et [lje]. Un cas voisin est celui de l'infinitif *clouer*, généralement prononcé [klue]. Pour cette forme, la prononciation [kluwe] est inusitée, mais on observe des prononciations intermédiaires entre [klue] et [kluwe], c'est-à-dire avec un [w] léger.

Ces prononciations intermédiaires posent un problème théorique et méthodologique : il existe souvent des formes phonétiques intermédiaires entre deux variantes phonétiques, mais elles sont difficiles à observer systématiquement, car les nuances de prononciation trop fines ne peuvent être reconnues à l'oreille d'une façon reproductible. Par rapport à l'observation à l'oreille, l'utilisation d'instruments d'analyse du signal de parole a des avantages et des inconvénients. Elle assure une meilleure reproductibilité des descriptions, mais elle implique le choix d'un corpus et elle ne

permet pas l'accès à des données aussi nombreuses et aussi variées, car seule l'observation à l'oreille est immédiate et se répète quotidiennement. Ces difficultés, celles qui sont liées à l'observation directe comme celles qui résultent de l'utilisation d'instruments, expliquent que les observations phonétiques ne puissent pas à la fois être fiables et précises, et couvrir un champ d'investigation étendu. Et effectivement, les observations phonétiques ne sont pas absolument fiables, même lorsqu'il y a un consensus entre les auteurs. Nous choisissons l'observation à l'oreille, ce qui permet de faire des observations immédiates, donc nombreuses et variées, et aussi d'émettre des jugements d'acceptabilité à volonté sur n'importe quelle prononciation, sans attendre de la voir apparaître dans un corpus. En revanche, pour que ces observations soient suffisamment reproductibles, nous limitons leur précision *a priori*, en excluant les prononciations intermédiaires telles que celles qui comprennent un [w] léger ou celles pour lesquelles le nombre de syllabes phonétiques n'est pas net. Seules les prononciations extrêmes sont donc prises en compte et jugées. Ainsi, nous utilisons des observations aussi reproductibles que possible, et tout de même suffisamment précises pour distinguer les différents types de comportements dans le lexique et pour classer les mots correspondants.

Le français offre de nombreux exemples de variations entre synérèse et diérèse. Ces variations font intervenir plusieurs facteurs. Dans les pages qui suivent, nous décrivons ces variations et nous discutons de leur représentation formelle. Tout d'abord, l'étude des types *louer*, *tuer* et *lier* fournit des exemples de variations libres entre synérèse et diérèse. Dans une deuxième partie, avec l'analyse des types *manier*, *il manie* et *plier*, nous rencontrerons des variations entre synérèse et diérèse conditionnées par le contexte phonémique. Enfin, l'inventaire des variations entre synérèse et diérèse à la limite des mots illustrera le rôle joué par la syntaxe. Ces trois études montreront en outre l'importance des facteurs lexicaux, c'est-à-dire la diversité des comportements phonétiques en fonction des mots envisagés.

4. Variations libres entre synérèse et diérèse : les types *louer*, *tuer*, *lier*

G. Gougenheim (1935, p. 27) évoque les variations libres entre synérèse et diérèse en français : "il y a hésitation entre les variantes voyelle et voyelle-consonne, lorsque la variante voyelle existe (...) dans d'autres formes du mot ou dans des mots de la même famille. Les habitudes de prononciation individuelle et le rythme plus ou moins accéléré de la phrase jouent un grand rôle dans le choix entre les deux variantes. (...) Tel est le cas de *lier* (*lié*), d'après *il lie* ; *nier* (*nié*), d'après *il nie*, en face de *nièce* (*niès*).". Nous décrivons l'extension lexicale de cette variation et nous en proposons une représentation formelle en conséquence. Dans un premier temps, nous considérons séparément les faits relatifs à [u] dans *louer*, à [y] dans *tuer* et à [i] dans *lier*.

4. a) Le type *louer*

Le verbe *louer*, indépendamment de la variété de ses emplois syntaxiques, se prononce soit en deux syllabes (diérèse), soit en une (synérèse) :

louer [lue] [lwe]

Les deux formes sont équivalentes. Comme cette situation se retrouve dans de nombreux mots, on peut formaliser en tant que telle la relation d'équivalence observée régulièrement entre la forme en [u] et celle en [w] :

(1) *louer* [lue] = [lwe]

Chaque classe d'équivalence est une paire telle que {[lue], [lwe]}.

L'étude du lexique permet de recenser les mots dans lesquels on observe cette variation, de préciser dans quelles conditions elle a lieu et de voir si elle est en corrélation avec d'autres faits. Une liste des mots caractérisés par une variation libre entre [u] et [w] est donnée dans l'annexe 7. Cette liste a été établie à l'aide du dictionnaire de groupes de consonnes évoqué dans le chapitre III, section 9, p. 69, et par une extraction automatique à partir du DELAP, suivie d'un codage manuel. Il en est de même des autres listes de l'annexe 7. Dans ces mots, [u] et [w] sont précédés d'une consonne et suivis d'une voyelle⁵. Celle-ci peut être l'une des voyelles [i e ε a ã y o œ ɔ̃] :

<i>jouir</i>	[ui]	=	[wi]
<i>louer</i>	[ue]	=	[we]
<i>mouette</i>	[uɛ]	=	[wɛ]
<i>rouage</i>	[ua]	=	[wa]
<i>louange</i>	[uã]	=	[wã]
<i>nouure</i>	[uy]	=	[wy]
<i>boueux</i>	[uø]	=	[wø]
<i>joueur</i>	[uœ]	=	[wœ]
<i>jouons</i>	[uɔ̃]	=	[wɔ̃]

Devant les autres voyelles [u o ɔ̃ ẽ], nous n'avons rencontré aucun exemple de l'alternance (1). Lorsque [u] ou [w] est précédé d'un groupe constitué d'une consonne obstruante (*Obs*⁶) et d'une consonne liquide (*Liq*⁷), on n'observe jamais cette équivalence entre synérèse et diérèse. On ne rencontre que des synérèses obligatoires et des diérèses obligatoires :

<i>croire</i>	*[ua] [wa]
<i>groin</i>	*[uẽ] [wẽ]
<i>clouage</i>	[ua] *[wa]

G. Gougenheim (1935, p. 27) notait que "devant *e* la variation extraphonologique a seulement un caractère individuel : pour le mot *groin* les deux prononciations *gruẽ* et *gruẽ* coexistent". Pour lui, même dans cette position, les deux variantes équivalentes existaient donc. Toutefois, le français standard actuel ne connaît pas [gruẽ] pour *groin*.

⁵ Il n'existe en français aucun exemple de mot comportant un [u] entre voyelles, même d'origine étrangère : *cacahuète*, *caoua*, etc., se prononcent avec [w].

⁶ *Obs* = : [p t k b d g f s v z ʒ].

⁷ *Liq* = : [l r].

Les conditions phonémiques de l'alternance entre [u] et [w] se résument donc ainsi : elle n'a lieu qu'après une consonne et devant une voyelle, et elle n'a jamais lieu après un groupe *Obs-Liq*.

Tant que ces conditions phonémiques restent remplies, la conjugaison des verbes ne modifie pas l'acceptabilité des variantes en [u] et en [w]. Par exemple, ces variantes restent toutes les deux acceptables dans toutes les formes conjuguées du verbe *louer* dont le suffixe de conjugaison commence par une voyelle :

<i>il louait</i>	[luɛ]	=	[lwɛ]
<i>en louant</i>	[luã]	=	[lwã]
<i>il loua</i>	[lua]	=	[lwa]

Il en est généralement de même pour les dérivations : celles-ci, dans la plupart des cas, ne modifient pas l'acceptabilité des variantes en [u] et en [w], tant que les conditions phonémiques ci-dessus restent remplies. Ainsi, parmi les dérivés du verbe *louer*, ceux qui ont le radical *lou-* présentent également l'alternance entre synérèse et diérèse :

<i>re</i> <u>louer</u> <i>sur</i> <u>louer</u> <i>sous</i> - <u>louer</u>	[luɛ]	=	[lwe]
<i>louange</i>	[luãZ]	=	[lwãZ]
<i>loueur</i>	[luœr]	=	[lwœr]
<i>louable</i> <i>louage</i>	[lua]	=	[lwa]

L'alternance observée dans *jouir* se retrouve dans *jouissance*, *réjouir*, *réjouissances*, mais semble plus difficile dans *jouissif* et *jouisseur*. Si ces deux mots sont bien des dérivés de *jouir*, la règle des dérivés n'est donc pas absolue.

On note par ailleurs que dans beaucoup de mots qui présentent cette alternance libre, les éléments équivalents [u] et [w] constituent la fin du radical et sont immédiatement suivis d'un suffixe de conjugaison ou de dérivation. La limite entre le radical et le suffixe coïncide alors avec le lieu où se produisent la synérèse et la diérèse. C'est le cas pour le verbe *louer* : ce mot s'analyse en *lou-er*. De telles corrélations entre des phénomènes phonétiques et la structure morphologique des mots ont été remarquées depuis longtemps ; N. Chomsky et M. Halle (1968) ont élaboré des solutions formelles pour les expliciter. Toutefois, la notion de suffixe n'est pas nettement définie. Dans *lou-er*, il est clair que la finale [e] est un suffixe de conjugaison, mais dans le cas général, la présence d'un suffixe ne se déduit pas automatiquement de la forme du mot ni d'autres informations formelles. De nombreux mots ont un élément final phonétiquement semblable à un suffixe tel que *-isme*, *-iste*, *-eur*, mais manifestement sans valeur suffixale :

prisme triste bonheur

Dans des exemples tels que *mouette* et *rouage*, on ne voit pas d'arguments autres qu'historiques pour attribuer une valeur suffixale aux éléments finaux *-ette* et *-age*. Dans le cas du nom *bouée*, même les arguments historiques font défaut. Quant au verbe *jouir*, la valeur suffixale de l'élément final *-ir* y est également douteuse. La conjugaison donne la valeur d'un suffixe de conjugaison non pas à [ir] mais à [r], car le [i] se retrouve dans toutes les formes conjuguées et dérivées. Les relations entre *rougir* et *rouge*, et dans de nombreuses autres paires analogues, font de [ir] un suffixe de dérivation dans ces mots.

mais cet argument est sans valeur pour le verbe *jouer*. En conclusion, on constate que l'alternance (1) est corrélée à la valeur suffixale de l'élément final, mais que cette corrélation n'est pas parfaite.

Lorsque la synérèse et la diérèse alternent librement dans une forme comme *louer*, il est courant de rencontrer, dans une forme apparentée, un [u] suivi d'une consonne ou situé en fin de mot (G. Gougenheim, 1935, p. 27). La forme apparentée peut être, par exemple, une forme conjuguée du même verbe :

louer en face de *il loue, il louera*

Toutefois, si l'on cherche à caractériser les mots présentant l'alternance (1), cette propriété est aussi peu opératoire que la valeur suffixale de l'élément final. En effet, d'une part, la notion de formes apparentées est aussi mal définie que celle de suffixe, et pour les mêmes raisons. D'autre part, dans certains mots tels que *jouer*, la synérèse et la diérèse alternent librement sans qu'il existe de formes apparentées en [ʒu] dans lesquelles [u] serait suivi d'une consonne ou situé en fin de mot : le verbe *jouer* est sans relation, par exemple, avec les noms *joue* et *joug* ni avec la forme *il joue* du verbe *jouer*.

En revanche, on peut utiliser d'autres données pertinentes à l'étude de cette alternance : l'existence de mots qui présentent un comportement différent dans des conditions analogues. On constate que, dans une série de mots, la synérèse est obligatoire, bien que les conditions de l'alternance (1) soient remplies :

loi [lwa] * [lua]

En effet, de nombreux mots comportent un [w] obligatoire situé après une consonne et suivi d'une voyelle. Il s'agit de mots en [wa] ou en [wɛ], comme *loi* et *soin*, et de quelques mots en [we we wi wo wɔ] qui présentent des particularités historiques :

*quais bqesse bqette bqetter débqetter serfouette douelle douellière couenne
couenneux couette marquette qued seringero hquaiche huerta cuesta
rastaquère mqere zarzuela oui quighour quistiti bouif embabouiner ribouis
cambouis fouine fouiner fouinard fouineur fouir fouisseur enfouir
enfouissement enfouisseur serfouir serfouissage gouine baragouineur
baragouinage baragouiner couic couinement couiner malouine mouise
shampouineur shampouineuse shampouiner méchoui marsouiner vouivre
linguauz statu quo quantum*

Contrairement à ce qui se passait dans les exemples de l'alternance (1), l'élément final qui suit le [w] obligatoire, c'est-à-dire *-enne* dans *couenne*, n'a généralement pas de valeur suffixale. Les mots *fouir*, *lingual*, *malouin*, *serfouette* constituent peut-être quelques contre-exemples à cette règle. De même, les formes à [w] obligatoire n'ont pas de formes apparentées dans lesquelles [w] serait remplacé par un [u] suivi d'une consonne ou situé en fin de mot.

Dans certains de ces exemples, toutes les conditions phonémiques de l'équivalence (1) sont remplies⁸. L'existence de ces exemples implique que la

⁸ Il s'agit des mots à synérèse obligatoire où 1° [w] est précédé d'une consonne, où 2° il n'est pas précédé par un groupe *Obs-Liq.* et où 3° il est suivi par une des voyelles [i e ε a ā]. Le nom *loi* fait partie de cet ensemble.

distribution entre le type *louer*, avec alternance, et le type *loi*, avec [w] obligatoire, n'est pas conditionnée par le contexte phonémique. Ce résultat est en contradiction avec le sentiment qui prévalait avant qu'on ne dispose de listes d'exemples suffisamment représentatives : "il est des gens qui disent *pwa* et *pwi*, mais qui font de *bouée* et *muer* des dissyllabes ; comme toutefois l'emploi de *u* et *w* (...) dépend chez eux de la voyelle qui suit, ils ne peuvent, pas plus que les autres Français, utiliser fonctionnellement dans ce cas la distinction entre la voyelle et la semi-voyelle" (A. Martinet, 1933). Cette dépendance vis-à-vis de la voyelle qui suit n'a que quelques contre-exemples : en effet, devant la voyelle [a], le type *louer* est exceptionnel par rapport au type *loi*, et devant les voyelles [i e ɛ ā], ce sont les formes à [w] obligatoire qui sont rares par rapport aux représentants du type *louer*. C'est par l'exploitation d'un dictionnaire phonémique électronique que ces quelques contre-exemples ont été mis en évidence. En ce sens, nous parlons de l'exploitation d'un dictionnaire électronique à des fins théoriques (cf. chapitre II, section 6, p. 42). La distinction entre les types *louer* et *loi* a un autre intérêt. Elle offre l'exemple d'une différence phonétique à faible rendement. Le fait que la diérèse soit acceptable pour le premier type et interdite pour l'autre constitue une différence phonétique qui est partiellement redondante, puisqu'elle dépend, à quelques exemples près, de la voyelle qui suit. Il est à noter que les distinctions phonétiques à faible rendement ont été peu étudiées.

Considérons maintenant les faits observés après un groupe *Obs-Liq*. Nous avons vu qu'après *Obs-Liq* et devant *Voy*, on n'a jamais l'alternance entre [u] et [w] : il n'existe que des [u] obligatoires (*clouage*) et des [w] obligatoires (*gloire*). La distribution entre [u] et [w] dans ce contexte permet donc également de distinguer deux types d'exemples. Le premier type est constitué par des mots en *Obs-Liq*-[wa] ou en *Obs-Liq*-[wĕ], dans lesquels [w] est obligatoire :

<i>gloire</i>	*[gluar] [glwar]
<i>trois</i>	*[trua] [trwa]
<i>groin</i>	*[gruĕ] [grwĕ]
<i>Floing</i>	*[fluĕ] [flwĕ]

Les exemples du second type comprennent une séquence *Obs-Liq*-[u]-*Voy*, dans laquelle la diérèse est obligatoire :

<i>éblouir</i>	[ebluir] *[eblwir]
<i>trouer</i>	[true] *[trwe]
<i>brouette</i>	[bruɛt] *[brwɛt]
<i>clouage</i>	[kluaZ] *[klwaZ]
<i>en clouant</i>	[kluā] *[klwā]
<i>Drouot</i>	[druo] *[drwo]
<i>clouons</i>	[kluō] *[klwō]

Dans les représentants de ce deuxième type, l'élément final qui suit *Obs-Liq*-[u], c'est-à-dire *-er* dans *trouer* et *-age* dans *clouage*, a souvent une valeur suffixale. Ces formes sont souvent en relation avec des formes dans lesquelles [u] est suivi d'une consonne ou situé en fin de mot : ainsi, *trouer* est en relation avec *il troue* et avec *faire un trou*. Au contraire, dans les représentants du type *gloire*, l'élément final n'a jamais de valeur

suffixale, et on ne trouve jamais de formes apparentées dans lesquelles [w] serait remplacé par un [u] suivi d'une consonne ou situé en fin de mot. La situation est donc parallèle à celle observée en l'absence d'un groupe *Obs-Liq*.

Le tableau 1 ci-dessous résume les propriétés des quatre types de mots que nous avons distingués : *louer*, *loi*, *trouer*, *gloire*. Chaque ligne de ce tableau représente donc non pas une seule entrée lexicale, mais un ensemble de mots qui correspondent à un contexte phonémique et à un comportement phonétique donnés. La ligne *louer* concerne tous les exemples de l'alternance (1) ; la ligne *loi*, les cas de synérèse obligatoire en l'absence d'un groupe *Obs-Liq* ; la ligne *trouer*, les cas de diérèse obligatoire ; et la ligne *gloire*, les cas de synérèse obligatoire après un groupe *Obs-Liq*. Des échantillons des quatre listes de mots correspondantes sont donnés dans l'annexe 7. Les colonnes marquées [u] et [w] donnent respectivement l'acceptabilité de la diérèse et de la synérèse.

Parmi les représentants du type *trouer*, ceux en *Obs-Liq*-[ua] s'opposent aux exemples en *Obs-Liq*-[wa] tels que *gloire* et montrent que la distribution entre ces deux types n'est pas conditionnée par le contexte phonémique. Il s'agit de quelques mots qui sont, mis à part le nom *brouhaha*, des formes de verbes en *-ouer* et des dérivés de ces verbes par un suffixe tel que *-age* :

Guy renfloua alors les finances d'Anne [ua] * [wa]

Le clouage des planches a duré une heure [ua] * [wa]

Dans ces formes, l'élément final *-a* ou *-age* est manifestement un suffixe de conjugaison ou de dérivation. Dans ce cas particulier, la distribution entre [u] et [w] est donc plus nettement corrélée avec la valeur suffixale de l'élément final.

Tableau 1. Synérèses et diérèses en [u] et [w]

Groupe <i>Obs-Liq</i> précédent	Exemple	[u]	[w]	Valeur suffixale de l'élément final	Formes apparentées en [u] suivi d'une Cons ou situé en fin de mot
-	louer	+	+	+	(généralement)
-	loi	-	+	-	(généralement)
+	trouer	+	-	+	(généralement)
+	gloire	-	+	-	(généralement)

Le tableau 1, p. 114, met en évidence les propriétés qui séparent les quatre types de mots. Le travail de description linguistique résumé dans ce tableau nous a permis de rendre explicites dans le dictionnaire les différences entre ces types, en intégrant les informations du tableau dans les transcriptions phonémiques des mots concernés. La variation libre entre [u] et [w] est ainsi spécifiée formellement, ce qui était notre objectif. Remarquons que la seule façon de spécifier cette variation libre était d'incorporer les informations adéquates dans les transcriptions phonémiques des mots. En particulier, l'orthographe n'en offre pas une spécification satisfaisante. Certes, les mots orthographiés en *-oi-*, *-qua-* et *-gua-* ne sont jamais sujets à la variation. Mais parmi les mots en *-ou-*, quelques-uns relèvent du type *loi*, comme *douane*, *bivouac*, *couard*, *couenne*, *fouine*, *oui...*, et les autres du type *louer*, l'orthographe ne permettant pas de distinguer entre eux. A cet argument s'en ajoute un autre, plus fondamental : la raison d'être des représentations phonémiques est de spécifier certaines propriétés, même si elles sont également spécifiées dans l'orthographe, et *a fortiori* si elles y sont spécifiées d'une manière insuffisante.

Nous allons donc voir comment donner une représentation formelle aux propriétés de variation étudiées. Il ressort du tableau 1 que la propriété de l'acceptabilité de la diérèse, donnée dans la colonne [u], rapproche les types *louer* et *trouer*, et les éloigne des types *loi* et *trois*. Les deux critères donnés dans les colonnes de droite sont la valeur suffixale de l'élément final et la présence d'un [u] suivi d'une consonne ou situé en fin de mot dans des formes apparentées. Ces deux critères corroborent la parenté mise en évidence par l'acceptabilité de la diérèse, même s'ils sont parfois mal définis, et même si leur corrélation avec la distribution de [u] et [w] est seulement approximative. L'acceptabilité de la diérèse semble donc être un bon critère pour classer tous ces mots.

De plus, la comparaison de la première colonne (présence d'un groupe *Obs-Liq* précédent) et de la colonne marquée [u] (acceptabilité de la diérèse) montre que ces deux propriétés suffisent à distinguer les quatre types de mots. Or, la présence d'un groupe *Obs-Liq* précédent est une propriété des représentations phonémiques des mots. Donc, si l'acceptabilité de la diérèse est elle aussi explicitée dans les représentations phonémiques, la classification en quatre types de mots sera déductible de l'examen de ces représentations. Cet argument, confirmé par l'étude du type *tuer*, nous a amené à expliciter dans les représentations phonémiques la propriété d'acceptabilité de la diérèse. Pour cela, nous avons introduit une marque abstraite, /+/, à l'endroit où la diérèse est acceptable, c'est-à-dire généralement à la limite entre le radical et le suffixe⁹ :

<i>louer</i>	/lu+e/	→	{[lue], [lue]}
<i>trouer</i>	/tru+e/	→	{[true]}

Avec l'introduction de cette information, l'acceptabilité de la synérèse, indiquée dans la colonne [w], devient prévisible en fonction de la représentation phonémique : la synérèse est interdite lorsqu'on a à la fois un groupe *Obs-Liq* et une marque /+/,

⁹ Les formes citées entre barres obliques ne sont pas des formes observées, mais des formes phonémiques, c'est-à-dire des formes abstraites qui symbolisent des formes observées.

comme dans *trouer* /tru+e/ ; elle est acceptable dans les trois autres cas, c'est-à-dire *louer*, *loi*, *gloire*. Dans les quatre types de mots représentés dans le tableau, la distribution de [u] et de [w] est alors conditionnée par le contexte phonémique. Ces deux phonèmes sont donc équivalents et peuvent être remplacés sans perte d'information par un phonème abstrait. Nous notons ce phonème /u/ plutôt que /w/, en raison des relations que les mots des types *louer* et *trouer* entretiennent avec des formes en [u] telles que *il loue* et *il troue*. Nous obtenons donc le système de codage suivant :

<i>louer</i>	/lu+e/	→	{[lue], [lwe]}
<i>trouer</i>	/tru+e/	→	{[true]}
<i>loi</i>	/lua/	→	{[lwa]}
<i>gloire</i>	/gluar/	→	{[glwar]}

Il est à noter que ce système de représentation n'emploie le phonème /w/ dans aucun des quatre types, puisque les [w] phonétiques sont symbolisés par le phonème /u/. L'emploi de la marque /+ / présente l'avantage de pouvoir être étendu à des mots sans [u] ni [w] : notamment, la même marque /+ / peut être réutilisée pour les types *tuer* et *lier* qui ressemblent à *louer*.

La marque /+ / utilisée dans ce système est un phonème non prononcé. Traditionnellement, les phonologues se servent de phonèmes non prononcés pour représenter des frontières, comme la limite de mot (N. Chomsky et M. Halle, 1968), ou certaines propriétés combinatoires, comme le *h* aspiré en français. Nous reprenons cet usage pour expliciter l'acceptabilité de la diérèse, et aussi (cf. plus haut section 2) l'effacement d'une consonne finale. Une description phonémique, par définition, ne se limite pas à une description phonétique et inclut des informations sur les propriétés des mots. Il nous semble donc justifié de consacrer certains phonèmes uniquement à ce genre de fonction, ce qui conduit à la notion de phonème muet. Nous ne voyons donc pas d'obstacles, ni de nature théorique ni de nature empirique, à l'utilisation de phonèmes muets dans des représentations abstraites. Toutefois, pour certains auteurs, étant donné qu'une chaîne de phonèmes symbolise une chaîne parlée, chaque phonème doit symboliser un son prononcé, et la notion de phonème muet est abusivement abstraite. C'est peut-être en partie pour cette raison théorique, associée à des faits dialectaux et historiques, que S. Schane (1968) considère le schwa final comme une voyelle facultative, et non comme un phonème muet ; de même, F. Dell (1973) considère le *h* aspiré comme une représentation abstraite de l'occlusive glottale facultative [ʔ]. B. de Cornulier (1978) critique cette position de F. Dell. Dans le cadre des variétés de français prises en compte ici, le *e* dit "muet" final et le *h* aspiré ne semblent pas pouvoir être considérés autrement que comme des phonèmes muets.

Nous avons vu que le système de codage ci-dessus n'emploie pas le phonème /w/, puisque les [w] phonétiques y sont symbolisés par le phonème /u/. Si nous élargissons à nouveau notre champ d'investigation dans le lexique et si nous considérons les autres [u] et les autres [w] que ceux abordés jusqu'ici, nous constatons encore une distribution quasi-complémentaire entre [u] et [w]. Devant voyelle, on n'a que [w] :

<i>cacahuète</i>	{[kakawɛt]} *{[kakauɛt]}
<i>ouate</i>	{[wat]} *{[uat]}

Devant consonne et en fin de mot, on n'a que [u] :

caoutchouc *[kawtsu] [kautsu]
caillou *[kajw] [kaju]

Les seules exceptions sont des onomatopées et des mots d'origine étrangère :

miaou [mjau] [mjaw]
out ?[aut] [awt]

Ces cas mis à part, la distinction entre /u/ et /w/ apparaît comme redondante compte tenu de la distribution de ces phonèmes, ce qui justifie qu'on leur attribue une source unique :

cacahuète /kakauet/ ———> {[kawkɛt]}
loi /lua/ ———> {[lwa]}

Dans les mots tels que *miaou* et *out*, le phonème /w/ reste indispensable, mais leur petit nombre et leur marginalité autorisent à les considérer comme hors système, et à dire que la disparition de /w/ dans les autres mots réalise l'économie d'un phonème.

Cette solution permet de réduire la redondance des représentations phonémiques, au prix d'une abstraction accrue. Cependant, nous ne voyons pas là l'application d'un principe absolu. Ainsi, pour rendre explicite la différence entre les types *louer* et *loi*, l'emploi d'une marque abstraite /+/ n'est pas la seule solution envisageable¹⁰. Plus généralement, il ne semble pas indispensable d'éliminer toute redondance des formes phonémiques. En effet, la langue regorge de redondances plus ou moins complètes, qui ont pour fonction de corriger des erreurs de transmission ou de compenser l'existence d'ambiguïtés. Or l'histoire de la phonologie montre qu'on ne peut faire disparaître ces redondances de la description qu'au prix de systèmes extrêmement abstraits, inévitablement basés sur des choix parfois arbitraires. D'ailleurs, de tels choix, même sur des points de détail, ont fait l'objet de controverses et de changements de position radicaux de la part des chercheurs. C'est pourquoi nous tenons à reconnaître le caractère arbitraire de nombreux choix, pourtant nécessaires.

4. b) Le type *tuer*

Parallèlement à ce que nous avons constaté à propos du verbe *louer* et des mots du même type, le verbe *tuer* se prononce en une syllabe ou en deux syllabes :

tuer [tye] [tɥe]

Les deux formes sont équivalentes :

(2) *tuer* [tye] = [tɥe]

La synérèse et la diérèse alternent ainsi dans de nombreux mots où [y] et [ɥ] sont

¹⁰ On peut notamment représenter par /u/ le [u] et le [w] de *louer* et *trouer*, et par /w/ le [w] de *loi* et *gloire*. Cette solution est moins abstraite.

précédés d'une consonne et suivis d'une voyelle :

<i>perpétuité</i>	[yi]	=	[ɥi]
<i>tuer</i>	[ye]	=	[ɥe]
<i>duel</i>	[yɛ]	=	[ɥɛ]
<i>diminuendo</i>	[yɛ̃]	=	[ɥɛ̃]
<i>respectueux</i>	[y0]	=	[ɥ0]
<i>sueur</i>	[yœ]	=	[ɥœ]
<i>duo</i>	[yo]	=	[ɥo]
<i>tuons</i>	[yɔ̃]	=	[ɥɔ̃]
<i>ruade</i>	[ya]	=	[ɥa]
<i>en tuant</i>	[yã]	=	[ɥã]

Contrairement à ce qui se passait avec [u] et [w], la variation entre synérèse et diérèse peut aussi apparaître après un groupe *Obs-Liq*. Dans une douzaine de mots, *Obs-Liq-[yi]* semble alterner avec *Obs-Liq-[ɥi]* :

<i>altruisme</i>	[tryi]	=	?[trɥi]			
<i>fluide -ifier -ifiant -ification -ique -iser -ité superfluide</i>				[flyi]	=	?[flɥi]
<i>incongruité</i>	[gryi]	=	?[grɥi]			
<i>superfluité</i>	[flyi]	=	?[flɥi]			
<i>truisme</i>	[tryi]	=	?[trɥi]			

Mis à part la présence du groupe *Obs-Liq*, et une préférence pour la diérèse, ces mots ne semblent pas notablement différents des exemples tels que *tuer*.

Dans la plupart des mots qui présentent cette alternance libre, les éléments équivalents [y] et [ɥ] constituent la fin du radical et sont immédiatement suivis d'un suffixe de conjugaison ou de dérivation. La synérèse et la diérèse se produisent alors entre le radical et le suffixe. C'est le cas pour le verbe *tuer* qui s'analyse en *tu-er*. Dans ce mot, il est clair que la finale [e] est un suffixe de conjugaison. Mais dans d'autres exemples tels que le mot *fluide*, le caractère suffixal de l'élément terminal n'est pas clairement établi. Dans des emprunts étrangers tels que *duo* ou *diminuendo*, cela n'a guère de sens de parler de suffixes à propos des finales -o et -endo. La situation est donc remarquablement semblable à celle du type *louer* : l'alternance libre entre synérèse et diérèse est corrélée à la valeur suffixale de l'élément final, mais cette corrélation n'est pas parfaite. De même, lorsque [y] et [ɥ] alternent librement dans une forme donnée, il est courant de rencontrer, dans une forme apparentée, un [y] suivi d'une consonne ou situé en fin de mot :

tuer en face de *il tue, il tuera*

Cette règle admet encore plus de contre-exemples que la précédente : ainsi, *altruisme* et *duel*, qui admettent la synérèse comme la diérèse, ne sont pas en relation avec des formes en [altry] ou en [dy] dans lesquelles [y] serait suivi d'une consonne ou situé en fin de mot.

Pour suivre le même plan que dans l'étude du type *louer*, considérons maintenant les mots qui présentent un comportement différent dans des conditions analogues. On

constate que dans d'autres mots, la synérèse est obligatoire, bien que les conditions de l'équivalence (2) soient remplies :

lui [lɥi] * [lyi]

Ces mots comportent un [ɥ] obligatoire situé après une consonne et devant une voyelle. Il s'agit tout d'abord de nombreux mots en [ɥi], dans lesquels [ɥi] peut ou non être précédé d'un groupe *Obs-Liq* :

bruyère pluie truite [ɥi] * [yi]

lui puis suie [ɥi] * [yi]

A ces exemples, dont nous donnons une liste dans l'annexe 7, s'ajoutent quelques mots dans lesquels la synérèse se fait entre [ɥ] et l'une des voyelles [a e ε Ē] :

persuader -sion -sif, dissuader -sion -sif [ɥa] * [ya]

lingual linguatule marijuana quichua [ɥa] * [ya] [wa]

puéril -ement -isme -ité -cultrice -culture [ɥɛri] * [pyeri]

puerpéral huerta [ɥɛ] * [yɛ]

suint -er -ement, chuinter -ant -ement, quindécemvir, quinto [ɥĒ] * [yĒ]

Dans ces exemples, [y] n'est jamais précédé d'un groupe *Obs-Liq*. Un certain nombre d'entre eux sont d'origine savante ou étrangère.

Il est à noter que l'élément final commençant après le [ɥ] obligatoire, c'est-à-dire *-ite* dans *suite* et *-inter* dans *suinter*, n'a généralement pas de valeur suffixale. L'adjectif *lingual*, avec sa finale *-al*, constitue peut-être une exception à cette règle. Par ailleurs, les formes à synérèse obligatoire ne sont pas en relation avec des formes dans lesquelles [ɥ] serait remplacé par un [y] suivi d'une consonne ou situé en fin de mot, sauf dans quelques paires telles que *pluie* et *il a plu*.

Les effectifs des types *tuer* et *lui* montrent que la distinction entre ces deux types est approximativement corrélée à la voyelle qui suit. En effet, devant les voyelles [i] et [Ē], le type *tuer* est exceptionnel par rapport au type *lui* ; et devant les autres voyelles, ce sont les formes à synérèse obligatoire qui sont rares par rapport aux représentants du type *tuer*.

Nous avons déjà évoqué deux cas où [y] ou [ɥ] est précédé d'un groupe *Obs-Liq* : d'une part, des mots en *Obs-Liq-[yi]* et *Obs-Liq-[ɥi]*, qui présentent une variation libre entre ces deux formes (*truisme*), et d'autre part des mots en *Obs-Liq-[ɥi]* où la synérèse est obligatoire (*truite*). Dans ces exemples, la voyelle est [i]. Il existe un troisième cas où la voyelle n'est pas [i] et où on constate que la diérèse est obligatoire :

affluer [aflɥe] * [aflɥe]

Dans ces mots, l'élément final qui suit *Obs-Liq-[y]*, c'est-à-dire *-er* dans *affluer*, a souvent une valeur suffixale. De plus, ces formes sont souvent en relation avec des formes dans lesquelles [y] est suivi d'une consonne ou situé en fin de mot : ainsi, *affluer* est en relation avec *ils affluent*.

Le tableau 2, p. 120, résume les propriétés des différents types que nous avons rencontrés. Il semble plus complexe que celui consacré au type *louer*, car nous avons pris

un critère supplémentaire pour classer les mots : à chaque fois, nous avons considéré séparément les mots dans lesquels la voyelle qui suit est [i] et ceux dans lesquels cette voyelle n'est pas [i]. Ainsi, deux lignes différentes sont consacrées aux types *perpétuité* et *tuer*, la première ligne pour les exemples où la synérèse et la diérèse ont lieu avec [i], et la deuxième pour ceux où elles ont lieu avec une autre voyelle. On obtient ainsi sept types, au lieu de quatre dans le tableau précédent. Cependant, le tableau montre que ce critère n'est en général pas corrélé aux autres. Il n'est donc pas souhaitable de choisir la voyelle qui suit comme critère pour classer les exemples. En revanche, comme pour le type *louer*, l'acceptabilité de la diérèse, indiquée dans la colonne [y], est corrélée approximativement aux critères donnés dans les deux dernières colonnes. Elle rapproche les exemples *perpétuité*, *tuer*, *incongruité*, *affluer*, et les oppose à *lui*, *persuader*, *truite*. Comme dans le cas du tableau de la p. 114, l'acceptabilité de la diérèse semble être un bon critère de classification pour ces mots. De plus, ce critère, conjugué avec ceux des colonnes n° 1 (présence d'un groupe *Obs-Liq* précédent) et 5 (la voyelle qui suit est [i]), suffit à distinguer les sept types de mots représentés dans le tableau. Or, les critères indiqués par les colonnes n° 1 et 5 concernent des informations qui figurent dans les représentations phonémiques des mots. Donc, si l'acceptabilité de la diérèse est également explicitée dans ces représentations, la classification en sept types est déductible de ces représentations, et en particulier la distribution de [y] et [ɥ] dans ces mots est prévisible en fonction du contexte phonémique. La situation est donc semblable à ce qui se passait pour [u] et [w]. Comme dans les mots en [u] et [w], nous explicitons par la marque abstraite /+/- la propriété d'acceptabilité de la diérèse, et nous

Tableau 2. Synérèses et diérèses en [y] et [ɥ]

Groupe <i>Obs-Liq</i> précédent	Exemple	[y]	[ɥ]	Voy suivante =: [i]	Valeur suffixale de l'élément final	Formes apparentées en [y] suivi d'une <i>Cors</i> ou situé en fin de mot		
-	perpétuité	+	+	+	+	(généralement)	+	(généralement)
-	tuer	+	+	-	+	(généralement)	+	(généralement)
-	lui	-	+	+	-	(toujours)	-	(toujours)
-	persuader	-	+	-	-	(généralement)	-	(toujours)
+	incongruité	+	+	+	+	(généralement)	+	(généralement)
+	affluer	+	-	-	+	(généralement)	+	(généralement)
+	truite	-	+	+	-	(toujours)	-	(généralement)

symbolisons par le phonème /y/ les variantes phonétiques [y] et [ɥ] :

<i>perpétuité</i>	/perpety + ite/	→	{{pɛrpetyite}, {pɛrpetɥite}}
<i>tuer</i>	/ty + e/	→	{{tɥe}, [tɥe]}
<i>lui</i>	/lyi/	→	{{liɥi}}
<i>persuader</i>	/persyade/	→	{{pɛrsɥade}}
<i>incongruité</i>	/enkongry + ite/	→	{{ɛkɔgryite}, [ɛkɔgryite]}
<i>affluer</i>	/aɥly + e/	→	{{aɥlye}}
<i>truite</i>	/tryit/	→	{{tryit}}

La distribution de [y] et de [ɥ] dans ces mots est donc conditionnée par le contexte phonémique, dès lors qu'on inclut dans ce contexte le phonème muet /+/. Il en est de même dans le reste du lexique. En effet, mis à part les [y] et les [ɥ] déjà évoqués, on a une distribution simple. Devant voyelle, on ne rencontre que [ɥ] :

<i>huître</i>	[ɥitr]	*[yitr]
---------------	--------	---------

Devant consonne et en fin de mot, on n'a que [y] :

<i>ahurir</i>	[ayrir]	*[aɥrir]
<i>bahut</i>	[bay]	*[baɥ]

Cette distribution conditionnée par le contexte introduit une équivalence entre [y] et [ɥ] et justifie qu'on les représente indifféremment par le phonème /y/. On fait alors l'économie du phonème /ɥ/.

4. c) Le type *lier*

Dans une petite série de mots, dont le verbe *lier*, on observe des variations comparables à celles que subissent *louer* et *tuer*. Cette série de mots est donnée dans l'annexe 7. Comme dans *louer* et *tuer*, la synérèse et la diérèse sont toutes les deux acceptables, et les prononciations correspondantes sont équivalentes :

(3)	<i>lier</i>	[lije] = [lje]
-----	-------------	----------------

Lorsqu'on fait la diérèse, on prononce généralement un [j] de transition entre le [i] et la voyelle qui suit¹¹. Il s'agit donc d'une variation libre entre [ij] et [j]. Dans les exemples de l'équivalence (3), les éléments [ij] et [j] sont précédés d'une consonne et suivis d'une des voyelles [e ɛ a ǣ y o œ ɔ ɔ̃], mais jamais d'une des voyelles [i ɛ̃ u]. Les mots concernés sont, mis à part quelques mots comme *hier*, *rhyolithe*, *Ryad* ou *ria*¹², une quinzaine de verbes et certains de leurs dérivés :

fiancer fier skier lier surlier nier expier rire orienter sourire scier chier obvier

¹¹ Toutefois, ce [j] ne semble pas aussi nettement obligatoire qu'après un groupe *Obs-Liq* ou *Cons-[y]* : la prononciation [lie], avec diérèse mais sans [j], est peut-être utilisée pour *lier*, contrairement au cas de *plier* et *appuyer* qui se prononcent toujours [plije] et [apɥije].

¹² On observe aussi une variation entre [j] et [ij] au cours de la conjugaison des verbes dont le radical se termine par un groupe *Obs-Liq*, avec les désinences *-ions* et *-iez* de l'imparfait et du subjonctif :

Nous (savions, voulons) que vous redoubliez. [blije] [blje]

C'est le seul cas où on observe une équivalence comparable à celle de (3) après un groupe *Obs-Liq*.

Tant que les conditions phonémiques de l'équivalence (3) sont remplies, c'est-à-dire tant que les éléments [ij] et [j] sont précédés d'une consonne et suivis d'une voyelle, toutes les formes conjuguées de ces verbes admettent la synérèse comme la diérèse :

en riant, il liait, il nia...

Quant aux dérivés de ces verbes, ceux dont le suffixe de dérivation commence par une des voyelles énumérées ci-dessus remplissent les conditions phonémiques de l'équivalence (3). La plupart sont également des exemples de (3) :

<i>skieur</i>	[skijœr]	=	[skjœr]
<i>liant</i>	[lija]	=	[lja]
<i>expiation</i>	[ɛkspijasjɔ̃]	=	[ɛkspjasjɔ̃]

Certains mots sont morphologiquement liés à un verbe de la série ci-dessus, mais semblent se prononcer avec une synérèse obligatoire : c'est le cas du verbe *confier*. Toutefois, les autres niveaux de l'analyse linguistique montrent que la parenté entre les verbes *fier* et *confier*, qui a probablement existé, s'est perdue au cours de l'histoire de la langue. Si l'on emploie ces verbes dans des phrases simples, on constate que rien ne laisse supposer l'existence d'une relation de dérivation vivante entre ces phrases :

Guy se fie à son flair

Guy a confié à Luc qu'il est amoureux

Nous ne considérons donc pas ces mots comme des dérivés des verbes considérés.

Quoi qu'il en soit, la plupart des mots qui admettent la variation libre de (3) sont des formes de verbes et des dérivés de verbes, dans lesquels la limite entre le radical et le suffixe coïncide avec le lieu où se produisent la synérèse et la diérèse. C'est le cas du verbe *lier*, analysable en *li-er*. Toutefois, cette règle n'est pas absolue : le verbe *orienter* et l'adverbe *hier*, qui admettent la synérèse comme la diérèse, sont des contre-exemples, car ils ne s'analysent pas en *ori-enter* et *hi-er*. On peut seulement dire que l'équivalence (3) se produit surtout dans des mots où l'élément final est un suffixe.

De même que dans les types *louer* et *tuer*, les formes telles que *lier* sont souvent apparentées à des formes dans lesquelles [i] est suivi d'une consonne ou situé en fin de mot. Ces formes peuvent être, par exemple, des formes conjuguées d'un même verbe :

<i>lier</i>	en face de	<i>il lie, il liera</i>
<i>il riait</i>	en face de	<i>il rit, rire</i>

Cette règle admet également des contre-exemples : les verbes *fiancer* et *orienter* admettent la synérèse comme la diérèse, sans qu'on trouve des formes apparentées en [fi] ou en [ori] dans lesquelles [i] serait suivi d'une consonne ou situé en fin de mot.

Finalement, l'équivalence (3) présente des ressemblances avec (1) et (2), y compris par les conditions dans lesquelles elle se produit.

Nous avons vu que cette équivalence concerne un petit nombre de mots. Dans des conditions analogues, d'autres mots présentent des comportements différents. Comme nous l'avons fait pour *louer* et *tuer*, nous allons étudier ces comportements qui s'opposent à celui de *lier*.

Tout d'abord, dans de nombreux composés comme *semi-aride*, où un élément terminé par [i] est suivi d'un élément commençant par une voyelle, on fait obligatoirement une diérèse à la limite entre les deux éléments, et on n'entend pas nettement un [j] à la transition entre le [i] et la voyelle :

semi-aride [s0miarid] ?*[s0mijarid] *[s0mjarid]

Cette diérèse obligatoire est phonétiquement semblable à celle qui se produit souvent à la limite de deux mots :

Paris est là *[pariela] ?*[parijela] *[parjela]

Les composés à diérèse obligatoire sont nombreux. D'une part, ils comprennent des formes de la langue courante, constituées d'un préfixe productif et d'un mot commençant par une voyelle :

<i>anti-</i>	<i>anti-<u>q</u>tomique</i>	[ia] ?*[ija] *[ja]
<i>archi-</i>	<i>archi-<u>é</u>tourdi</i>	[ie] ?*[ije] *[je]
<i>semi-</i>	<i>semi-<u>a</u>ride</i>	[ia] ?*[ija] *[ja]
<i>demi-</i>	<i>demi-<u>a</u>bricot</i>	[ia] ?*[ija] *[ja]

D'autre part, la langue savante, technique ou scientifique, utilise de nombreux éléments terminés par [i] (A. Cailleux et J. Komorn, 1981 ; H. Cottez, 1985). Lorsque ces éléments se combinent avec un élément commençant par une voyelle, le composé qui en résulte a souvent une diérèse obligatoire, comme s'il s'agissait de deux mots :

<i>bi-</i>	<i>biunivoque</i>
<i>di-</i>	<i>dialcool</i>
<i>poly-</i>	<i>polyester</i>
<i>multi-</i>	<i>multiovulé</i>
<i>pluri-</i>	<i>pluriovulé</i>
<i>hémi-</i>	<i>hémièdre</i>
<i>uni-</i>	<i>uniovulé</i>
<i>quadri-</i>	<i>quadriennal</i>
<i>sesqui-</i>	<i>sesquiangle</i>
<i>amphi-</i>	<i>amphiarthrose</i>
<i>oxy-</i>	<i>oxyhémoglobine</i>
<i>équi-</i>	<i>équiangle</i>
<i>péri-</i>	<i>périarthrite</i>
<i>quasi-</i>	<i>quasi-aphasique</i>
<i>archi-</i>	<i>archiépiscopal</i>
<i>calci-</i>	<i>calciurie</i>
<i>rachi-</i>	<i>rachianalgésie</i>
<i>etc.</i>	

Enfin, les quelques autres cas de diérèse obligatoire sans [j] à la transition sont des emprunts étrangers ou dialectaux :

behaviorisme [biavjorism] ?*[bijavjorism] *[bjavjorism]

Tous ces mots à diérèse obligatoire ont été représentés comme de véritables mots composés, dans lesquels les différents éléments sont séparés par des limites de mots. L'introduction de ces limites de mots se justifie donc peut-être sur des bases morphosyntaxiques, mais surtout sur des bases phonémiques, ce qui est souvent le cas pour les expressions figées (L. Danlos, 1981).

Les composés savants dont nous avons parlé seraient à étudier plus précisément. En effet, d'une part, si la majorité d'entre eux font la diérèse obligatoire, quelques-uns font la synérèse obligatoire, bien qu'il s'agisse clairement de composés. C'est le cas de *centiare*, *dièdre* et *périoste*, qu'on peut opposer à *milliampère*, *dialcool* et *périarthrite* respectivement : les premiers se prononcent avec une synérèse et les seconds avec une diérèse. Par ailleurs, la prononciation des termes savants est peu normalisée, notamment dans le domaine médical, et l'usage hésite entre synérèse et diérèse pour certains mots tels que *pollakiurie*, *kaliémie*, *rachialgie*, *myalgie*, *myasthénie*, *biopsie*, ... Enfin, lorsqu'un mot courant est issu historiquement de la langue savante, il peut soit avoir gardé la diérèse obligatoire, comme *polyester*, soit admettre la synérèse en plus de la diérèse, comme dans le nom composé *la biennale de Venise*. Lorsque l'usage n'est pas fixé, que ce soit dans la langue savante ou dans la langue courante, il en résulte une variation libre entre synérèse et diérèse. Les mots concernés sont alors difficiles à distinguer de ceux du type *lier*. Toutefois, nous avons maintenu cette distinction. En effet, nous considérons que les variations phonétiques observées dans *rachialgie* et dans *la biennale de Venise* reflètent une hésitation entre deux séries de mots, les uns savants et à diérèse obligatoire, comme *rachianalgésique* et *biunivoque*, et les autres à synérèse obligatoire, comme *centiare*. Au contraire, les variations de *lier* ne semblent pas découler d'une hésitation entre deux types. Nous n'avons donc pas codé de la même façon les variations de *rachialgie* et celles de *lier*. Dans le premier cas, les deux variantes ont été représentées comme deux éléments lexicaux distincts reliés par une équivalence d'emploi :

1 = *rachialgie*, /raSi alZi/, .N21

2 = *rachialgie*, /raSialZi/, .N21

Une autre série de mots s'oppose aux représentants du type *lier* et compte des effectifs beaucoup plus nombreux que les types *lier* et *dialcool* réunis. Il s'agit des mots à synérèse obligatoire en [j], comme *pied*. En effet, les séquences *Cons-[j]-Voy*, avec [j] obligatoire, sont fréquentes en français. Dans ces séquences, le [j] n'est jamais précédé d'un groupe *Obs-Liq* ni d'un groupe *Cons-[y]*. La voyelle qui suit le [j] peut être n'importe quelle voyelle, sauf peut-être [i] :

[e]	<i>pied</i>	[pje] *[pije]
[ɛ]	<i>miette</i>	[mjɛt] *[mijɛt]
[ɛ̃]	<i>lien</i>	[ljɛ̃] *[lijɛ̃]
[a]	<i>liane</i>	[ljan] *[lijan]
[ã]	<i>science</i>	[sjãs] *[sijãs]
[y]	<i>sciure</i>	[sjyr] *[sijyr]
[o]	<i>mieux</i>	[mj0] *[mij0]
[œ]	<i>pieuvre</i>	[pjœvr] *[pijœvr]

[u]	<i>fuel</i>	[fjul] * [fijul]
[o]	<i>national</i>	[nasjonal] * [nasijonal]
[ɔ]	<i>fiote</i>	[fjɔl] * [fijɔl]
[ɔ̃]	<i>pion</i>	[pjɔ̃] * [pijɔ̃]

Dans certains de ces exemples, toutes les conditions phonémiques de la variation (3) sont réunies, mais cette variation n'a pas lieu. Certains ont des propriétés grammaticales et morphologiques semblables à celles des mots du type *lier*. C'est le cas de plusieurs dizaines de verbes en *-ier*, comme *manier*. Dans ces verbes, l'élément final [e] a une valeur suffixale, et certaines formes conjuguées comportent un [i] suivi d'une consonne ou situé en fin de mot :

manier en face de *il manie, il maniera*

Ces propriétés rapprochent *manier* de *lier*, mais la variation libre caractéristique de *lier* ne se retrouve pas dans *manier*. Voici d'autres exemples de mots à [j] obligatoire dont l'élément final a une valeur suffixale :

association -ation
vérifiable -able
mendiant -ant

Les mots suivants ont des formes apparentées qui comportent un [i] suivi par une consonne ou situé en fin de mot, ce qui les rapproche du type *lier*, mais ils ont un [j] obligatoire :

italien en face de *Italie*
expédier en face de *expéditeur, il expédie*

La distribution entre les types *lier*, avec variation libre, et *pied*, avec synérèse obligatoire, n'est donc pas corrélée à la valeur suffixale de l'élément final, ni à l'existence de formes apparentées comportant un [i] suivi par une consonne ou situé en fin de mot. Cette distribution semble d'ailleurs variable suivant les régions : notamment, le type *lier* est probablement plus étendu dans le lexique du français méridional qu'en français standard. Par exemple, le verbe *parier* se rattache probablement au type *lier* dans certaines variétés méridionales du français, et au type *pied* en français standard. Le verbe *concilier*, lui, semble prendre un [j] obligatoire dans un cas comme dans l'autre. Mentionnons pour mémoire que dans la diction poétique traditionnelle, la diérèse est courante, même dans les mots tels que *manier* et *nation* qui, en-dehors de la récitation des vers, se prononcent toujours avec une synérèse.

Le tableau 3, p. 126, reprend les propriétés des tableaux consacrés précédemment aux types *louer* et *tuer*, mais cette fois-ci pour les synérèses et diérèses en [j] et [ij]. Ce tableau montre des résultats assez différents de ceux des pp. 114 et 120. Dans ces derniers, on constatait une corrélation approximative entre l'acceptabilité de la diérèse et les deux dernières colonnes, qui indiquent respectivement la valeur suffixale de l'élément final et l'existence de formes apparentées comportant la voyelle fermée suivie d'une consonne ou située en fin de mot. Ce n'est pas le cas ici : la propriété d'acceptabilité de la diérèse n'est en corrélation avec aucune autre propriété, si ce n'est la présence d'un groupe *Obs-Liq* précédent. Ce critère semble donc jouer un rôle moins

important dans ce tableau que dans les deux précédents. Par ailleurs, ce tableau est compliqué par l'existence d'une relation entre le type *plier* et le type *pied*, qui ont des représentants nombreux. La relation entre ces types fait l'objet de la section suivante. La propriété d'acceptabilité de la diérèse, qui différencie les deux types en dépit de cette relation, ne semble pas un critère adéquat pour classer les mots traités dans le tableau 3. C'est pourquoi nous ne l'explicitons pas directement dans les représentations phonémiques. Nous réutilisons seulement la marque abstraite /+/, évoquée ci-dessus à propos de *louer* et *tuer*, et nous l'introduisons dans les exemples de la relation (3) :

$$\text{lier} \quad /li+e/ \quad = \quad \{[lje], [lje]\}$$

Ce mode de représentation reflète l'analogie remarquable qui existe entre les équivalences (1), (2) et (3).

La représentation formelle des types *pied* et *plier* est abordée dans la section suivante. Quant au type *vous redoublez*, caractérisé par la variation libre entre *Obs-Liq-[ij]-Voy* et *Obs-Liq-[j]-Voy*, il est à la fois peu représenté, puisqu'il ne concerne que quelques formes de quelques verbes, et isolé par rapport au reste du système, puisque les groupes *Obs-Liq-[j]-Voy* intérieurs n'existent en français que dans ces formes. C'est ce qui explique la difficulté à trouver une représentation formelle satisfaisante pour les formes de ce type. Pour l'instant, dans le DELAP, le type *vous redoublez* est confondu avec le type *plier*.

Tableau 3. Synérèses et diérèses en [ij] et [j]

Groupe Obs-Liq ou Cons-[y] précédent	Exemple	[ij]	[j]	Valeur suffixale de l'élément final	Formes apparentées en [j] suivi d'une Cons ou situé en fin de mot
- lier		+	+	+	+
- pied		-	+	variable	variable
+ vous redoublez		+	+	?	-
+ plier		+	-	variable	variable

5. Variations contextuelles entre synérèse et diérèse : les types *manier*, *il manie*, *plier*

Les relations qu'on observe entre [i], [j] et [ij] dans des mots comme *manier*, *plier*, *italien*, *ombrien*, offrent de nouveaux exemples de variations entre synérèse et diérèse.

5.a. Une relation entre [i] et [j] : *manier*

De nombreux mots offrent une variation systématique entre [i] et [j]. Ainsi, le verbe *manier* et ses dérivés constituent une série de formes dans lesquelles divers suffixes de conjugaison et de dérivation s'associent à un même radical. Ce radical se présente sous la forme de deux variantes [mani] et [manj], et on constate une distribution quasi-complémentaire de ces deux formes en fonction du suffixe. Si celui-ci est nul ou commence par une consonne, on emploie la forme en [i], et s'il commence par une voyelle, on emploie la forme en [j] :

<i>Guy (manie, maniera) bien ces pinc</i>	[mani] * [manj]
<i>Ces pinc</i> sont d'un <i>maniement</i> facile	[mani] * [manj]
<i>Guy maniait bien ces pinc</i>	*[mani] [manj] ?*[manij]
<i>Ces pinc</i> sont <i>maniables</i>	*[mani] [manj] ?[manij]

La distribution complémentaire n'est en défaut qu'avec les suffixes *-ions* et *-iez* de l'imparfait et du subjonctif :

<i>Nous (savions, voulons) que vous maniez bien ces pinc</i>	[manje] ?*[manije]
--	--------------------

En dehors de ces cas, les formes en [i] et en [j] entrent dans une relation d'équivalence :

(4)	<i>manier</i>	[manje]	=	<i>il manie</i>	[mani]
-----	---------------	---------	---	-----------------	--------

Pour noter cette équivalence, nous choisirons une représentation formelle commune aux deux variantes.

On observe cette variation entre [i] et [j], avec les mêmes conditions d'emploi, dans environ 250 verbes et leurs dérivés, ainsi que dans plusieurs dizaines d'adjectifs en *-ien*, *-iaque*, *-iard*, *-iatique* formés sur des noms propres tels que *Italie*, *Chili*, *Libye*, *Bosnie*, *Chamonix*, *Asie*. Dans tous ces exemples, le [i] ou le [j] est précédé d'une consonne simple ou d'un groupe de consonnes, mais jamais d'un groupe *Obs-Liq* ni d'un groupe *Cons-[y]* intérieurs au mot.

5.b. Une relation entre [ij] et [j] : *italien*, *ombrien*

Dans la prononciation des suffixes *-ien*, *-ienne* et *-ier*, *-ière*, [ij] alterne avec [j] suivant les mots auxquels on les applique :

<i>italien</i>	*[italijē] [italjē]		
<i>ombrien</i>	[ɔbrijē] *[ɔbrjē]		
<i>poirier</i>	*[pwarije] [pwarje]	<i>plombier</i>	*[plɔ̃bije] [plɔ̃bje]
<i>néflier</i>	[nefije] *[nefje]	<i>plâtrier</i>	[platrije] *[platrje]

Lorsque le suffixe est précédé d'un groupe *Obs-Liq*¹³, la forme en [ij] est obligatoire, alors que dans les autres contextes, on prononce toujours [j]. Dans ces mots, il y a équivalence entre [ije] et [je], ou entre [ijĕ] et [jĕ] :

(5) (poir)ier [je] = (néfl)ier [ije]

Les conditions d'emploi des variantes prennent la forme d'une distribution complémentaire en fonction des caractéristiques phonémiques du radical¹⁴.

Les conditions d'emploi de la forme en [j] dans (5) (*poirier*) correspondent à celles de la forme en [j] dans la relation (4) étudiée en 5a, p. 127 (*manier*). Dans les deux cas, [j] est précédé soit par une consonne, soit par une voyelle, mais pas par un groupe intérieur *Obs-Liq* ou *Cons-[y]* ; il est suivi d'une voyelle. Un autre fait permet de rapprocher les relations (4) et (5) : dans quelques exemples, tels que certains adjectifs en *-ien*, cet élément est de plus en relation avec un élément *-ie*. Cela donne un troisième terme à la relation :

(5) (ital)ien [jĕ] ————— [ijĕ] (ombr)ien
 (4) /
 (Ital)ie [i]
 (Ombr)ie [i]

Le nom en *-ie* est en relation avec une forme en [jĕ] ou avec une forme en [ijĕ], suivant que le radical se termine en un groupe *Obs-Liq* ou *Cons-[y]*, ou non. La relation entre [i] et [j] dans *Italie* et *italien* s'identifie à la relation (4) entre *maniera* et *manier*. La relation entre [i] et [ij] dans *Ombrie* et *ombrien* a également une certaine généralité. On la retrouve notamment dans la conjugaison du verbe *plier*.

5.c. Une relation entre [ij] et [i] : *plier*

De la même manière que dans le type *manier*, on observe dans de nombreux mots une variation systématique entre [i] et [ij] après un groupe intérieur *Obs-Liq* ou *Cons-[y]*. Ainsi, dans les formes conjuguées et les dérivés du verbe *plier*, le radical auquel s'ajoutent les suffixes prend deux formes [pli] et [plij] qui se distribuent suivant les mêmes règles que pour *manier* :

Guy (<i>plie</i> , <i>pliera</i>) la carte	[pli] * [plij]
Cette carte a un seul pli	[pli] * [plij]
Guy <i>pliait</i> la carte	* [plie] [plije]
Cette carte a une <i>pliure</i>	* [pliy] [plijyr]

¹³ Ces suffixes ne sont jamais précédés d'un groupe *Cons-[y]*, car ils ne s'appliquent guère à des mots en *Cons-[y]* ou en *Cons-[y]*. Les paires *écu* - *écuyer* et *copahu* - *copayer* ne sont pas des cas nets d'application du suffixe *-ier*.

¹⁴ L'alternance entre [j] et [ij] en fonction de la présence d'un groupe *Obs-Liq* intérieur est à rapprocher du fait que les groupes *Obs-Liq-[j]* sont rares en français. A l'intérieur des mots, on ne les rencontre qu'avec les suffixes *-ions* et *-iez* de l'imparfait et du subjonctif :

Nous (*savions*, *voulons*) que vous redoubliiez. [blje] [blije]

La prononciation [pli] sans [j] devant voyelle, c'est-à-dire par exemple [plie] ou [pliyɾ], est interprétable mais inusitée. Vis-à-vis de cette distribution entre [i] et [ij], les suffixes *-era* du futur et *-erais* du conditionnel, ainsi que le suffixe de nominalisation *-ement*, se comportent comme des suffixes à initiale consonantique, puisqu'ils s'emploient obligatoirement avec la forme en [i]. Le suffixe *-oir*, au contraire, se comporte comme un suffixe à initiale vocalique, puisque les noms *plioir* et *oubloir* se prononcent obligatoirement [plijwar] et [ublijwar]. Les suffixes *-ions* et *-iez* de l'imparfait et du subjonctif semblent admettre les deux variantes :

Nous (savions, voulons) que vous pliez la carte [plije] ?[plije]

C'est la seule exception à la distribution complémentaire entre les radicaux [pli] et [plij].

Cette situation se prête à deux descriptions formelles différentes. Soit on considère les deux radicaux comme des variantes phonétiques et on leur donne une représentation phonémique commune, soit on considère que le radical est sélectionné par la conjugaison ou par la dérivation, comme pour le verbe *prendre* : dans la conjugaison de ce verbe, l'imparfait *prenait* et le futur *prendra*, par exemple, se construisent sur des radicaux distincts. Appliquée à *plier*, cette solution revient à le ranger dans une classe de conjugaison à plusieurs radicaux, donc dans une autre classe que le verbe *briller*. Il nous semble plus naturel de considérer les deux radicaux comme des variantes phonétiques, car leurs conditions d'emploi tiennent aux caractéristiques phonémiques du suffixe. Toutefois, cet argument n'est pas décisif, et le choix entre les deux solutions est une décision globale à prendre au niveau du lexique.

Nous parlerons d'une relation d'équivalence entre les formes en [i] et les formes en [ij] :

(6) *pli* [pli] = *plier* [plije]

Les conditions d'emploi sont déterminées essentiellement par les caractéristiques phonémiques du suffixe. On retrouve cette relation, avec les mêmes conditions d'emploi, dans les quelque 25 verbes dont le radical se termine en *Obs-Liq-[i]* ou en *Cons-[y]*, et dans les dérivés de ces verbes :

*approprier crier décrier démultiplier dépatrier écrier expatrier exproprier
multiplier oublier plier prier rapatrier récrier remplier strier supplier
surmultiplier appuyer enfuir ennuyer essayer fuir ressuyer*

Comme nous l'avons vu en 5b, p. 127, les adjectifs en *-ien* ou en *-iard* liés à des noms propres en *-ie* participent également à la relation (6) lorsque le *i* est précédé par un groupe intérieur *Obs-Liq* ou *Cons-[y]* :

Ombrie [ɔbri] = *ombrien* [ɔbrijɛ]
Brie [bri] = *briard* [brijar]

Dans tous ces exemples, le [i] ou le [ij] est précédé d'un groupe *Obs-Liq* ou *Cons-[y]* intérieur au mot. On observe également une relation entre [i] et [ij] dans d'autres verbes comme *fier* et *lier* (cf. section 4c, p. 121), mais cette relation est différente car la variante en [j], qui existe aussi dans ces exemples, s'emploie dans les mêmes conditions

que la variante en [ij] :

Guy se fie à son flair [fi] *[fij] *[fj]
Guy se fiait à son flair ?*[fie] [fije] [fje]

Ces exemples sans groupes *Obs-Liq*-[i] ni *Cons*-[y] intérieurs au mot ne relèvent donc pas de la relation étudiée ici.

Un autre exemple difficile est celui des préfixes savants de la section 4c terminés par un groupe *Obs-Liq*-[i] ou *Cons*-[y], comme *tri-* ou *équi-*. Devant consonne, ils se prononcent en [i] :

tricuspide [trikyspid] *équimolaire* [ekyimolær]

mais devant voyelle, ils admettent parfois deux prononciations, avec ou sans [j] :

triacide [triasid] ?[trijasid] *équiangle* [ekyiägl] ?[ekyijägl]

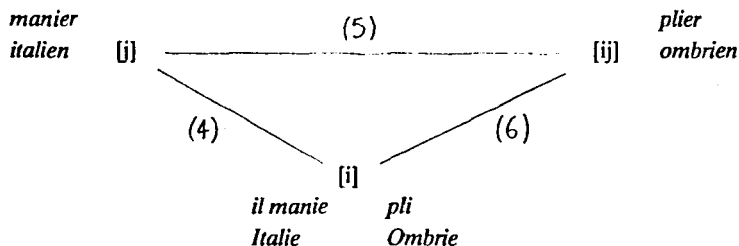
Des deux formes équivalentes, seule celle en [ij] peut être rattachée à la relation (6) entre [i] et [ij]. L'autre est caractéristique des composés savants.

On constate finalement que dans la relation (6), les conditions d'emploi de la forme en [i] (*pliera*) correspondent à celles de la forme en [i] dans la relation (4) (*maniera*) : ces formes sont employées obligatoirement en fin de mot ou devant un suffixe à initiale consonantique. La différence réside dans la présence d'un groupe intérieur *Obs-Liq* ou *Cons*-[y] dans la relation (6).

Quant à la forme en [ij] de (6) (*plier*), ses conditions d'emploi correspondent à celles de la forme en [ij] dans la relation (5) (*ombrien*) : [ij] est précédé d'un groupe intérieur *Obs-Liq* ou *Cons*-[y], et il est suivi d'une voyelle.

5.d. Représentation formelle

Les relations (4), (5) et (6) s'organisent donc selon le schéma suivant :



Les trois éléments [i], [ij] et [ij] de cette relation sont en distribution quasi-complémentaire, et leurs conditions d'emploi tiennent aux caractéristiques phonémiques du contexte gauche et du contexte droit :

- En fin de mot (*pli*) ou devant une consonne (*maniement*), on a [i].
- Devant une voyelle, on a [ij] après un groupe intérieur *Obs-Liq* (*plier*) ou *Cons*-[y] (*appuyer*), et [ij] dans les autres cas (*manier postier hawaïen*).

Cette règle n'est en défaut que devant le suffixe *-oir*, dont le début [wa] se comporte comme une voyelle et non comme une consonne, et devant les suffixes *-ions* et *-iez* de l'imparfait et du subjonctif :

<i>Nous (savions, voulons) que vous <u>maniez</u> bien ces pinc</i>	[manje] ?[manije]
<i>Nous (savions, voulons) que vous <u>pliez</u> la carte</i>	[plije] ?[plijje]

Les trois termes du schéma ci-dessus étant deux à deux en relation, nous parlerons d'une équivalence entre [j], [ij] et [i] dans ces mots. Nous les représenterons tous les trois d'une même façon. Pour cela nous pouvons choisir l'un des trois éléments. Il sera alors écrit /j/, /ij/ ou /i/ pour rappeler qu'il symbolise plusieurs formes phonétiques équivalentes. Dans cet exemple, le choix entre les trois éléments n'est pas arbitraire, car ce choix peut entraîner une perte d'information. C'est ce qui arrive, par exemple, si on choisit la séquence /ij/ pour représenter [j], [ij] et [i].

En effet, les formes telles que *il manie* s'opposent à des formes telles que *il habille*. Ces deux exemples ont des finales différentes : [i] pour *il manie*, [ij] pour *il habille* ; mais les deux finales apparaissent dans le même contexte phonémique : l'une comme l'autre est précédée d'une consonne et située en fin de mot. Donc, les deux finales présentent une différence phonétique non déductible du contexte, et on doit faire figurer cette différence dans la représentation phonémique. Il en résulte que le [i] de *il manie* et le [ij] de *il habille* ne peuvent être symbolisés tous les deux par la même séquence de phonèmes. Comme /ij/ est la solution la plus simple pour représenter le [ij] de *il habille* (/il abij/), on ne peut pas représenter *il manie* par /il manij/. Ainsi, /ij/ est inadéquat pour symboliser [j], [ij] et [i] dans (4), (5) et (6).

En revanche, /j/ et surtout /i/ sont plus appropriés. Considérons d'abord le cas de /j/. Les formes telles que *il manie* s'opposent à des formes telles que *il gagne* : dans ces deux formes, les finales [i] et [j] apparaissent dans le même contexte, ce qui implique qu'on doit leur donner des représentations phonémiques distinctes. Toutefois, si on représente le [nj] de *il gagne* par /N/, comme on le fait généralement¹⁵, on peut représenter *il manie* par /il manj/. Le choix de /j/ n'introduit pas d'autres difficultés.

Quant à /i/, un problème pourrait résulter du fait que les formes telles que *manier*, avec [j], s'opposent à des mots composés tels que *polyester*, avec [i]. Mais la représentation phonémique peut expliciter le fait que ces mots sont des composés, si on introduit une limite de mot entre les deux éléments du composé (/poli ester/), ce qui exclut toute ambiguïté (cf. section 4c, p. 121).

¹⁵ Cette représentation abstraite se justifie par le fait qu'on rencontre des [N] phonétiques dans des formes apparentées :

<i>Guy <u>gagne</u> la bouteille</i>	[gaNla]
--------------------------------------	---------

Le choix entre /i/ et /j/ pour représenter [j], [ij] et [i] semble donc relativement arbitraire. Nous avons opté pour /i/, ce qui donne les formes phonémiques suivantes :

<i>manier</i>	/manie/	→	[manje]
(il) <i>manie</i>	/mani/	→	[mani]
<i>plier</i>	/plie/	→	[plije]

Dans tous les exemples que nous avons pris, les équivalences (4), (5) et (6) ont lieu entre formes conjuguées d'un même verbe (*manier*, *il manie*), entre formes dérivées d'un même mot (*associer*, *association*), entre formes d'un même suffixe (*italien*, *ombrien*), etc., c'est-à-dire toujours entre des formes apparentées. C'est ce qui rend ces relations particulièrement nettes. Mais on peut se demander si ces relations sont plus générales et si, dans le reste du lexique, on les retrouve entre des formes non apparentées. En effet, les trois termes de l'équivalence constituent en fait trois configurations phonémiques abondamment représentées dans le lexique :

- La première configuration, celle de *manier*, correspond à la séquence *Cons-[j]-Voy*, avec la condition que le [j] n'est pas précédé d'un groupe *Obs-Liq* ni *Cons-[y]*. Cette situation de retrouve dans de nombreux mots comme *pied*, même en l'absence de toute relation avec l'une des deux autres configurations.
- La deuxième configuration, illustrée par *plier* et *appuyer*, est celle des séquences *Obs-Liq-[ij]-Voy* et *Cons-[yij]-Voy*. Elle figure également dans les mots comme *client* et *tuyau*, mais dans ces deux exemples, la configuration se retrouve dans toutes les formes apparentées.
- La troisième configuration est celle de *il manie*. Il s'agit des séquences *Cons-[i]* dans lesquelles [i] est suivi d'une consonne ou en fin de mot. Cette configuration est illustrée par des exemples relativement nombreux qui ne sont pas en relation avec des mots présentant l'une des deux autres configurations : ainsi, *nid* et *vide*.

Considérons, pour chacune des trois configurations, l'ensemble des mots dans lesquels elle figure, mais seulement lorsqu'elle figure également dans toutes les formes apparentées. Par exemple, pour la première configuration, le mot *pied* appartient à l'ensemble considéré, car il n'a pas de dérivés vivants, mais le verbe *signer* n'y appartient pas puisque la forme conjuguée *il signe* peut comporter une séquence *Cons-[j]* en fin de mot. Pour la deuxième configuration, la définition inclut *client* mais pas le verbe *griller*, car dans la forme conjuguée *il grille*, la séquence *Cons-Liq-[ij]-Voy* est remplacée par *Cons-Liq-[ij]* en fin de mot.

On obtient ainsi trois séries de mots : les séries *pied*, *client* et *vide*, comparables phonétiquement à *manier*, *plier* et *il manie*. Peut-on généraliser entre les séries *pied*, *client* et *vide* les relations (4), (5) et (6) qui concernent les types *manier*, *plier* et *il manie*, et représenter d'une même façon le [j] de *pied*, le [ij] de *client* et le [i] de *vide* ? ou doit-on laisser à ces mots des représentations proches de leur forme phonétique ? Cette dernière possibilité est simple : elle consiste à représenter *pied*, *client* et *vide* par les formes /pje/, /klijāt/ et /vid/. L'autre est plus abstraite : elle revient à symboliser par un même phonème le [j] de *pied*, le [ij] de *client* et le [i] de *vide*. Pour les mêmes raisons que celles données ci-dessus à propos de *manier*, *plier* et *il manie*, seuls /i/ et /j/ sont susceptibles de jouer ce rôle, sans qu'on puisse choisir localement. Si nous optons pour

/i/, cela donne les formes /pie/, /kliät*/ et /vid/, qui sont proches de l'orthographe, et cohérentes avec la représentation des formes /manie/, /plie/ et /mani/. Nous choisissons arbitrairement cette solution plutôt que la précédente.

6. Variations entre synérèse et diérèse à la limite de deux mots

Lorsqu'un mot normalement terminé par une voyelle fermée est suivi d'un mot commençant par une voyelle, cela produit généralement un hiatus phonétique¹⁶, c'est-à-dire que la voyelle fermée garde sa valeur syllabique et que la transition entre les deux voyelles se fait sans semi-consonne :

Guy arrive [gia] ?*[gija] *[gja]

Toutefois, dans certaines conditions, on observe trois autres comportements :

- Effacement de la voyelle fermée. On observe cet effacement dans quelques cas isolés qui font intervenir *qui, si, tu, y* :

C'est lui qu'a parlé

Luc ne sait pas s'il viendra

T'as raison

Luc aime Paris, il ira souvent

Ces cas sont étudiés plus loin, bien qu'il ne s'agisse pas de synérèses proprement dites.

- Diérèse avec transition en [j] entre les deux voyelles. Cette diérèse peut se produire lorsque la voyelle finale est [i] :

Luc a vu Max y aller [i] [ij] [j]

- Synérèse. Si la voyelle finale en hiatus est [i] ou [y], elle devient parfois [j] ou [y] :

Luc va y arriver [j]

Tu as raison [tʲa] [tʲa]

On peut donc distinguer, outre le cas général qui est celui d'un hiatus phonétique, trois comportements différents, qui sont d'ailleurs en concurrence dans certains exemples : plusieurs prononciations sont alors possibles et équivalentes. Chacun de ces trois comportements apparaît dans des conditions déterminées que nous allons examiner. Remarquons que l'observation exhaustive des faits est plus difficile ici que précédemment, car il s'agit de synérèses et de diérèses à la limite de deux mots, et non plus à l'intérieur des mots. Il ne suffit donc plus de considérer les mots séparément : on doit envisager leurs combinaisons.

¹⁶ Nous excluons de cette étude le cas où les deux mots en hiatus sont séparés par une pause d'une longueur proche de 1s due à une reprise de souffle ou à une hésitation. En effet, les mots séparés par la pause sont alors phonétiquement indépendants dans une large mesure.

6. a) Effacement de la voyelle fermée finale

Nous avons observé cet effacement dans les mots *qui, si, tu, y*.

Variante [k] du relatif *qui* devant voyelle

On rencontre la prononciation de niveau familier [k] pour *qui* relatif sujet avec antécédent, s'il est immédiatement suivi du verbe de la relative ou de pronoms atones précédant celui-ci :

Ceux (qu'ont essayé, qu'en ont pris) le disent

Luc a pris (tout, E) ce qu'est venu

Si *qui* est suivi d'un adverbial, cette prononciation est interdite :

C'est lui qui en une heure a tout fini *[kã]

Lui qui avant-hier encore est venu ici ! *[ka]

Par ailleurs, alors que *qui* relatif sujet avec antécédent n'est généralement pas mis en relief par la prosodie, *qui* admet d'autres emplois dans lesquels il joue un rôle prosodique particulier. Dans ces emplois, on n'observe pas de variante [k], qu'il s'agisse du relatif précédé d'une préposition, du relatif sans antécédent ou de l'interrogatif direct ou indirect :

Voici ceux (à, de, à cause de) qui on parle *[k5]

Qui essaye réussit *[ke]

Qui est venu ? *[ke]

Guy sait qui est venu *[ke]

Enfin, on ne trouve cette forme que si le mot suivant commence par une voyelle :

**Ceux qu'veulent peuvent venir*

La variante [k] de *qui* relatif sujet avec antécédent n'est donc acceptée que si certaines conditions syntaxiques et phonémiques sont réunies.

Variante [s] de *si*

Cette variante de niveau légèrement recherché n'est acceptée que devant le pronom *il* :

Luc ne sait pas (s'il viendra, s'ils viendront)

Luc demande si (Isabelle, elle) viendra *[siz], *[sɛl]

Elle apparaît aussi bien à la place de *si* conjonction de subordination, dans tous ses emplois, que de *si* interrogatif indirect :

S'il n'a pas réussi, du moins a-t-il essayé

S'il ne réussit pas, il demande de l'aide

Luc ne sait pas s'il viendra

Elle a donc des conditions d'emploi purement lexicales.

Variante [t] du pronom *tu*

Cette variante familière n'est acceptée que si le pronom est atone, donc immédiatement suivi du verbe dont il est sujet ou de pronoms atones précédant celui-ci :

<i>T<u>a</u>s raison</i>	[ta]
<i>T'<u>e</u>n as trop</i>	[tã]
<i>As-t<u>u</u> à boire ?</i>	*[ta]

Par ailleurs, elle n'apparaît que devant voyelle :

<i>T<u>u</u> restes là</i>	*[tr]
----------------------------	-------

Ses conditions d'emploi sont donc syntaxiques et phonémiques.

Présence ou absence de la *Ppv* y devant le verbe *aller*

Avec le verbe *aller*, la *Ppv* y est muette ou interdite aux temps en *ir-* mais se prononce toujours aux autres formes, y compris au futur antérieur (M. Gross, 1968, p. 36) :

*Jean aime Paris, il (y va, *va, *y ira, ira) souvent*

On peut rendre compte de cette alternance par un effacement de y sous des conditions lexicales et phonémiques. L'effacement est obligatoire, contrairement au cas de *qui*, *tu* et *si*.

En conclusion, ces quatre effacements de voyelles fermées finales en hiatus ont lieu dans des conditions lexicales très restrictives auxquelles s'ajoutent des conditions phonémiques. Les quatre effacements sont subordonnés à des conditions syntaxiques voisines : l'effacement n'a lieu que dans un mot atone immédiatement suivi d'un verbe ou de pronoms atones précédant un verbe. Suivant le cas, le mot atone est le sujet de ce verbe, une de ses *Ppv* ou une conjonction qui l'introduit. Comme ces effacements ne concernent que quelques mots, et que certains sont facultatifs et spécifiques à des niveaux de langue particuliers, on peut les représenter formellement en dédoublant les entrées lexicales de ces mots.

6. b) Diérèse avec transition en [j] à la limite de deux mots

L'extension de ce phénomène est difficile à cerner, car l'observation des faits est particulièrement malaisée en raison de l'existence de prononciations intermédiaires telles que celles qui comprennent un [j] léger (cf. section 3, p. 108). Certains faits se dégagent toutefois assez clairement :

- Devant [i], la prononciation [ij] est difficile, même dans le cas où la *Ppv* y est effacée entre les deux mots :

<i>Ce savon lui<u>i</u>rite la peau</i>	[ii] ?*[iji]
<i>Luc ne va pas à Pau, mais Guy<u>i</u>ra</i>	[ii] ?*[iji]

- Devant une autre voyelle, la prononciation en [ij] est généralement difficile, même dans un mot composé :

Guy ne sait pas si elle vient [siɛ] ?*[sije]

Paris est grand [rie] ?*[rije]

demi-abricot [ia] ?*[ija]

anti-atomique [ia] ?*[ija]

Guy lui a dit oui [lɥia] ?*[lɥija]

Le nom composé *demi-heure* est un des contre-exemples à cette règle : même très net, le [j] y est acceptable entre [i] et [ɛ].

- Toutefois, la *Ppv* y semble faire exception à cette règle. Les prononciations [i] et [ij] de y sont caractéristiques de la diction poétique traditionnelle, d'un style recherché ou d'un débit lent, mais sont assez acceptables, surtout après un groupe *Obs-Liq* :

Guy va sans doute y ouvrir un bar [tiu] ?[tiju]

Guy va peut-être y ouvrir un bar [triu] ?[triju]

- Les formes figées *il y avoir* constituent des cas particuliers pour lesquels les faits sont encore moins clairs. Lorsque l'expression est à l'infinitif, le terme qui précède immédiatement la *Ppv* y n'est jamais *il* atone. Les prononciations [i] et [ij] de y semblent alors également assez acceptables, surtout après un groupe *Obs-Liq* :

Il va (E, sans doute, presque) y avoir du vent *[i] *[ij] [j]

Il (va peut-être, semble) y avoir du vent ?[i] ?[ij] [j]

- Lorsque l'expression est conjuguée à un temps fini, la *Ppv* y est toujours précédée de *il* atone, du moins orthographiquement, car, à un niveau de langue familier, ce mot peut ne pas se prononcer. Phonétiquement, on observe les trois prononciations [ilj], [ij] et [j], auxquelles s'ajoutent [ilij] et peut-être [ili] qui caractérisent un débit particulièrement lent :

On dit qu'il y (a, aurait, eut) du vent [ilj] [ij] [j] *[i] ?[ilij] ?*[ili]

On dit qu'il y en a [ilj] [ij] [j] *[i] ?[ilij] ?*[ili]

6. c) Synérèse à la limite de deux mots

Les conditions dans lesquelles une synérèse peut se produire à la limite de deux mots sont de deux ordres : d'une part, des conditions lexicales et syntaxiques ; d'autre part, des conditions phonémiques.

Conditions lexicales et syntaxiques

L'apparition de [j] pour [i], et de [ɥ] pour [y], n'a lieu que dans un mot atone, c'est-à-dire sans relief prosodique. Dans la phrase suivante, la syllabe finale du mot *Paris* a une place importante dans la structure prosodique, et la synérèse est interdite :

Paris est grand [rie] *[rje]

En particulier, le relatif *qui*, lorsqu'il est sujet et qu'il a un antécédent, est généralement atone. Il admet alors la prononciation [kj] :

C'est Luc qui a dormi [kj]

Dans ses autres emplois, *qui* a des propriétés prosodiques spécifiques (cf. section 6a, p. 134) et n'admet pas cette prononciation :

C'est Luc à qui elle a parlé *[kje]

Qui essaye réussit *[kje]

Qui est venu ? *[kje]

Luc ne sait pas qui est venu *[kje]

De même, seuls les emplois atones de *si* admettent [sj] :

Luc (ne sait pas, partira) si elle vient ?[sjε]

Léa ne dort pas. - Si elle dort *[sjε]

Quant à la *Ppv* y, comme toutes les *Ppv*, elle est atone sauf quand elle est placée après le verbe à l'impératif. Cette distinction est nettement corrélée à l'acceptabilité de la prononciation [j] :

Luc y a mis un cendrier [kja]

Mets-y un cendrier *[zjε]

Les conditions dans lesquelles on observe cette synérèse sont proches des conditions dans lesquelles la voyelle finale des mêmes mots peut s'effacer (cf. section 6a, p. 134). Lorsque les conditions coïncident, trois variantes sont acceptables :

C'est Luc qui a dormi [kia]

C'est Luc qui a dormi [kja]

C'est Luc qu'a dormi [ka]

Dans les exemples où on observe la synérèse à la limite de deux mots, le mot atone est toujours suivi d'un verbe. Les deux mots peuvent se suivre immédiatement ou être séparés par des particules atones, mais s'ils sont séparés par un adverbial ou par un groupe nominal, la synérèse est plus difficile :

C'est Luc qui (a dormi, en a pris) [kj]

C'est Luc qui (en une heure, un jour) a tout fini ?[kj]

On retrouve la même différence pour la conjonction *si* :

Luc (ne sait pas, partira) si (elle, on) vient ?[sj]

Léa ne sait pas si (une heure lui suffira, à trois heures elle aura fini) ?*[sj]

Pour décrire plus précisément cette situation, on peut inventorier ces particules atones qui peuvent apparaître entre le mot où se produit la synérèse et le verbe. Ce sont les pronoms sujets atones (*je, tu, il, elle, nous, vous, ils, elles, on*) et les *Ppv* (*me, te, le, lu, lui, nous, vous, les, leur, en, y, ne*) :

Luc partira si (elle, on) (vient, en vient, ne vient pas, y va, lui parle, ...) ?[sj]

Par ailleurs, il existe toujours une relation syntaxique étroite entre le mot où se produit la synérèse et le verbe. Dans la phrase suivante, ce mot est le sujet du verbe :

Tu as raison [tʁa]

De même, il peut s'agir du pronom relatif sujet *qui* et du verbe de la relative :

Ceux qui ont essayé le disent [kj]

On peut considérer comme voisin le cas d'une extraction du sujet par *c'est... qui* :

C'est Luc qui a dormi [kj]

Il peut aussi s'agir de la *Ppv* y et du verbe correspondant :

Luc y a mis un cendrier [kj]

ou d'une conjonction et du verbe qu'elle introduit. On rencontre la prononciation [j] avec *si* conjonction de subordination, dans tous ses emplois, et avec *si* conjonction introduisant une complétive interrogative indirecte, mais pas avec la conjonction de coordination *ni* :

Luc (ne sait pas, partira) si elle vient ?[sj]

Luc n'a ni aimé ni détesté ce film *[nj]

Lorsqu'un mot atone en [i] ou [y] est en relation syntaxique avec un nom ou un adjectif, et non plus un verbe, la synérèse est exclue :

Il n'y a pas d'étrangers parmi eux *[mj]

Ce jardin est si abrité qu'un citronnier y pousse *[sja]

On peut conclure de toutes ces observations que l'acceptabilité de la synérèse à la limite de deux mots obéit à des conditions lexicales et syntaxiques. En particulier, la notion de mot atone intervient. Il s'agit d'une notion prosodique, puisqu'elle concerne l'intensité, l'intonation et le rythme. Nous y avons accès plutôt par la perception que par des critères formels, c'est pourquoi la notion de mot atone reste floue. Toutefois, elle est corrélée à des phénomènes phonémiques comme la prononciation de [i] final en hiatus, et à la structure syntaxique de la phrase. Des critères phonémiques et syntaxiques peuvent donc contribuer à préciser la notion de mot atone.

Conditions phonémiques

L'apparition de la synérèse admet également des restrictions en fonction du contexte phonémique. Bien sûr, la première condition est que le mot qui suit commence par une voyelle. Considérons le cas limite où le mot qui suit commence par une semi-consonne, comme dans les phrases suivantes :

La tension électrique va y ioniser les molécules [aijo] *[ajio] *[ajjo]

Luc va y huiler la serrure [aiyi] *[ajyi] *[ajji]

L'évêque y oignait les fidèles [iwa] *[jua] *[jwa]

On constate que le [i] final reste syllabique. Si la semi-consonne initiale est représentée formellement, à un certain niveau d'abstraction, par une voyelle fermée, on doit donc

ordonner les règles de telle sorte que la synérèse intérieure dans le deuxième mot ait lieu avant qu'une synérèse ne puisse se produire à la limite des deux mots.

Par ailleurs, la synérèse est pratiquement interdite si les voyelles en hiatus sont deux [i] ou deux [y] précédés d'une consonne :

Luc ne sait pas si il viendra *{sji}
C'est lui qui imagine tout ça ?*{kji}
Tu uses tes chaussures *{tuy}

alors qu'elle est acceptable lorsque deux [i] sont précédés d'une voyelle :

Luc va yimiter Max {aji}
Yimiter Max amuserait Luc ?{ji}

Acceptabilité de la diérèse

Nous avons examiné les conditions dans lesquelles la synérèse peut se produire. Nous allons maintenant préciser dans quelles conditions elle est facultative ou obligatoire. Le problème ne se pose que pour la *Ppv* y. En effet, partout ailleurs, la diérèse reste acceptable, que la synérèse le soit ou non :

C'est Luc qui a dormi [kia] [kja]
Luc ya mis un cendrier [kia] [kja]

Pour y, la prononciation en [j], lorsqu'elle est possible, est la plus courante. C'est celle de la langue familière. On la rencontre même après un groupe *Obs-Liq* :

Guy a vu l'autre y aller ?{trja}
Il semble y avoir du vent ?{blja}

La prononciation en [i], lorsqu'elle est en concurrence avec celle en [j], correspond à un débit plus lent et à une élocution plus soutenue, comme lorsqu'on lit à haute voix avec soin. Elle est plus acceptable après consonne qu'après voyelle :

Luc (vient de, va) y ouvrir un bar {iu}

Dans les phrases figées *il y avoir*, la prononciation [i] de la *Ppv* y est légèrement plus déviante, ainsi d'ailleurs que la prononciation [ij] (voir 6b, p. 135).

Représentation formelle

Comment rendre compte de cette synérèse facultative dans une représentation formelle ? Nous avons vu que la synérèse à la limite de deux mots est soumise à la fois à des conditions lexicales et syntaxiques et à des conditions phonémiques. Elle ne peut avoir lieu que lorsque ces deux types de conditions sont réunies. Il est naturel qu'une variation phonétique soit soumise à des conditions phonémiques. En revanche, les conditions lexicales et syntaxiques posent un problème de formalisation : un phénomène phonétique dépend d'informations qui relèvent d'un autre niveau. Plus précisément, la limite de mot où intervient le hiatus possède la propriété de permettre ou non la transformation de [i] en [j], mais cette propriété phonémique dépend de considérations lexicales et syntaxiques. Il serait intéressant que cette propriété figure dans les

représentations phonémiques, mais cela reviendrait à intégrer des informations syntaxiques dans les représentations phonémiques, ce qui ne peut pas se faire systématiquement. Le formalisme phonémique semble donc peu adapté pour rendre compte de la synérèse à la limite de deux mots. Un système formel ne peut traiter ce phénomène que s'il comporte à la fois des informations phonémiques et des informations syntaxiques suffisamment élaborées, et s'il peut accéder à ces deux types d'informations conjointement. Le formalisme utilisé en phonémique ne répond pas à ces exigences, puisqu'il comprend seulement des représentations phonétiques et phonémiques constituées en dictionnaires, et des transformations entre ces représentations.

6. d) Conclusion

En cas de rencontre entre une voyelle fermée et une voyelle à la limite de deux mots, le comportement régulier est un hiatus phonétique. Quelques mots présentent un comportement différent, dépendant entre autres de la structure syntaxique : la voyelle finale peut se consonantiser. Enfin, on observe plusieurs autres comportements dont l'extension dans le lexique est ponctuelle, et qui peuvent être représentés formellement en séparant les variantes observées en éléments lexicaux distincts.

Chapitre VI. Applications informatiques

La description phonétique et la construction de dictionnaires électroniques débouchent sur la réalisation de plusieurs sortes d'applications informatiques dans le domaine de la communication entre l'homme et la machine. L'aide à la correction orthographique de mots constitue une première application, envisageable à court terme et qui revêt une importance considérable pour les systèmes qui analysent des textes tapés par l'utilisateur, car ces textes ne sont jamais exempts d'erreurs.

A plus long terme se prépare l'utilisation informatique de la langue naturelle et de la parole : bien des données fondamentales restent à réunir avant d'atteindre le stade de l'utilisation industrielle. Cependant, nous nous efforçons de réunir et d'organiser ces données d'une façon qui soit compatible avec l'utilisation industrielle qui pourra en être faite le moment venu.

1. L'aide à la correction orthographique

La description phonétique apporte une solution partielle mais réaliste au problème de corriger l'orthographe des textes mis sur support informatique. Pour le montrer, nous allons d'abord considérer le problème de la correction orthographique dans son entier.

Généralités

L'aide à la correction orthographique de textes se situe dans une perspective de saisie sur ordinateur. Le problème des erreurs de saisie est très général : il concerne tous les textes destinés à être traités manuellement ou automatiquement, imprimés, stockés ou transmis. Que les textes saisis soient réduits à des mots ou des groupes de mots, comme dans le cas de réponses à des formulaires, ou qu'ils constituent des phrases, des erreurs sont inévitablement commises lors de la saisie : même un personnel spécialisé ne peut effectuer une saisie correcte à grande échelle sans aucune aide.

Peut-on automatiser la correction d'erreurs ? Remarquons d'abord que dans le cadre d'un système informatique, la détection des erreurs et leur correction sont nécessairement deux opérations distinctes. La première consiste à sélectionner les éléments dans lesquels figurent des erreurs, ou qu'on peut soupçonner d'être erronés. La deuxième consiste à produire des éléments corrects susceptibles de correspondre aux éléments sélectionnés lors de la détection. La détection automatique d'une erreur n'apporte pas en elle-même l'élément correct qui doit remplacer l'élément erroné. Pour trouver automatiquement cet élément correct, une ou plusieurs procédures de correction doivent s'ajouter aux procédures qui détectent les erreurs.

On détecte des erreurs par des procédures de vérification qui analysent les textes saisis. En cas d'échec de cette analyse, on peut conclure que le texte ou le mot saisi comporte soit une ou des erreurs, soit un élément inconnu du système, par exemple un

mot nouveau. A ce stade, on ne peut pas choisir entre ces deux hypothèses, ni *a fortiori* préciser de quel type d'erreur il s'agit. C'est seulement après la correction des erreurs éventuelles qu'on peut répondre à ces questions. Lors de la détection d'une erreur, cette partie du système se borne donc à appeler le ou les systèmes de correction.

Plusieurs types d'erreurs peuvent se produire. On peut établir une typologie des erreurs afin de concevoir des procédures de vérification et de correction qui soient adaptées aux différents types d'erreurs. Pour établir cette typologie, nous nous basons sur des critères formels.

Un premier critère permet de définir les erreurs lexicales. Il s'agit des erreurs qui produisent un mot qui ne fait pas partie du vocabulaire français. Ainsi, les erreurs suivantes sont des erreurs lexicales :

<i>éfac</i> er	pour	<i>eff</i> acer,
<i>ilôt</i>	pour	<i>ilot</i> ,
<i>le sautres</i>	pour	<i>les autres</i> ,
<i>résoud</i>	pour	<i>résout</i> ,
<i>sortie</i>	pour	<i>sortie</i> ,
<i>avecun</i>	pour	<i>avec un</i> ,

car elles produisent les formes *éfac*er, *ilôt*, *sautres*, *résoud*, *sort*, *avecun* qui n'existent pas en français. Il arrive qu'au contraire, une erreur produise un mot distinct mais correct. C'est le cas des erreurs suivantes :

<i>peaux</i>	pour	<i>pots</i> ,
<i>pou ravis</i> er	au lieu de	<i>pour aviser</i> ,
<i>carre</i>	pour	<i>carte</i> ,
<i>correcteur</i>	pour	<i>connecteur</i> ¹ ,
<i>couse</i>	pour	<i>cause</i>

(*couse* est une forme du verbe *coudre*). La notion d'erreur lexicale peut s'élargir dans le cas particulier d'une application donnée, si on sait à l'avance que certains mots n'apparaîtront jamais. Quant au cas général, les erreurs lexicales constituent une catégorie d'erreurs particulièrement faciles à détecter automatiquement. Cette opération nécessite toutefois d'utiliser une liste aussi exhaustive que possible des mots du vocabulaire français, y compris les formes fléchies : les mots qui comportent des erreurs lexicales ne seront pas trouvés dans cette liste. Tel est le principe du vérificateur orthographique lexical de M. Silberstein (1987) : ce programme utilise le dictionnaire DELAF-L, présenté dans le chapitre II, section 2, p. 29.

Mis à part le cas des erreurs lexicales, la détection requiert une analyse du texte. La détection des fautes d'accord, notamment, demande fréquemment des données grammaticales formalisées. Dans l'état actuel des connaissances, une telle analyse ne peut être faite que d'une façon très dépendante des textes, et en particulier de leur vocabulaire, de leur nature, de leur origine et de leur domaine. Par exemple, dans le

¹ Cette erreur est attestée : elle figure en première page de couverture dans le titre du n° 77 de la revue *Langages* (février 1988).

cadre d'applications où les textes saisis sont réduits à des termes isolés, et où on dispose de la liste des termes corrects susceptibles d'apparaître, des erreurs peuvent être détectées dans ces termes, même s'il ne s'agit pas d'erreurs lexicales. Par exemple, *peaux d'échappement* pour *pots d'échappement* n'est pas une erreur lexicale, mais elle introduit un terme incorrect, ce qui permet de la détecter automatiquement en recherchant le terme dans la liste des termes corrects.

Un deuxième critère nous semble important pour la classification des erreurs. Toute erreur est une déformation de la graphie, mais seules certaines erreurs déforment également la prononciation. Ainsi, *éfacier* pour *effacer* et *peaux* pour *pots* ne déforment pas la prononciation, tandis que *plus* pour *puis* et *couse* pour *cause* la déforment. Cette distinction a son intérêt car on constate qu'en français une proportion importante d'erreurs ne déforment pas la prononciation ou la déforment peu. Il peut s'agir aussi bien d'erreurs lexicales, comme *éfacier* pour *effacer*, que d'erreurs non lexicales, comme *peaux* pour *pots*, de fautes d'accent, comme *râcler* pour *racler*, de fautes de segmentation, comme *beaucoou pplus* pour *beaucoup plus*, ou de fautes d'accord, comme *exclus* pour *exclu*. Le fait que ces erreurs ne déforment pas la prononciation n'a pas d'intérêt pour leur détection, puisque la détection se fait à partir de l'orthographe. En revanche, nous allons voir qu'elle peut faciliter la correction.

Cette typologie met en évidence les difficultés de la détection comme de la correction des erreurs. En effet, en ce qui concerne la détection, on n'a aucun moyen de distinguer à coup sûr un mot nouveau d'un mot erroné ; de plus, certaines fautes d'accord ne pourraient être détectées qu'à la suite d'une analyse syntaxique approfondie.

Quant à la correction, on dispose de plusieurs méthodes pour retrouver la forme correcte à partir de la forme erronée. Pour cela, on tire profit de la "ressemblance" entre la forme erronée et la forme correcte : une erreur n'apporte qu'une déformation limitée par rapport à la forme correcte. Toutefois, cette "ressemblance" peut se situer à plusieurs niveaux, dont voici trois exemples.

- (1) Dans le cas des erreurs dues à l'utilisation d'un clavier, la forme correcte et la forme erronée sont proches en termes de chaînes de caractères : elles ne diffèrent souvent que par un ou deux ajouts, omissions, remplacements ou interversions de lettres.
- (2) Nous avons déjà parlé des erreurs qui ne déforment pas la prononciation, ou qui la déforment peu. Dans ce cas, la forme correcte et la forme erronée ont une ressemblance phonétique.
- (3) Certaines erreurs déforment peu l'image visuelle du mot, comme *borne* pour *borne*. Peu visibles, elles échappent facilement à un relecteur humain, et constituent probablement une catégorie non négligeable dans certaines situations, en particulier dans l'industrie de l'édition.

Ainsi, les méthodes de correction correspondent souvent à des hypothèses sur la cause ou le type de l'erreur. Mais les causes d'erreurs possibles sont nombreuses, et lorsqu'une erreur a été détectée, la cause de l'erreur ne peut être connue qu'après la correction. C'est pourquoi, dans le cas général, les différentes méthodes de correction amènent

plusieurs propositions de correction, parmi lesquelles on ne peut pas choisir automatiquement. Ainsi, un correcteur orthographique entièrement automatique ne peut être fiable. Or les conséquences d'un choix erroné de la part d'un système de correction sont difficilement acceptables : un tel système introduirait des erreurs dans des textes corrects, ou remplacerait des erreurs par d'autres erreurs.

Toutefois, un système de détection et de correction d'erreurs peut apporter une aide à son utilisateur, en détectant des erreurs et en lui proposant des corrections. Cette aide est appréciable si la forme correcte figure parmi les propositions de corrections, et si ces propositions ne sont pas trop nombreuses à consulter. Ces objectifs peuvent être atteints en ciblant la recherche des corrections à partir d'hypothèses sur la cause ou le type de l'erreur. Nous avons exploré une méthode pour proposer des corrections adaptée à un type particulier d'erreurs : celles qui ne déforment pas la prononciation.

L'aide à la correction par phonétisation

Les erreurs qui ne déforment pas la prononciation ou qui la déforment peu constituent un type d'erreurs fréquent en français. Par conséquent, lorsqu'une erreur a été détectée, l'hypothèse qu'il s'agirait d'une erreur de ce type est une des hypothèses à considérer en priorité. Un système informatique peut explorer cette hypothèse en produisant, à partir du mot erroné, un ou plusieurs mots corrects qui se prononcent de la même façon ou d'une façon voisine. Cette méthode a ses ancêtres dans certaines procédures manuelles utilisées par la police judiciaire : comme les noms de personnes transmis à ces services le sont souvent de bouche à oreille, un code phonétique rudimentaire a été mis en place pour identifier plus facilement l'orthographe de ces noms.

Une technique analogue est employée dans l'annuaire électronique du téléphone, pour l'aide à la correction des noms propres, ainsi que dans les systèmes d'aide à la correction construits au CERFIA par G. Pérennou, F. Lahens, P. Daubèze et M. de Calmès (1986) et au GRTC par J. Véronis (1987). Ces trois systèmes utilisent des règles de réécriture entre séquences orthographiques, par exemple *o* → *au*, *o* → *eau* et *au* → *eau*. En face d'une forme erronée telle que *Pau*, le correcteur de l'annuaire électronique produit des "formes réduites" telles que *Po*, qui est obtenue par application d'une règle *au* → *o*. Ces formes réduites sont ensuite recherchées dans l'annuaire.

G. Pérennou et al. (1986) utilisent un dictionnaire spécialisé construit de manière semi-automatique à partir de BDLEX. Les mots y sont découpés en sous-séquences orthographiques qui ont chacune une valeur phonétique : par exemple, *orthodoxe* est segmenté *o-r-th-o-d-o-xe*. Des règles de réécriture telles que *th* → *t* et *o* → *ho* permettent d'engendrer des formes erronées qui se prononcent de la même façon, comme *hortodoxe*. D'autres règles simulent les erreurs caractéristiques de l'utilisation d'un clavier : omissions, ajouts, remplacements et interversions de lettres. Etant donnée une forme erronée, un programme qui utilise un modèle probabiliste permet d'extraire du dictionnaire les mots à partir desquels ces règles reproduisent la forme erronée. Les performances que l'on peut attendre de ce correcteur dépendent essentiellement (1) du

dictionnaire de mots segmentés, et (2) du système de réécriture. Il serait intéressant de connaître les principes qui ont présidé à l'élaboration de ces deux éléments.

Le système de J. Véronis (1987), devant une forme erronée telle que *pau*, extrairait d'un éventuel dictionnaire une liste de mots tels que *peau*, qu'on peut obtenir à partir de *pau* en appliquant une des règles de réécriture. D'autres mots pourraient également être extraits, malgré leur dissemblance phonétique avec *pau* : par exemple, *pua*, *tau* ou *veau*. Les performances de ce système ne se limitent donc pas aux erreurs qui préservent la prononciation. En ce qui concerne celles-ci, l'élément central du système est constitué par une matrice de règles de réécriture entre séquences orthographiques. Cette matrice a été élaborée après l'examen d'une liste de mots réduite : moins de 4.000 formes. L'auteur estime que le degré de généralité atteint est suffisant pour rendre compte de presque toutes les erreurs orthographiques de ce type, mais seuls des tests en vraie grandeur permettraient de confirmer cette estimation, et cela supposerait l'utilisation d'un dictionnaire représentatif du français.

Ainsi, si plusieurs systèmes d'aide à la correction orthographique utilisent des méthodes phonétiques, très peu manipulent explicitement des transcriptions phonétiques, et la plupart mettent en jeu des règles de réécriture qui associent directement des orthographes correctes et des orthographes erronées.

Toutefois, dans quelques systèmes, dont le nôtre, la proposition de corrections se fait par l'intermédiaire d'une ou plusieurs transcriptions phonémiques du mot erroné. Dans l'hypothèse où le mot correct et le mot erroné se prononcent de la même façon, la transcription phonétique commune aux deux mots permet de retrouver le mot correct à partir du mot erroné. Cette technique est encore peu utilisée, car la plupart des auteurs reculent devant la complexité d'un phonétiseur par règles. D'autres ont recours à des phonétiseurs par règles rudimentaires (L. Davidson, 1962 ; Ch. Fluhr, communication personnelle). Quant à nous, nous utilisons les algorithmes de phonétisation et les dictionnaires évoqués dans le chapitre II.

Le programme est modularisé en deux étapes :

- la production d'une ou plusieurs transcriptions phonétiques du mot erroné,
- la recherche du ou des mots corrects qui correspondent à ces transcriptions phonétiques.

La première étape consiste à phonétiser le mot erroné. Cette opération ne peut se faire que par un algorithme, et non à l'aide d'un dictionnaire, car on ne peut lister dans un dictionnaire tous les mots erronés qui peuvent apparaître. Par ailleurs, pour phonétiser le mot erroné, on ne dispose pas d'autres informations que son orthographe, qui n'est pas toujours suffisamment explicite pour que la phonétisation puisse se faire sans ambiguïté. En particulier, on ne peut écarter systématiquement l'hypothèse où le mot erroné se prononcerait suivant une règle exceptionnelle, et non suivant les règles locales les plus générales. Ainsi, pour que le système puisse proposer *raconter* quand on lui soumet le mot erroné *racompter*, il doit émettre l'hypothèse que le *p* ne se prononce pas dans ce mot. Or rien dans la forme de ce mot n'indique que le *p* ne se prononce

pas : par exemple, dans *dompter*, le *p* peut se prononcer. C'est pourquoi il est souhaitable, dans le cadre de cette application, que le phonétiseur donne plusieurs transcriptions phonétiques de certains mots. Cette première étape est effectuée par le phonétiseur de mots inconnus présenté dans le chapitre II, section 5, p. 41.

La seconde étape consiste à rechercher le ou les mots existants qui correspondent aux transcriptions phonétiques trouvées. Ainsi, pour [rakɔ̃te], les mots trouvés sont :

<i>racontaient</i>	<i>raconté</i>	<i>raconter</i>
<i>racontais</i>	<i>racontée</i>	<i>racontés</i>
<i>racontait</i>	<i>racontées</i>	<i>racontez</i>

Cette recherche peut se faire dans un dictionnaire approprié. Les clés d'entrée du dictionnaire doivent être des transcriptions phonétiques, et pour chaque transcription, le dictionnaire doit donner la liste des mots corrects correspondants. Cette liste doit inclure les mots fléchis, c'est-à-dire les verbes conjugués, les pluriels et les féminins. En effet, ce sont des mots fléchis qui figurent dans les textes. Le dictionnaire voulu est donc un dictionnaire d'homonymes dont les entrées sont rangées par ordre alphabétique des formes phonémiques. Ce dictionnaire a été construit sous le nom de DELAP-FH à partir du DELAP-F qui est le dictionnaire de formes phonémiques fléchies. La maintenance du dictionnaire DELAP-FH peut se faire à partir du DELAP : il est commode de faire la mise à jour et les corrections sur le DELAP et sur le programme de flexion ; après chaque modification, le DELAP-FH peut être reconstruit automatiquement à l'aide du programme de flexion.

Les deux étapes de la proposition de corrections sont donc effectuées respectivement par l'algorithme de phonétisation de mots inconnus, et par une procédure de recherche dans le dictionnaire DELAP-FH. Ces deux algorithmes sont implantés en langage C sur un Vax 730.

2. La synthèse de la parole

Les systèmes à sorties vocales actuellement sur le marché sont capables de produire un jeu de quelques dizaines de messages oraux. Les messages sont enregistrés puis stockés sous forme de parole compressée, car ils sont connus à l'avance et suffisamment peu nombreux pour que cela soit possible. De tels systèmes fonctionnent parfaitement en l'absence de synthèse de parole. Nous n'en parlerons pas ici.

La raison d'être de la synthèse de la parole est de permettre des sorties vocales exprimées en français ou en d'autres langues, et suffisamment nombreuses et variées pour que leur contenu sémantique puisse être précis et informatif.

Lorsque les messages à synthétiser sont trop nombreux, on est amené à les assembler en combinant des éléments linguistiques plus courts : mots, syllabes ou segments phonétiques. C'est donc à partir d'un texte phonétique que la synthèse peut avoir lieu. Outre les difficultés liées à la synthèse de la parole elle-même, se pose alors le problème de la production automatique de ces textes phonétiques. Il est essentiel qu'ils correspondent à la prononciation courante des messages produits, ou du moins à

une des prononciations possibles. En cas de variations phonétiques libres, comme celles qui apparaissent dans les mots du type *louer* (cf. chapitre V, section 4, p. 109), les différentes variantes sont acceptables et n'importe laquelle peut figurer dans le texte phonétique. On peut donc ignorer ces variations libres. Il en va autrement des variations phonétiques dans lesquelles des conditions imposent l'une ou l'autre des variantes. Considérons par exemple les liaisons. Certaines sont obligatoires et d'autres interdites, comme le montrent les exemples suivants :

*C'est son ([], *[n]) jouet favori*

C'est son ([], [n]) amusement favori*

Un texte phonétique correct comportera donc la liaison dans le premier cas et ne la comportera pas dans le second. L'enjeu n'est pas négligeable, car les liaisons obligatoires et interdites jouent un rôle important dans l'intelligibilité à l'oreille.

De plus, on doit déterminer les valeurs des paramètres prosodiques associés à ces éléments : la durée, l'intensité et la fréquence fondamentale. Ces paramètres contribuent d'une façon cruciale à l'intelligibilité et au naturel des messages oraux. Lorsque nous lisons des textes écrits, nous nous passons d'informations prosodiques complètes sans inconvénient pour l'intelligibilité. Cette habitude a pu faire croire qu'un traitement sommaire de la prosodie est suffisant, du moins pour la synthèse de messages oraux de moyenne qualité. Les résultats des systèmes réalisés convergent pour montrer qu'au contraire, l'élaboration d'une prosodie de bonne qualité est une condition préalable à l'utilisation industrielle de la parole synthétique.

Peut-on automatiser la production de textes phonétiques et des paramètres prosodiques à rattacher à ces textes ? Dans une application informatique où on cherche à produire des messages nombreux et variés, on est inévitablement confronté à une variété de situations et de combinaisons qui est caractéristique de la langue naturelle. En même temps, et malgré cette complexité, les systèmes ne seront commercialisables que s'ils satisfont à une exigence de fiabilité. Pour mettre au point des techniques générales et robustes dans ce domaine, une approche réaliste est d'envisager d'une manière aussi exhaustive que possible les cas particuliers et les cas difficiles qui peuvent apparaître dans cette variété de situations. Telle est l'une des difficultés auxquelles est confrontée la production de messages synthétiques, dès que les messages à produire sont variés, nombreux ou imprévisibles.

Le calcul de la prosodie aussi bien que l'établissement d'un texte phonétique correct reposent sur l'utilisation de connaissances syntaxiques précises sur le texte des messages. L'exemple du système MITalk, développé au MIT pour la synthèse de la parole à partir de textes, montre dans le cas de l'anglais que des informations syntaxiques permettent un calcul des valeurs des paramètres prosodiques et que des améliorations pourront être obtenues en partant d'informations syntaxiques plus complètes (J. Allen, M.Sh. Hunnicutt et D. Klatt, 1987). Quant à la production du texte phonétique, elle met également en jeu l'utilisation d'informations grammaticales et syntaxiques. Par exemple, la détermination des liaisons obligatoires et facultatives en français repose sur plusieurs paramètres linguistiques. Rappelons ces paramètres que

nous avons déjà rencontrés dans le chapitre I, section 2, p. 12 :

- l'initiale du mot qui suit ;
- les relations syntaxiques entre les deux mots ;
- les traits flexionnels du premier mot ;
- etc.

Un autre exemple est fourni par les homographes non homophones du français, c'est-à-dire les mots qui s'écrivent de la même façon mais se prononcent différemment. La désambiguïsation automatique de la plupart des paires d'homographes peut se faire si on connaît leurs catégories grammaticales. Ainsi, la forme orthographique *poster* est prononcée différemment s'il s'agit d'un nom ou s'il s'agit d'un verbe. Dans les cas résiduels, qui sont peu nombreux, les deux formes homographes correspondent à une même catégorie grammaticale, et ce sont des informations syntaxiques plus élaborées qui seraient utiles. Par exemple, le nom *os* peut se prononcer différemment au singulier et au pluriel ; le nom *punch* se prononce différemment dans les phrases suivantes :

<i>Guy a eu le punch de répondre</i>	[pœnS]	*[pɔ̃S]
<i>Guy a bu du punch</i>	?*[pœnS]	[pɔ̃S]

La production automatique de textes phonétiques corrects et d'une prosodie acceptable suppose donc qu'on dispose d'informations grammaticales et syntaxiques formalisées sur le texte des messages à synthétiser. Cette conclusion a des conséquences importantes, étant donné les difficultés qu'on rencontre à trouver automatiquement et de façon fiable des informations syntaxiques formalisées sur un texte écrit : la synthèse de la parole à partir du texte se heurte à des problèmes analogues à ceux de l'analyse syntaxique automatique.

Toutefois, la synthèse de la parole à partir de textes n'est pas la seule utilisation envisageable de la parole synthétique. En effet, un message oral peut être synthétisé pour exprimer le contenu d'une représentation sémantique. Cela suppose l'utilisation d'un générateur automatique de textes en langue naturelle. Or, pour exprimer dans un texte le contenu de la représentation sémantique, un générateur automatique est amené à élaborer la structure syntaxique du texte (L. Danlos, 1986) : les informations grammaticales et syntaxiques nécessaires à l'établissement du texte phonétique et de sa prosodie sont donc disponibles sous la forme de données formalisées utilisables par les programmes. Nous avons participé à la réalisation d'un système de génération de messages oraux à partir de représentations sémantiques (L. Danlos, F. Emerard et E. Laporte, 1986), dans lequel les éléments suivants ont été utilisés pour mettre sous

forme phonétique le texte produit par le générateur :

- le programme de flexion automatique, qui étant donné la forme canonique d'un mot, par exemple l'infinitif d'un verbe, calcule une transcription phonémique de la forme fléchie voulue ;
- un extrait du dictionnaire DELAP, qui donne pour chaque entrée une représentation phonémique et des indications nécessaires à la flexion ;
- une procédure de placement des liaisons obligatoires qui utilise les informations syntaxiques sur le texte, ainsi que la forme phonémique des mots, pour déterminer quels mots du texte font obligatoirement la liaison.
- l'algorithme de calcul du texte phonétique à partir du texte phonémique (cf. chapitre III, section 8, p. 68).

Ce système pourra être étendu afin de prendre en compte davantage de variations phonétiques, notamment en ce qui concerne les liaisons et les liaisons.

3. La reconnaissance de la parole

Les variations phonétiques constituent un problème majeur en reconnaissance automatique de la parole. Certaines de ces variations sont particulièrement fines et difficiles à observer à l'oreille, mais d'autres sont nettement perceptibles et suffisamment importantes pour qu'on attribue aux différentes variantes des transcriptions phonétiques distinctes. Ainsi, l'élément *-endre* admet plusieurs prononciations naturelles dans l'expression *prendre du temps*. On peut les représenter par [ãdrɔ] (où la voyelle [ɔ] admet des variations), [ãd] et [ãn]. Elles correspondent aux prononciations [prãdrɔdytã], [prãddytã] et [prãndytã], qui est la plus courante. Comme les variantes sont suffisamment différentes pour justifier qu'on leur attribue plusieurs transcriptions, cette variation peut être représentée dans un dictionnaire phonémique. Dans la suite, nous parlerons de variations phonétiques pour désigner des variations de ce type. Leur étude systématique est possible, et nous allons voir l'utilité d'une telle étude pour la reconnaissance de la parole. Pour cela nous envisagerons successivement plusieurs techniques de reconnaissance.

Cas de la reconnaissance globale

Lors de la constitution d'un système de reconnaissance globale, des références sont soit construites à partir d'échantillons prononcés par le locuteur, soit synthétisées au moyen de connaissances acoustiques générales. Le signal de parole à reconnaître est mis en correspondance avec ces références, cette opération mettant généralement en jeu au moins une fois un alignement temporel dynamique. Dans ce cadre, le problème est qu'en cas de variabilité, le locuteur peut utiliser des variantes très différentes lors de l'apprentissage et pour la reconnaissance, par exemple [illegal] et [ilegal]. Si c'est le cas, il y aura des différences phonétiques importantes entre les références et le signal de parole à reconnaître. Deux approches ont été proposées pour résoudre ce problème, et de nombreux compromis sont d'ailleurs possibles entre elles :

a) La première consiste à prévoir des références distinctes pour des variantes distinctes. Ceci est généralement fait au hasard des diverses prononciations que donne le locuteur lors de l'apprentissage, par exemple en lui faisant prononcer chaque énoncé un certain nombre de fois, pour faire apparaître les variabilités libres, ou en lui faisant prononcer des mots en contexte, pour faire apparaître les variabilités contextuelles. Ainsi, T. Watanabe (1986) utilise comme échantillons pour l'apprentissage des énoncés de plusieurs syllabes, prononcés de manière naturelle, car cela fait apparaître des variations phonétiques (assourdissements, nasalisations, neutralisations, chutes, abrègements de voyelles), alors que dans une expérience antérieure (T. Watanabe, 1983), des échantillons d'une ou deux syllabes prononcés avec un souci de clarté ne faisaient pas apparaître ces variantes. Une limite de cette technique est que, même avec des prononciations répétées d'un même énoncé, il est difficile d'inciter le locuteur à utiliser lors de l'apprentissage toutes les variantes phonétiques libres et contextuelles qu'il est susceptible d'employer par la suite. S'il ne les utilise pas toutes, la description des variantes phonétiques est insuffisante. En fait, elle ne pourrait s'approcher de l'exhaustivité que si on disposait de données systématiques sur les variations phonétiques, et si on concevait l'apprentissage en tenant compte de ces données. Il en est de même, *a fortiori*, si les références sont synthétisées au lieu d'être construites à partir d'échantillons.

b) La deuxième solution consiste à ne pas prévoir de variantes au niveau des références, et à contrôler les variations phonétiques, comme les variations acoustiques fines, à l'aide de l'algorithme de reconnaissance acoustique. Ce composant du système permet en effet de mettre en correspondance une référence et une portion de parole même si elles présentent une ressemblance seulement approximative. Il fonctionne généralement par alignement temporel dynamique. Or les différences qu'on peut neutraliser par alignement temporel dynamique sont surtout les variations acoustiques fines et les distorsions temporelles, alors que de nombreuses variations phonétiques ne consistent pas seulement en une distorsion temporelle (cf. les prononciations de *prendre* évoquées ci-dessus). L'algorithme de reconnaissance devrait donc être spécifiquement tolérant à ce type de variations. Ici encore, rien ne peut être entrepris sans données sur les variations phonétiques, et en particulier sur leur extension lexicale.

Cas de la reconnaissance analytique

Dans un système de reconnaissance analytique, une analyse acoustique du signal de parole permet de détecter des segments phonétiques. Les séquences ainsi obtenues sont ensuite comparées aux entrées d'un dictionnaire phonétique. Cette dernière opération est une mise en correspondance temporelle et le problème des variantes phonétiques se retrouve à ce niveau : en cas de variabilité, toutes les variantes rencontrées doivent pouvoir être reconnues. Ici encore, deux moyens sont envisageables pour cela, ainsi que des compromis entre ces deux moyens :

a) Décrire les variantes exhaustivement et prévoir des entrées de dictionnaire distinctes pour des variantes distinctes. Ainsi, dans l'exemple précédent, les variantes en [adro], [ad] et [an] sont connues du système, et subissent une reconnaissance acoustique indépendamment les unes des autres. Lorsque l'une d'elles a été reconnue, il faut

ensuite passer de cette forme à l'élément lexical, qui, dans ce cas, est probablement unique et a reçu une forme phonémique de base représentant toutes les variantes phonétiques. Le dictionnaire doit mettre chaque variante en correspondance avec une forme phonémique. Cette méthode suppose donc l'existence d'un système de données phonétiques élaboré, avec notamment plusieurs niveaux de représentation : formes de base et variantes éventuelles.

b) L'autre solution consiste à concevoir un algorithme de mise en correspondance temporelle qui comporterait une description implicite des variantes (R. Vivès, 1985) et qui associerait directement les séquences phonétiques détectées à des formes de base représentant toutes les variantes. Un tel algorithme, pour être spécifiquement tolérant aux variations phonétiques qui peuvent se produire, inclurait des données systématiques sur ces variations.

En conclusion, aussi bien dans le cadre de l'approche analytique que dans le cadre de l'approche globale, la prise en compte des variations phonétiques suppose une description de ces variations, et en particulier de leur extension lexicale. Plus généralement, et en dehors de tout choix technique, il semble incontournable qu'on doive disposer de ces données dès lors qu'on désire reconnaître un vocabulaire étendu (P.S. Cohen et R.L. Mercer, 1974).

Reconnaissance globale avec comme unité la portion de parole comprise entre deux centres de syllabes consécutifs

Dans cette section, nous envisageons la prise en compte des variations phonétiques dans le cadre d'une technique particulière : la reconnaissance globale avec alignement temporel dynamique, l'unité de reconnaissance étant la portion de parole comprise entre deux centres de syllabes consécutifs (T. Watanabe, 1983). Nous rappelons tout d'abord les grandes lignes de cette approche.

Dans les systèmes de reconnaissance globale existants, l'unité de reconnaissance acoustique ou unité de décision est soit le mot, soit une unité de type syllabique telle que la syllabe, soit une unité de type phonémique telle que le diphonème (M. Decker, J.L. Gauvain, J. Mariani, 1985). De nombreux auteurs, comme J.E. Shoup (1980), se sont penchés sur le problème du choix entre ces solutions. Le fait que les variations contextuelles de part et d'autre d'un centre de syllabe soient limitées suggère de prendre comme unité de reconnaissance acoustique la portion de parole comprise entre deux

centres de syllabes² consécutifs, ce qui donne, pour le discours

Guy est ici. Sa voiture est en panne,

prononcé sans pause³, le découpage suivant :

[#gi - ie - eti - isi - isa - avwa - aty - yre - etã - ãpa - an#]

La seule contrainte combinatoire à laquelle obéit une suite d'unités de ce type est la contrainte "domino" : la deuxième voyelle d'une unité correspond à la première voyelle de l'unité qui la suit.

La raison évoquée pour choisir une unité de reconnaissance acoustique amène à choisir la même unité pour constituer les références servant à la reconnaissance acoustique. De plus, si ces références correspondent à des unités plus courtes que le mot, par exemple à des unités de type syllabique, elles forment un dictionnaire supplémentaire dont la taille n'est pas proportionnelle à la taille du vocabulaire reconnu : elle atteint un maximum de quelques milliers d'entrées et n'évolue plus guère lorsque le vocabulaire reconnu dépasse un certain seuil. Une liste pratiquement exhaustive des références distinctes peut donc être obtenue à partir de dictionnaires phonétiques couvrant l'ensemble de la langue courante. En d'autres termes, un examen extensif du lexique permet d'élaborer des données moins dépendantes du vocabulaire. Or, si la liste des références est rendue indépendante du vocabulaire, le dictionnaire de références ne dépend plus que de la langue et du locuteur. Cela facilite sa constitution et sa maintenance lors de l'évolution du vocabulaire ; de plus, cela accroît sa portabilité.

Après la reconnaissance acoustique, on dispose d'une suite ou d'un treillis d'unités de reconnaissance acoustique du type considéré. Ce treillis peut être transformé en un treillis de segments phonétiques qui lui est équivalent. Il reste alors à identifier les mots prononcés, en mettant le treillis en correspondance avec les mots d'un dictionnaire phonétique. Même pour une séquence phonétique donnée, plusieurs solutions sont souvent possibles : par exemple, à la suite [Oropeé] peut correspondre aussi bien *européen* que la séquence *eux roi paix un*. Le choix entre solutions concurrentes met en jeu des informations non phonétiques telles que la structure syntaxique du discours.

Le système comprend autant de références distinctes qu'il existe d'unités du type choisi susceptibles d'apparaître dans les discours à reconnaître. La liste de ces unités peut être établie pour l'ensemble de la langue courante à partir d'un dictionnaire phonétique approprié. Elle est alors pratiquement indépendante du vocabulaire connu par le système. Le calcul de cette liste met en jeu le découpage des mots du dictionnaire. Par exemple, le verbe *participer* fournit les unités [arti], [isi] et [ipe]. Les débuts de mots et les fins de mots, comme [pa] et [e] pour *participer*, requièrent un traitement spécial : *a priori*, toute fin de mot peut être combinée avec un début de mot pour former une unité,

² Plusieurs définitions de la notion de centre de syllabe sont envisageables, par exemple le centre temporel de la voyelle, le maximum d'énergie, ou le minimum de variabilité par rapport au temps. Par ailleurs, le découpage envisagé peut tenir compte non seulement des centres de syllabes mais aussi des silences, car les variations contextuelles de part et d'autre d'un silence sont limitées.

³ Plusieurs variations phonétiques peuvent intervenir par rapport à la prononciation donnée ici, en particulier pour ce qui est des deux liaisons facultatives et du timbre des deux [e]. Ces variations se répercutent sur la suite des unités à reconnaître.

comme [ipSwa] dans *type choisi*. D'autres combinaisons font intervenir des mots non syllabiques tels que *te, que, de, se, ce, je, le, me, ne, y* : ainsi [ɛski] figure dans plusieurs des prononciations possibles de *Qu'est-ce qu'il y a ?* Ces opérations fournissent une liste d'unités parmi lesquelles il reste à regrouper les unités identiques.

A partir de quel dictionnaire ce calcul peut-il intervenir ? Tout d'abord, on doit tenir compte des flexions : conjugaisons, pluriels, féminins. Comme ce sont des mots fléchis qui sont employés dans les textes, la liste des unités doit être calculée à partir d'un dictionnaire de formes fléchies. Par ailleurs, en cas de variations phonétiques, les différentes variantes qui apparaissent dans le discours se répercutent sur la suite des unités à reconnaître. Par exemple, la phrase *Guy repart* sera reconnue par l'intermédiaire de l'unité [irpa] si le *e* "muet" tombe, et des unités [ir0] et [0pa], ou [irœ] et [œpa], s'il est prononcé. Or, certaines unités n'apparaissent que par le jeu de variations phonétiques : c'est le cas de [ʃɛdy], observable dans la séquence *longue durée*. Pour établir la liste des unités, on doit donc tenir compte de toutes les formes, y compris les variantes. Le calcul de la liste des unités à partir du dictionnaire phonémique comprend ainsi la production automatique des variantes phonétiques, ce qui demande des connaissances précises sur les variations de ce type.

Cette recherche d'une plus grande précision devrait permettre de mieux définir les références, et contribuer à résoudre le problème du compromis entre le bruit (références reconnues à tort) et le silence (références non reconnues alors qu'elles devraient l'être) dans les performances du système. En effet, à défaut de données précises sur les variations phonétiques, on doit choisir un algorithme de reconnaissance acoustique plus tolérant aux dissemblances. Le problème revient alors à comparer des paroles bruitées ou déformées à des références imprécises à l'aide d'un algorithme approximatif. S'il est inévitable que le signal de parole à reconnaître puisse être bruité ou déformé, les deux autres parties du système (les références et l'algorithme de reconnaissance acoustique) peuvent être rendues plus précises, d'une part en affinant la description phonétique, d'autre part en utilisant des algorithmes de reconnaissance acoustique moins approximatifs, qui sont d'ailleurs plus efficaces.

Conclusion

Un système d'algorithmes et de données a été constitué. Les principaux éléments en sont récapitulés ci-dessous.

Les programmes de phonétisation automatique comprennent notamment un phonétiseur de mots inconnus. On soumet à ce phonétiseur une chaîne orthographique et il produit une ou plusieurs chaînes phonétiques qui correspondent aux différentes façons de la lire. Cet algorithme a été utilisé pour construire un logiciel d'aide à la correction orthographique.

Le dictionnaire DELAP donne la prononciation de 64.000 mots. Le dictionnaire DELAP-PL comporte des données formalisées sur les irrégularités de la correspondance entre l'orthographe et la prononciation.

Des programmes reproduisent les variations morphologiques des mots et, entre autres, calculent la prononciation des formes conjuguées des verbes.

Des programmes de comparaison et de contrôle servent à l'entretien et à l'évolution des dictionnaires.

Tous ces programmes et dictionnaires sont en relation avec les autres éléments du système de données linguistiques du LADL. C'est la première fois qu'un dictionnaire électronique consacré à la prononciation et comportant un nombre d'entrées comparable est intégré à un système de données orthographiques, grammaticales, morphologiques et syntaxiques.

Nous avons apporté un soin particulier à maintenir le système et à en faire la mise à jour. Les outils informatiques et les méthodes de travail sont en place pour la poursuite de ce développement.

Les programmes et les données ont été exploités dans deux applications informatiques : la génération de messages parlés à partir de représentations sémantiques, et la correction orthographique de mots.

De ces travaux se dégagent des conclusions et des directions de recherche.

L'exploitation industrielle d'un système informatique, quel qu'il soit, suppose que ce système satisfasse à des contraintes importantes. Dans le cas de systèmes de traitement automatique de la parole ou de la langue naturelle, nous pouvons discerner deux de ces contraintes et les prendre en compte.

Tout d'abord intervient une exigence de fiabilité. L'informatique linguistique vise à automatiser des opérations qui mettent en jeu des mots ou des textes, par exemple corriger des textes écrits, ou produire des messages parlés pour transmettre des informations. Cette automatisation n'a d'intérêt que si elle aboutit à des systèmes qui

effectuent ces opérations d'une façon fiable : que dire, par exemple, d'un correcteur orthographique qui corrigerait des erreurs dans certains mots mais en introduirait dans d'autres ? Nous avons essayé de réaliser des procédures et des applications informatiques qui répondent à un objectif de fiabilité. Les résultats obtenus montrent que cet objectif peut être atteint lorsque les algorithmes utilisent des données concrètes et explicites.

Par ailleurs, pour justifier une exploitation à grande échelle, les systèmes informatiques qui manipulent une langue naturelle devront incorporer des possibilités d'évolution et d'extension. Cela a des conséquences sur les données linguistiques utilisées par les programmes. Si ces données avaient été réunies à partir d'un corpus délimité *a priori*, elles ne seraient guère conformes qu'aux exemples du corpus, et les performances du système pourraient difficilement être étendues à de nouveaux exemples. En effet, en constituant les algorithmes et les dictionnaires, nous avons constaté que la prise en compte d'un grand nombre d'exemples rend caduques les conclusions qui se dégagent de l'examen de quelques exemples. La seule garantie que les programmes puissent évoluer est qu'ils soient construits à partir d'une étude préalable sur un large vocabulaire, sans *a priori* restrictifs sur le vocabulaire utile lors de l'utilisation. C'est pourquoi les études sur le lexique constituent un investissement important, apte à produire des résultats concrets, explicites et utilisables par des programmes. Ces conclusions s'appliquent également, à notre avis, à d'autres branches du traitement automatique des langues naturelles.

Nous ne savons pas construire des algorithmes qui manipulent des données non formalisées. En revanche, nous savons construire des algorithmes qui manipulent des données explicites. Ainsi, la prononciation des mots est une donnée indispensable, mais qui n'apparaît pas explicitement dans les textes et que, normalement, nous ne manipulons pas sous une forme explicite. Pourtant, c'est seulement sous une forme explicite qu'elle peut être traitée informatiquement. Nous avons donc besoin de réunir des données formalisées, par le moyen de descriptions formelles. Nous pensons avoir contribué à montrer que cette formalisation est réalisable, et que ses résultats peuvent être utilisés avec succès dans des applications informatiques.

Bibliographie

- J. ALLEN, 1973, "Speech Synthesis from Unrestricted Text", in J.L. FLANAGAN et L.R. RABINER (éd.), 1973, *Speech Synthesis*, Stroudsburg (Pennsylvanie) : Dowden, Hutchinson and Ross, pp. 416-428.
- J. ALLEN, M.Sh. HUNNICUTT et D. KLATT, 1987, *From text to speech. The MITalk system*, Cambridge : Cambridge University Press.
- J.P. BOONS, A. GUILLET et Ch. LECLERE, 1976, *La structure des phrases simples en français. Constructions intransitives*, Genève : Droz.
- A. CAILLEUX et J. KOMORN, 1981, *Dictionnaire des racines scientifiques*, Paris : CDU et SEDES.
- N. CATACH, 1984, *La phonétisation automatique du français*, Paris : CNRS.
- N. CHOMSKY et M. HALLE, 1968, *The Sound Pattern of English*, New York : Harper and Row.
- P.S. COHEN & R.L. MERCER, 1974, "The Phonological Component of an Automatic Speech-Recognition System", *Proceedings of IEEE Symposium on Speech Recognition*, pp. 177-188.
- C.H. COKER, N. UMEDA et C.P. BROWMAN, 1973, "Automatic Synthesis from Ordinary English Text", *IEEE Transactions on Audio and Electroacoustics* AU-21, pp. 293-397 ; in J.L. FLANAGAN et L.R. RABINER (éd.), 1973, *Speech Synthesis*, Stroudsburg (Pennsylvanie) : Dowden, Hutchinson and Ross, pp. 400-411.
- B. de CORNULIER, 1978, "Syllabe et suite de phonèmes en phonologie du français", in B. de Cornulier et F. Dell (éd.), *Etudes de phonologie du français*, Paris : CNRS.
- H. COTTEZ, 1985, *Dictionnaire des structures du vocabulaire savant*, Paris : Le Robert.
- B. COURTOIS, 1984, "Le DELAS : notice technique", *Rapport de recherches du LADL*, Université Paris 7.
- L. DANLOS, 1981, *Représentation informatique d'informations linguistiques : les constructions N être Prép X*, Thèse de troisième cycle, Université Paris 7.
- L. DANLOS, 1986, *Génération automatique de textes en langues naturelles*, Paris : Masson.
- L. DANLOS, F. EMERARD et E. LAPORTE, 1986, "Generation of Spoken Messages from Semantic Representations : A Semantic-Representation-to-Speech System", *Proceedings of COLING 1986*, Bonn.

L. DAVIDSON, 1962, "Retrieval of Misspelled Names in an Airline's Passenger Record System", *Comm. ACM*, 5, 3, pp. 169-171.

M. DECKER, J.L. GAUVAIN et J.J. MARIANI, 1985, "Reconnaissance de mots isolés par diphonèmes enchaînés", *Actes des 14^e Journées d'étude sur la parole*, Paris, pp. 271-274.

F. DELL, 1973, *Les règles et les sons, Introduction à la phonologie générative*, Paris : Hermann.

F. DELL et M. PLENAT, 1985, "Semi-voyelles et consonnes finales en français", *Rapport interne du GRECO "Communication parlée"*, Toulouse.

M. DIVAY et M. GUYOMARD, 1977, *Conception et réalisation sur ordinateur d'un programme de transcription graphémo-phonétique du français*, Thèse de 3^e cycle, Université de Rennes.

A. DUJOUR, 1987, *Conversion graphèmes-phonèmes avec variantes du français par règles*, Mémoire de DEA, Université Paris 7.

G. GOUGENHEIM, 1935, *Eléments de phonologie française. Etude descriptive des sons du français au point de vue fonctionnel*, Paris : Les Belles Lettres.

G. GOUGENHEIM, RIVENC, MICHEA et SAUVAGEAT, 1964, *L'élaboration du français fondamental*, Paris : Didier.

P. GOYER, D. DEGRYSE et B. GUERIN, 1979, "TROPs : un système de transcription orthographique phonétique et de synthèse du français", *Actes des 10^e Journées d'étude sur la parole*, Grenoble.

Gaston GROSS, 1986, "Typologie des noms composés", *Rapport de l'ATP "Nouvelles recherches sur le langage"*.

Maurice GROSS, 1968, *Grammaire transformationnelle du français : Syntaxe du verbe*, Paris : Larousse.

Maurice GROSS, 1975, *Méthodes en syntaxe*, Paris : Hermann.

Maurice GROSS, 1981, "Les bases empiriques de la notion de prédicat sémantique", *Langages* 63, Paris : Larousse, pp. 7-52.

Maurice GROSS, 1982, "Une classification des phrases "figées" du français", *Revue Québécoise de Linguistique*, vol. 11, n^o 2, Montréal : Presses de l'Université du Québec.

Maurice GROSS, 1986, *Grammaire transformationnelle du français : Syntaxe de l'adverbe*, Paris : Cantilène.

Z.S. HARRIS, 1951, *Methods In Structural Linguistics*, Chicago : The University of Chicago Press.

P. JUBAN, 1974, *Etude des types et fréquences syllabiques dans le français fondamental*, Rapport de fin d'année 1974, C'NET, Lannion.

M. LAMBERT et M. ROSSI, 1986, "Représentation et traitement des voyelles à timbre multiple", *Actes du séminaire "Lexiques et traitement automatique des langages" du GRECO "Communication parlée" et du GALF*, Université Paul Sabatier, Toulouse.

E. LAPORTE, 1984, *Transductions et phonologie*, Mémoire de DEA, Université Paris 7.

E. LAPORTE, 1986, "Applications de la morphophonologie à la production automatique de textes phonétiques", *Actes du séminaire "Lexiques et traitement automatique des langages" du GRECO "Communication parlée" et du GALF*, Université Paul Sabatier, Toulouse.

LE KIEN V., 1987, *Génération automatique de l'accord du participe passé*, Thèse, Université Paris 7.

A. MARTINET, 1933, "Remarques sur le système phonologique du français", *Bulletin de la Société de linguistique de Paris*, n° 34, pp. 191-202.

A. MARTINET et H. WALTER, 1973, *Dictionnaire de la prononciation française dans son usage réel*, Paris : France-Expansion.

Ph. MARTINON, 1913, *Comment on prononce le français*, Paris : Larousse.

E. MATER, 1966, *Deutsche Verben*, Leipzig.

F. NEEL, M. ESKENAZI et J.J. MARIANI, 1986, "Module de traduction phonétique avec variantes", *Actes du séminaire "Lexiques et traitement automatique des langages" du GRECO "Communication parlée" et du GALF*, Université Paul Sabatier, Toulouse.

G. PERENNOU et M. de CALMES, 1986, "BDLEX : une base de données et de connaissances du français parlé", *Actes du séminaire "Lexiques et traitement automatique des langages" du GRECO "Communication parlée" et du GALF*, Université Paul Sabatier, Toulouse.

G. PERENNOU et M. de CALMES, 1987, "Le traitement morpho-phonologique dans BDLEX 1", *Actes des 16^e Journées d'étude sur la parole*, Hammamet.

G. PERENNOU, F. LAHENS, P. DAUBEZE et M. de CALMES, 1986, "Rôle et structure du lexique dans le correcteur VORTEX", *Actes du séminaire "Lexiques et traitement automatique des langages" du GRECO "Communication parlée" et du GALF*, Université Paul Sabatier, Toulouse.

L. POIROT et X. RODET, 1976, "Transcription phonétique automatique du français", *Rapport du CEA 76-068*, Saclay.

B. PROUTS, 1979, "Traduction phonétique de textes écrits en français", *Actes des 10^e Journées d'étude sur la parole*, Grenoble.

M. SALKOFF, 1973, *Une grammaire en chaîne du français. Analyse distributionnelle*, Paris : Dunod.

E. SAPIR, 1933, "La réalité psychologique des phonèmes", *Journal de psychologie normale et pathologique*, n° 30, pp. 247-265.

S.S. SCHANE, 1968, *French Phonology and Morphology*, Cambridge (Mass.).

E.O. SELKIRK, 1972, *The Phrase Phonology of English and French*, Ph.D., MIT.

J.E. SHOUP, 1980, "Phonological Aspects of Speech Recognition", in Wayne A. LEA (éd.), *Trends in Speech Recognition*, Prentice-Hall, pp. 125-138.

M. SILBERZTEIN, 1987, "L'inversion de texte", *Rapport du Programme de recherches coordonnées "Informatique et linguistique"*, Université Paris 7.

D. TEIL, 1969, *Etude de génération synthétique de parole à l'aide d'un ordinateur*, Thèse, CNAM, Paris.

D. TEIL, 1974, "Prototype industriel d'une unité de réponse vocale autonome, à vocabulaire illimité et réponse immédiate", *Actes du 8^e Congrès international d'acoustique*, Londres.

G. TEP, 1979, "Système de génération des phrases phonétiques", *Actes des 10^e Journées d'étude sur la parole*, Grenoble.

G.L. TRAGER, 1944, "The Verb Morphology of Spoken French", *Language* 20, pp.131-141.

D. TREMBLAY, 1986, *Représentation sémantique et syntaxique de termes dans les dictionnaires électroniques*, Thèse, Université Paris VIII.

J. VERONIS, 1986, "Etude quantitative sur le système graphique et phono-graphique du français", *Cahiers de psychologie cognitive*, vol. 6, n° 5, pp. 501-531.

J. VERONIS, 1987, "Correction of Phonographic Errors in Natural Language Processing", *8th International Conference on Computers in the Humanities*, Columbia (South Carolina).

R. VIVES, 1985, "Mise en correspondance temporelle de descriptions phonologique et prosodique de mots dans le système de reconnaissance de la parole KEAL", *Actes des 14^e Journées d'étude sur la parole*, Paris, pp. 253-256.

L. WARNANT, 1964, *Dictionnaire de la Prononciation française*, Gembloux : Ducolot.

T. WATANABE, 1986, "Syllable Recognition for Continuous Japanese Speech Recognition", *ICASSP 86*, Tokyo, pp. 2295-2298.

T. WATANABE, 1983, "Segmentation-Free Syllable Recognition in Continuously Spoken Japanese", *ICASSP 83*, Boston, pp. 320-323.

Table des matières

Remerciements	1
Conventions et alphabets	2
Introduction	5
I. Les méthodes en phonétisation automatique	9
1. Les débuts de la phonétisation automatique	9
2. L'apparition de la phonétisation par dictionnaire	11
3. Les dictionnaires phonétiques récents	15
4. Un intérêt croissant pour les dictionnaires phonétiques	18
5. L'insuffisance des données phonétiques disponibles	21
II. Cadre général	25
1. Dictionnaires électroniques et dictionnaires usuels	25
2. Le dictionnaire DELAS	27
3. Le dictionnaire DELAP	30
4. Le phonétiseur de mots connus	31
5. Le phonétiseur de mots inconnus	41
6. L'exploitation d'un système de données linguistiques	42
7. L'élaboration du DELAS et du DELAP	44
III. Le dictionnaire DELAP	49
1. La structure du dictionnaire et des articles	49
2. Les entrées multiples	51
3. Les mots du DELAP	53
4. Brève introduction aux représentations phonétiques et phonémiques	54
5. L'utilisation de représentations phonémiques dans des applications informatiques	57
6. Le caractère arbitraire des représentations phonémiques	58
7. La lecture des formes phonémiques dans le DELAP	59
8. Le calcul de transcriptions phonétiques à partir de transcriptions phonémiques	68
9. Les groupes de consonnes phonétiques	69

IV. Les flexions : conjugaisons, pluriels, féminins	75
1. Les noms et les adjectifs	75
2. La conjugaison des verbes	80
3. Les notions de radical et de terminaison	81
4. Le radical	83
5. L'élément intermédiaire et la terminaison	86
6. L'élément intermédiaire	88
7. La terminaison	90
V. Variations phonétiques	97
1. Différences d'emploi des variantes phonétiques	97
2. Effacement de consonnes finales	100
3. Synérèse et diérèse : généralités	107
4. Variations libres entre synérèse et diérèse : les types <i>louer, tuer, lier</i>	109
5. Variations contextuelles entre synérèse et diérèse : les types <i>manier, il manie, plier</i>	127
6. Variations entre synérèse et diérèse à la limite de deux mots	133
VI. Applications informatiques	141
1. L'aide à la correction orthographique	141
2. La synthèse de la parole	146
3. La reconnaissance de la parole	149
Conclusion	155
Bibliographie	157
Table des matières	161