



Lexique, langue et tâche dans un dialogue homme-machine finalisé

Jean-Marie Pierrel, Laurent Romary

► To cite this version:

Jean-Marie Pierrel, Laurent Romary. Lexique, langue et tâche dans un dialogue homme-machine finalisé. *Sémiotiques*, 1997, 11, pp.95-117. hal-00521612

HAL Id: hal-00521612

<https://hal.science/hal-00521612>

Submitted on 5 Jan 2016

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution 4.0 International License

Lexique, langue et tâche dans le dialogue homme-machine finalisé

*Jean-Marie Pierrel⁺
Laurent Romary[°]*

1. Introduction

Dès que l'on parle du lexique, si tout le monde s'accorde pour lui donner une place de choix dans l'analyse ou le traitement du langage, le débat reste largement ouvert sur le type d'information et d'accès lexicaux à mettre en œuvre. Ce constat ne doit pas surprendre ; à notre sens il n'y a pas en effet de réponse globale ni de norme à imposer. La structure d'un lexique, tant du point de vue informationnel qu'opératoire, dépend fortement de la finalité qui préside à sa définition, même si les questions de fond que pose la constitution d'une composante lexicale restent les mêmes quel que soit son objectif — si l'on accepte du moins que le lexique doit regrouper l'ensemble des informations liées aux mots de la langue ou du langage utilisé.

“Accepter l'hypothèse d'un lexique [...] en tant que système dans lequel sont représentées toutes les informations concernant les mots de la langue conduit à soulever plusieurs questions importantes dont les suivantes :

1. Quelle est la nature de l'organisation de ce lexique ?
2. Quel est le format des représentations lexicales ?
3. Quelles sortes de procédures rendent possible l'accès à ces représentations ?” [Ségui, 1989].

Partant de cette considération deux voies s'offrent aux concepteurs de lexiques : soit s'orienter vers une définition très spécifique, voire *ad hoc*, d'un lexique donné, soit tenter de répondre aux trois questions ci-dessus en recherchant un maximum de validité tant du point de vue de la langue que de celui de l'application spécifique visée. Si la première approche privilégie l'efficacité en terme d'implantation informatique, la seconde est

indispensable pour assurer une pérennité et une réutilisabilité à toute composante lexicale. Pour notre part, dans le cadre du projet scientifique sur «Langue et Dialogue Homme-Machine» [Dialogue, 1996], c'est cette seconde voie qui nous guide même si, comme nous le verrons dans la suite, il n'est pas toujours aisé de définir un modèle propre intégrant dans un juste équilibre les contraintes langagières et applicatives.

Après une présentation rapide des contraintes imposées par notre perspective de dialogue homme-machine finalisé et de quelques problèmes qui se posent au niveau des descriptions infra-lexicales à inclure dans un lexique pour la reconnaissance automatique de la parole, nous nous focaliserons sur les rapports entre lexique, langue et tâche et leurs répercussions sur la composante lexicale de nos systèmes. Nous terminerons enfin en présentant les choix qui sont les nôtres dans la quête d'un modèle de lexique finalisé pour des dialogues langagiers homme-machine. Cela nous conduit, comme nous le verrons, à poser une certaine indissociabilité entre la sémantique et la pragmatique dans une composante lexicale — indissociabilité qui, trop souvent encore, fait défaut dans les bases de connaissances lexicographiques actuelles.

2. Lexique et dialogue homme-machine

2. 1. Place du lexique dans la compréhension de dialogues homme-machine

L'objectif essentiel des études que nous menons sur le dialogue oral homme-machine est de permettre à terme un dialogue le plus naturel possible entre un usager et un système automatisé [Pierrel *et al.*, 1990 ; Carré *et al.*, 1991]. En ce sens, nos recherches sont nettement finalisées, même si elles nécessitent d'aborder les problèmes fondamentaux liés à la langue orale, support de la communication. Nous ne prétendons pas étudier des systèmes capables de comprendre la langue dans son intégralité, néanmoins, même dans le cadre d'un univers pragmatique (ou application) restreint, réaliser un dialogue à forte composante langagière convivial et efficace nécessite de considérer diverses caractéristiques linguistiques pour que la communication puisse être réellement "naturelle". Plus qu'avoir une grande couverture de la langue (tant lexicale, syntaxique que sémantique), il faut qu'un système de dialogue homme-machine sache traiter des aspects aussi divers que :

— la perception et l'identification de phonèmes, la reconnaissance de phrases orales,

— l'interprétation contextuelle permettant de traiter les aspects de référence et de co-référence elliptique, anaphorique ou déictique,

— la gestion de la tâche ou de l'application.

La nécessité d'utiliser diverses sources d'informations pour mener à bien ces différents traitements est à la base d'un schéma théorique de système de compréhension (cf. **figure 1**) où le cœur du système, correspondant à une représentation interne de l'énoncé en cours de traitement, est une sorte de superviseur qui met en œuvre les diverses sources de connaissances S_i (acoustique, phonétique, lexicale, syntaxique, sémantique, etc.) grâce à un mécanisme propre à chaque type d'information M_i .

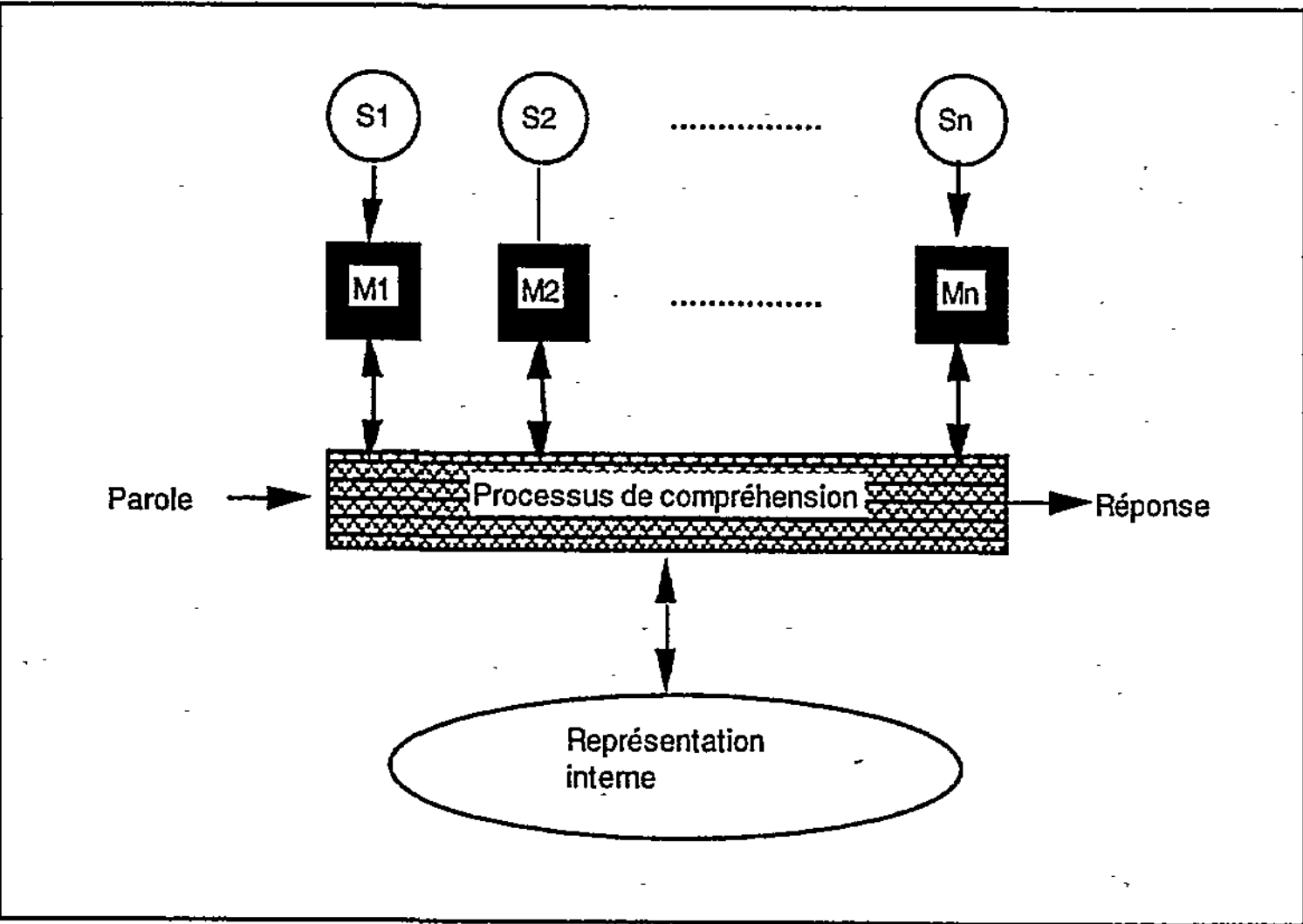


figure 1
Schéma théorique d'un système de compréhension

Dans de tels systèmes, le niveau lexical occupe une position clé comme pivot entre les niveaux linguistiques infra-lexicaux (en particulier phonétique et phonologique) et supra-lexicaux (syntaxique, sémantique) et la tâche ou l'application (pragmatique). Si nous laissons de côté les systèmes auto-organiseurs (neuromimétiques par exemple) ou statistiques (modèles markoviens) qui considèrent les énoncés en contexte comme autant de formes globales à reconnaître ou détecter, les stratégies de compréhension et de traitement les plus couramment admises dans les systèmes analytiques [Haton *et al.*, 1991] font en effet alterner des phases "ascendantes", du signal aux mots et des mots aux énoncés en contexte, et "descendantes", des énoncés en contexte aux mots et des mots au signal de parole. L'entité lexicale devient alors incontournable tant pour la

représentation des divers niveaux de connaissances que pour la structuration des traitements correspondants. Une des principales difficultés est donc de choisir le type de lexique à utiliser ; autrement dit, de définir quelles doivent être les informations à intégrer dans la composante lexicale d'un système de dialogue oral homme-machine [Pierrel, 1989].

Dans la suite, pour faciliter la présentation de notre approche, nous avons choisi de distinguer trois niveaux d'intervention du lexique (cf. figure 2) :

— le niveau *langagier infra-lexical* incluant phonétique, phonologie et morphologie,

— le niveau *langagier supra-lexical* essentiellement centré sur les structures syntaxico-sémantiques,

— le niveau *tâche ou application*.

¹Par langue finalisée ou artificielle, nous entendons un langage dont l'ensemble des structures lexicales, syntaxiques, sémantiques sont entièrement définies de façon spécifique. L'utilisation d'un tel langage passe donc nécessairement par une phase d'apprentissage assez longue de la part de l'utilisateur.

²Un langage opératif est un sous ensemble de la langue naturelle qui ne prend pas en compte la totalité des structures lexicales, syntaxiques et sémantiques dont dispose un locuteur de langue française, mais uniquement celles intervenant dans l'ensemble des énoncés qu'il peut être conduit à prononcer dans le cadre d'une situation de dialogue fortement finalisée. Un tel langage est généralement défini à partir d'études expérimentales de corpus de dialogues et de considérations plus linguistiques afin d'assurer la régularité langagière nécessaire à ce type de langage.

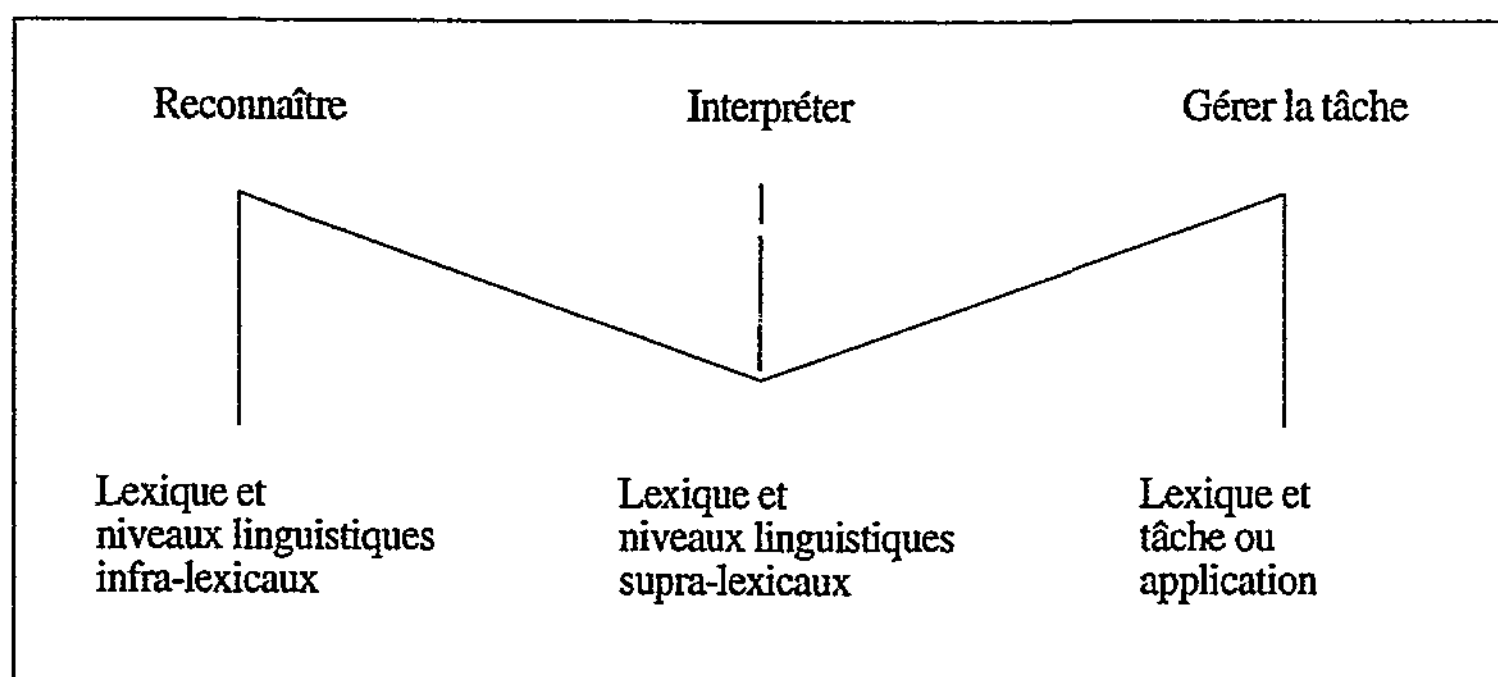


figure 2
Lexique et dialogue

Lorsqu'on a à définir une interface de dialogue sur une application donnée, la tâche et l'application sont entièrement définies et s'imposent comme un ensemble de contraintes externes prédéfinies. Ce n'est en effet pas du ressort du concepteur de l'interface homme-machine de modifier ou de faire évoluer la tâche à réaliser en termes de fonctions ou objets élémentaires. En revanche, c'est lui qui a la charge de définir le niveau langagier à introduire dans l'interface de dialogue. Deux grands choix s'offrent à lui : soit définir une langue finalisée ou artificielle¹ spécifique à l'application visée, soit prendre en compte le langage opératif mis en œuvre naturellement² par les utilisateurs de cette tâche.

Nos expériences, qui rejoignent en cela les conclusions de bon nombre d'ergonomes [Spérando, 1993], nous conduisent à penser que peu d'intermédiaires sont possibles entre ces deux choix. Si l'on opte pour une langue finalisée ou artificielle, il importe que, par sa structure même, cet

aspect artificiel s'impose toujours à l'utilisateur après une phase préalable d'apprentissage. Si ce n'est pas le cas, on risque d'être confronté, lors de son utilisation, à un glissement de cette langue vers certains aspects de la langue naturelle (lexicaux ou syntaxiques par exemple) et donc de sortir du strict cadre de la langue finalisée préalablement définie. A l'inverse, si on opte pour la prise en compte d'un véritable langage opératif, il convient de s'assurer pleinement de la complétude, de la cohérence et de la légitimité de ce langage par rapport à la langue et à l'application. Pour notre part nous pensons que, dès qu'on opte pour l'utilisation d'entrées lexicales (mots) et de structures syntaxico-sémantiques issues de la langue, seule l'utilisation de langages opératifs est possible et, comme nous allons le voir maintenant, cela influe considérablement sur le lexique.

2. 2. Lexique et connaissances langagières infra-lexicales

On peut distinguer deux aspects complémentaires au niveau infra-lexical :

— les connaissances morpho-lexicales : ce sont elles qui vont nous permettre de définir les unités d'accès et de codage du lexique : morphèmes, lexèmes ou mots, formes ;

— les connaissances phonétiques et phonologiques : dans le cadre d'une approche analytique de la reconnaissance automatique de la parole qui présuppose une reconnaissance (segmentation et identification) de segments minimaux de parole ou sons élémentaires (phonèmes, diphonèmes, syllabes) et une description des mots ou entités lexicales à partir de ces sons élémentaires, le niveau phonétique et phonologique est incontournable et une approche trop simpliste à ce niveau devient vite inopérante en reconnaissance automatique de la parole.

2. 2. 1. Composante morpho-lexicale

De façon générale, on distingue trois sortes de processus morphologiques : la composition, la dérivation et la flexion. Le problème qui se pose dans l'élaboration d'un lexique pour le dialogue oral finalisé homme-machine est alors de savoir s'il faut coder en machine l'ensemble des entités lexicales (ou formes) ou au contraire se limiter au codage des morphèmes de base et intégrer des procédures de reconstruction automatique des entités lexicales de surface à partir de ces morphèmes de base que sont les radicaux, préfixes, suffixes et flexions. Actuellement, hélas, nous ne disposons pas d'une théorie entièrement unifiée de la morphologie, suffisamment "calculable" pour permettre de définir et d'implanter en machine un analyseur morphématique complet et robuste. Pourtant, l'analyse morphologique est nécessaire pour limiter le plus possible la taille de nos lexiques.

La composition consiste à lier les éléments lexicaux pour obtenir une nouvelle entrée lexicale qui, comme le souligne Guy Pérennou [Haton, 1991], forme un tout non interprétable par la seule analyse de ses constituants. Dans un système de compréhension automatique, il nous faudra, compte tenu du contexte applicatif visé, considérer comme entrées lexicales spécifiques des expressions aussi diverses que : *pomme de terre*, *œil-de-bœuf*, *portemanteau*, *garde-robe*, etc. auxquelles il convient d'ajouter un certain nombre d'expressions figées qui, bien qu'analysables en composants élémentaires, seront le plus souvent considérées comme des entités lexicales uniques (par exemple les locutions adjectives, adverbiales, conjonctives, interjectives, propositionnelles ou verbales).

Si ce point de vue semble condamner tout essai d'analyse morphologique compositionnelle, il en va tout autrement lorsqu'on aborde la flexion ou la dérivation. Ainsi, pour les listes (1) et (2), faut-il considérer autant d'entrées lexicales que de mots ou au contraire ne conserver que les radicaux et l'ensemble des préfixes et suffixes ?

(1) observer, observé, observable, inobservable, observation, observatoire, observateur

(2) explorer, exploré, explorable, inexplorable, exploration, exploratoire, explorateur

Et comment faut-il prévoir dans un lexique l'apparition probable d'un mot tel "redénationalisation", rendue possible par un contexte pragmatique bien particulier, sans en considérer une analyse morphologique du type de celle donnée en (3) ?

(3) [re[de[[[nation]al]is]ation]]]

Si on ne veut pas conserver en mémoire l'ensemble des mots dérivés d'un même radical (en particulier par certains processus de dérivation, tels ceux mis en œuvre avec /re/ et /de/), il serait nécessaire de mettre en œuvre une "grammaire morphologique" qui, pour être opérationnelle et ne pas proposer de mots construits non attestés ou non interprétables, devrait inclure une forte composante sémantique voire pragmatique. De plus même si, comme c'est notre choix actuellement, on opte pour un codage manuel des divers dérivés d'un mot, un double problème apparaît : d'une part, dans un contexte applicatif donné, toutes les formes dérivées ne sont pas susceptibles d'être verbalisées, d'autre part il semble indispensable d'assurer une certaine régularité langagière et donc d'assurer, à contextes syntaxique, sémantique et pragmatique équivalents, une cohérence entre les dérivés de racines équivalentes. Quel utilisateur en effet comprendrait les mécanismes sous-jacents d'un système automatique qui accepterait *observer*, *observation* et *explorer* et refuserait *exploration*. Une solution peut conduire à définir des "types langagiers", ou classes d'équivalence

syntaxico-sémantico-pragmatiques, support de la cohérence dérivationnelle au sein d'un lexique donné. Concrètement, il faudrait introduire un type langagier correspondant à tous les mots susceptibles d'être dérivés de la même façon. D'où l'émergence de deux nouveaux problèmes : quels "types langagiers" utiliser dans un lexique et quelles procédures dérivationnelles attacher à un type donné ? Pour répondre proprement à ces deux questions, il est manifeste qu'il est nécessaire de justifier ou légitimer ces types et ces procédures par des considérations à la fois d'ordre linguistique (en aucun cas on ne peut transgresser les processus dérivationnels de la langue) et d'ordre applicatif (les limites dérivationnelles posées dans un lexique doivent pouvoir se justifier par des caractéristiques de la tâche et de l'application cible).

Quant aux flexions qui se distinguent des dérivations par leur caractère systématique et par l'interdiction qu'elles provoquent de toute nouvelle suffixation, si elles semblent à première vue plus simples à traiter (tout le monde s'accorde à dire qu'il est inutile de conserver dans le lexique l'ensemble des formes fléchies des verbes, des noms ou des adjectifs), il n'en demeure pas moins qu'elles sont fortement liées à des aspects syntaxiques et ont des incidences sur les aspects phonologiques (par exemple pour les liaisons).

Ainsi, dès ce premier niveau d'étude du lexique, on peut noter qu'une approche trop naïve, consistant par exemple à considérer chaque niveau du lexique comme indépendant des autres, conduit à une impasse et à se demander si seule une approche "unifiée" du lexique, allant de la phonétique à la pragmatique, est raisonnable.

2. 2. 2. Composante phonétique et phonologique

Dans un système de reconnaissance automatique de la parole, il est nécessaire pour le lexique d'inclure une représentation des mots sous forme de description phonétique standard telle que celle fournie dans tout bon dictionnaire. Mais coder dans le lexique une représentation phonétique de chaque mot ne suffit pas ; il est aussi nécessaire de prendre en compte les variantes phonologiques. On pourra trouver dans le chapitre 3 de [Haton, 1991] une présentation plus détaillée de ces aspects qu'il serait trop long d'exposer ici. La liste des altérations phonologiques pour une langue telle que le français est fort longue : assimilations, fusions, coarticulations, anticipations, etc. et leurs combinaisons fort complexes compte tenu du contexte d'énonciation. À ce niveau, il faut remarquer que les altérations phonologiques ne dépendent pas des seules transcriptions phonétiques théoriques des entrées lexicales mais aussi de leur catégorie syntaxique. Les nombreuses expériences de reconnaissances de parole que nous avons menées nous ont conforté dans cette optique. Ainsi *le* (lə) ou *un* ont un comportement phonologique foncièrement différent suivant qu'ils apparaissent en position de déterminant ou de

pronom... Cela tend à remettre en cause le niveau d'attachement des descriptions phonétiques des entrées lexicales qui, contrairement à ce qui est fait dans la plupart des dictionnaires, ne devrait pas être celui des entrées lexicales elles-mêmes mais plutôt celui des sous-catégories grammaticales, instanciations possibles de ces entrées lexicales. Quant aux descriptions du comportement phonologique de ces structures phonétiques, il nous faut reconnaître qu'elles furent trop souvent oubliées dans les descriptions lexicales et que nous manquons encore beaucoup de données expérimentales à ce sujet alors même que les règles de phonologie théorique ne rendent compte que partiellement des productions réelles des locuteurs.

Notons au passage que considérer comme acquis que l'unité de reconnaissance est et doit être le phonème relève d'une vision simpliste du problème remise en cause par une étude attentive de spectrogrammes de parole : la grande variabilité contextuelle des phonèmes montre qu'ils peuvent se traduire par différentes unités phonétiques et, qu'inversement, il est souvent difficile d'assigner un phonème à un segment élémentaire de parole. De nombreuses études tendent à prouver que la décision d'identification est fortement contextuelle, ce qui plaide pour l'utilisation d'unités plus larges : diphtongues, syllabes. On rejoint ici des résultats de psycholinguistique déjà anciens [Delattre, 1958 ; Ségui, 1980] qui ont montré l'importance des phénomènes contextuels en perception de parole. Dans l'état actuel de nos connaissances, ce problème du choix de l'unité de reconnaissance phonétique reste largement ouvert en reconnaissance automatique de la parole. Là encore on peut s'interroger à juste titre sur la légitimité de ne conserver dans nos lexiques qu'une représentation en termes de phonèmes même si cela demeure l'approche dominante aujourd'hui.

Un second choix se pose très rapidement lorsqu'on doit définir une composante phonétique et phonologique pour un lexique en reconnaissance automatique : faut-il coder toutes les formes phonologiques d'un mot ou au contraire ne coder qu'une représentation phonétique unique par mot à partir de laquelle grâce à des règles phonologiques seraient dérivées toutes les variantes phonologiques ? Si la seconde solution paraît être la plus satisfaisante intellectuellement, il faut reconnaître que les limites actuelles de nos connaissances des règles phonologiques nous conduisent souvent à opter pour la première solution.

2. 3. Lexique et connaissances langagières supra-lexicales

Si, comme nous l'avons montré, il existe des interactions fortes entre lexique et niveaux infra-lexicaux, les principales informations linguistiques constitutives d'un lexique portent en général sur les aspects supra-lexicaux que sont la *syntaxe*, la *sémantique* et la *pragmatique*.

L'absence de consensus sur la définition claire de ces trois niveaux de description linguistique³ nous amène à distinguer comme connaissances supra-lexicales nécessaires dans un lexique et décrites dans tout bon dictionnaire trois niveaux : les informations morphosyntaxiques, les informations structurelles (qu'on peut qualifier de syntaxico-sémantiques) et les définitions proprement dites et règles d'usage.

2. 3. 1. Informations morpho-syntaxiques

A priori, cela semble le niveau de description plus simple. Nous y incluons classiquement trois informations complémentaires : la classe grammaticale et, pour les noms, adjectifs, déterminants et pronoms, le genre et le nombre. Si ces deux dernières informations se discutent peu⁴, il en va tout autrement dès que l'on aborde le problème des classes grammaticales. Certes on peut facilement trouver un consensus sur un premier niveau minimal dans lequel apparaîtraient les classes majeures correspondant aux têtes de chapitres des parties du discours de [Grevisse, 1975] : *nom, article, adjectif, pronom, verbe, adverbe, préposition, conjonction, interjection*, mais chacun conviendra que ce type de classes est insuffisant car porteur de trop peu d'informations. Or, dès que l'on veut aller à un niveau de détail plus fin, se pose le problème difficile de la détermination tout à la fois du nombre de classes et de celles qu'il faut effectivement retenir et qualifier. Ainsi rien n'assure que dans le contexte d'un langage opératif particulier, le découpage de la classe des adverbes en adverbes de manière, quantité ou intensité, temps, lieu, affirmation, négation ou doute, tel que le propose Grevisse, soit opératoire. Il en va de même pour le redécoupage de l'ensemble des grandes classes précédemment listées.

Deux phénomènes doivent être pris en compte pour guider la détermination des classes morpho-syntaxiques les plus fines : le modèle syntaxique avec lequel devra interagir le lexique et la réalité du langage modélisé. À ce propos il convient de méditer la phrase de Littré dans la *Revue des Deux Mondes* du 1er juin 1842 : " le bon usage [...] doit être incessamment rajeuni aux sources vives dont il découle directement". Pour le concepteur de lexiques cela conduit à rechercher un juste équilibre dans l'introduction de classes syntaxiques entre celles issues de modèles théoriques de la langue et celles reflétant la réalité des corpus traités ou du langage opératif modélisé. On retrouve ici le vieux débat bien connu en linguistique entre compétence et performance qui s'applique aussi aux structures lexicales.

2. 3. 2 Informations structurelles ou syntaxico-sémantiques

Particulièrement importantes dans un lexique pour tous les lexèmes prédicatifs, les informations structurelles fournissent des données indispensables sur les structures régies par ces lexèmes. L'introduction de

³Il se pose un problème pratique de définition de ces niveaux de description linguistique... Où s'arrête la syntaxe ? Où commence la sémantique ? Faut-il, comme le proclame Chomsky à ses débuts poser une séparation hermétique entre syntaxe et sémantique ou, comme le suggère Fillmore, s'appuyer essentiellement sur des cas sémantiques ? Quant à la sémantique que l'on peut définir du point de vue linguistique comme la relation entre la forme des signes linguistiques, ou "signifiant", et ce qui est signifié, ou "signifiés", elle est trop souvent définie de façon négative comme regroupant tout ce qui n'est ni syntaxe, ni lexique, ce qui ne permet guère de résoudre le problème. Le plus souvent, on s'accorde à penser que la sémantique présuppose la syntaxe et le lexique, qu'elle regroupe l'ensemble des informations liées au sens, c'est-à-dire à la signification des mots et aux liaisons entre les mots autres que de type syntaxique.

⁴Pour le pluriel, on pourrait néanmoins distinguer le "dual" et le "véritable pluriel" comme le font bon nombre d'autres langues... de même qu'il convient de se poser la question du "neutre" pour le genre.

classes morpho-syntaxiques particulières (par exemple : verbes intransitifs, transitifs direct ou indirect avec telle proposition, verbes ditransitifs) répond en partie à cet objectif ; c'est d'ailleurs souvent ce type d'information qui est fourni dans un dictionnaire, complété par divers exemples. Une autre solution, utilisée dans un certain nombre de modèles lexicaux, consiste à associer à chaque lexème prédictif, directement ou par référence à un autre lexème prototypique, une structure précisant la construction possible des arguments : par exemple "*donner* x à y". Il reste alors à déterminer comment spécifier les variables correspondant aux arguments. Deux grandes approches peuvent être utilisées : soit à l'aide de classes syntaxico-sémantiques particulières soit à l'aide de classes morpho-syntaxiques associées à des traits sémantiques.

Dans les deux cas, un problème de choix se pose entre l'efficacité et la rapidité du traitement, qui font tendre vers la définition d'un nombre restreint de sous-classes ou de traits, et la précision de la description sémantique qui conduit à la définition d'un grand nombre de sous-classes ou de traits. Bien entendu, ces classes et traits ne doivent pas être choisis au hasard. Nous pensons que la légitimité de ces choix ne peut être que de deux ordres : soit ces classes sont générales et communes à plusieurs langages opératifs, donc issues de considérations purement langagières, soit au contraire spécifiques au langage opératif visé, c'est-à-dire issues d'une étude attentive de la tâche ou de l'application.

2.3.3. Définitions proprement dites et règles d'usage

Ce dernier point correspond à la plupart des informations que l'on trouve dans un dictionnaire de langue tel le *TLF*⁵, et pose un véritable problème aux concepteurs de lexique pour le traitement automatique. Que ce soit sous forme de définitions classiques ou de règles d'usage fournies le plus souvent par des exemples, ces informations deviennent vite inopérantes en traitement automatique du langage car elles pèchent par une formalisation insuffisante. Là encore deux solutions sont possibles : soit on ne définit ces informations que par référence à la tâche envisagée — et nous aborderons cet aspect dans le paragraphe suivant —, soit il convient de s'appuyer sur un modèle formel spécifique potentiellement applicable à l'ensemble de la langue.

Dans notre système DIAL, nous avons opté pour un modèle dérivé des grammaires fonctionnelles [Dik, 1979] dont une description détaillée est donnée dans [Deville, 1989]. Toute expression linguistique y était représentée par une structure comprenant un certain nombre de termes, définis par une fonction sémantique, et correspondant aux arguments d'un prédicat. Le prédicat, élément pivot de la phrase, correspondant à un lexème verbal, nominal ou adjectival, était lui même dérivé d'un nombre restreint de primitives prédictives qui déterminent ses caractéristiques et les relations entre ses arguments. La typologie des primitives prédictives

⁵ «*Trésor de la Langue Française : dictionnaire de la langue du 19^e et du 20^e siècle (1789-1960)*», réalisé par l'INaLF (Institut National de la Langue Française — CNRS) et publié chez Gallimard.

utilisées pour une application de renseignements administratifs, sur la base de deux paramètres de DYNAMISME (+/- DYN) et de CONTROLE (+/- CTR), faisait apparaître trois grandes classes :

- les états (STATE) : - DYN ;
- les actes (ACT) : + DYN et + CTR ;
- les processus : + DYN et - CTR.

L'approche fonctionnelle sémantique choisie nous permet de définir chacun des cas de notre système sur la base de cette typologie. Pour certains cas, une restriction sémantique s'impose : un terme devra alors satisfaire à une condition sémantique pour être candidat d'un cas spécifique. Les principaux cas utilisés dans notre système sont :

- AGENT (AGT) : entité contrôlante d' ACT / restriction sémantique : + animé
- PATIENT (PAT) : entité affectée par un PROCESS ou impliquée dans un ETAT / restriction sémantique : + animé
- BENEFICIAIRE (BEN) : entité à qui une autre entité est transférée lors d'un échange ou entité impliquée dans un état de possession (POSSESS)

Les autres cas sont : LOCATION (LOC), OBJET (OBJ), SOURCE (SCR), DIRECTION (DIR), MOYEN (MED), TEMPOREL (TEMP), INSTRUMENTAL (INST), MESURE (MES), BUT (GOAL), CAUSE (CAUSE), CONDITION (COND).

Compte tenu de ce modèle, notre lexique précise pour chaque mot, en le rattachant à un référent commun par classe sémantique, des informations syntaxiques catégorielles et, suivant les cas, des informations sémantiques sous forme de traits (**figure 3**) ou des informations sur la structure casuelle associée à la primitive dans le cas d'entité prédictive (**figure 4**).

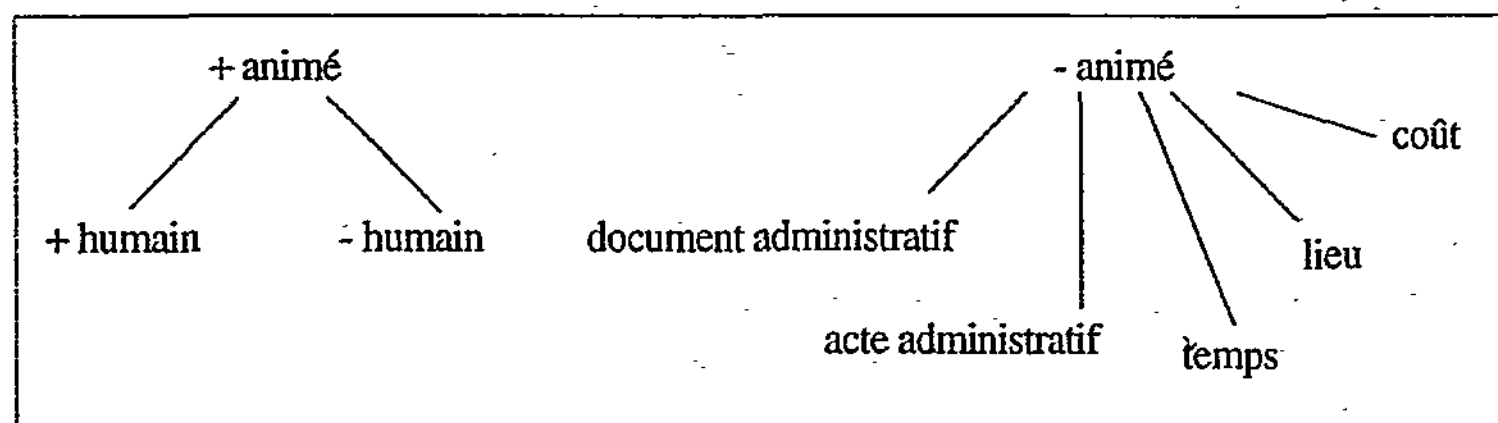


figure 3

Hierarchie des traits sémantiques
(application : renseignements administratifs)

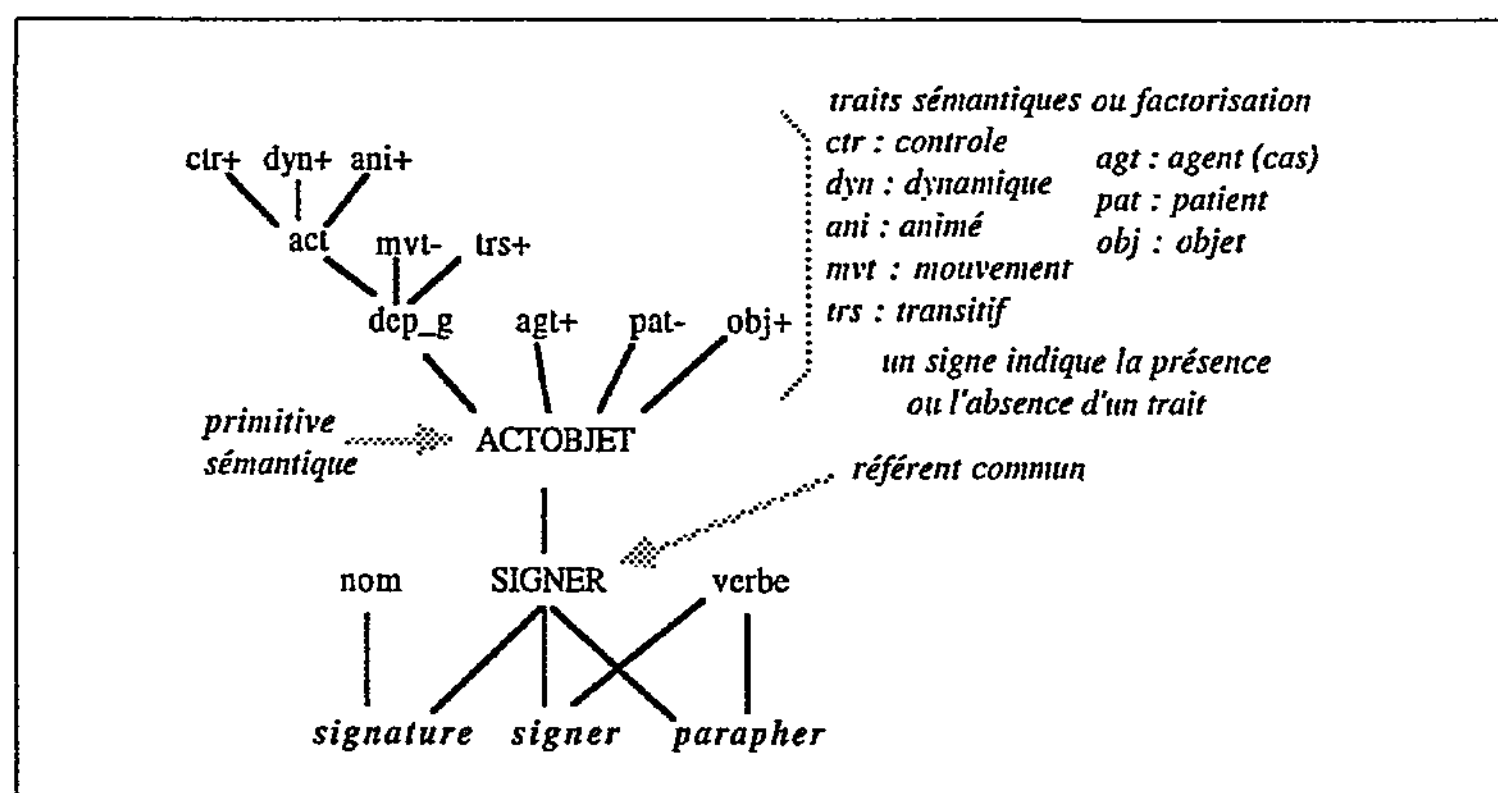


figure 4

Informations lexicales relatives à «signer»
et à certains de ses voisins (extraits)

2. 4. Lexique et tâche

Si l'on opte pour la définition des informations supra-lexicales détaillées ci-dessus pour une approche essentiellement langagière, la structure de notre lexique demeurera inopérante dans le cadre de compréhension de dialogue homme-machine. Pour pouvoir résoudre les problèmes de références, qui sont centraux dans ce cadre, il nous faudra enfin indiquer les liaisons entre les types langagiers et les types d'objets, de propriétés et d'actions de l'application. Cela revient donc à introduire au niveau de la sémantique lexicale une forte composante pragmatique structurée par la tâche ou l'application sous-jacente. Ces informations permettront de plus de filtrer les constructions syntaxico-sémantiques possibles. Ce sont de telles contraintes qui, dans un contexte de météorologie, conduisent à retenir, après *orage de*, des hypothèses telles que *pluie, neige, grêle...* et rejeter *soleil, température...* En fait, dans le cadre de langages opératifs, dès que nous parlons de sémantique, nous pensons essentiellement à la sémantique d'une application : c'est parce que nous pensons à une application météorologique, que nous rejetons "orage de soleil" ; en poésie, cela n'aurait certes pas rebuté un Boris Vian qui, au début de *L'Écume des jours*, écrit : "Son peigne d'ambre divisa la masse soyeuse en longs filets oranges pareils aux sillons que le gai laboureur trace à l'aide d'une fourchette dans de la confiture d'abricot".

Cette restriction de la sémantique à une application particulière n'est pas pénalisante dans le cadre de la compréhension de langages opératifs. Ces langages, qui apparaissent d'ailleurs dans une situation et un contexte d'application donnés de façon spontanée et progressive, sont en effet caractérisés essentiellement par des restrictions lexicales, tant d'un point de vue statistique (nombre de mots du lexique) que sémantique (liaison

signifiant-signifiés), et des restrictions syntaxiques (formes d'expression privilégiées, procédés syntaxiques spécifiques). Il apparaît donc qu'il n'est pas possible de se limiter aux seules informations linguistiques si le lexique doit s'intégrer dans un processus plus général qui ne peut se limiter à comprendre, c'est-à-dire à construire une représentation, mais doit interpréter les énoncés dans le contexte particulier de la tâche sous-jacente, c'est-à-dire le plus souvent exécuter une action de l'application. Il convient donc d'intégrer au mieux dans le lexique les aspects syntaxiques, sémantiques et pragmatiques.

3. Vers un lexique finalisé

Après avoir décrit les principales composantes qui, selon nous, sont indispensables à la gestion du lexique dans la mise en œuvre d'un système de dialogue homme-machine, nous nous devons d'observer différentes contraintes qui pèsent sur le lexique lui-même, dans la perspective d'en régulariser ses entrées. En effet, si l'on considère que toute opération de description lexicale est en définitive une abstraction par rapport au matériau écrit ou oral qu'elle est censée pouvoir représenter ou qu'elle doit permettre d'analyser automatiquement, la légitimation de ce même lexique passe par le degré de généralisation que l'on peut lui attribuer⁶. Concrètement cette légitimation *via* la généralisation peut se traduire par la réponse à la question suivante : dans quelle mesure un lexique L utilisé pour une application A1 peut-il servir de base pour décrire le lexique d'une nouvelle application A2 ? De fait, la question posée rejoint un peu le problème de la réutilisabilité d'un dictionnaire, mais le long de dimensions temporelles ou thématiques peut-être différentes. Ainsi un dictionnaire "généraliste" plus ancien va devoir être réactualisé pour suivre l'évolution d'une langue au fil du temps. Il peut aussi servir de référence pour la définition d'un dictionnaire spécialisé. Cependant, un dictionnaire spécialisé ne va pas directement fonder un autre dictionnaire dans un autre domaine par exemple. Ce serait pourtant une bonne analogie avec la situation d'un concepteur d'un nouveau système de dialogue qui ne souhaite pas redéfinir entièrement sa base lexicale. Regardons dans ce cadre plus précisément les paramètres que nous devons considérer avant d'essayer de trouver des réponses généralisables à d'autres activités tournant autour du lexique.

⁶Ceci faisant écho aux difficultés exprimées dans le paragraphe 2. 3. 2.

3. 1. Mots-vides et mots-pleins

L'une des premières intuitions qui vient à l'esprit lorsque l'on pense à "recycler" un lexique existant est de s'appuyer sur l'opposition entre mots "vides" (ou outils, ou grammaticaux) et mots "pleins". L'hypothèse qui

semblerait pouvoir être retenue serait que les premiers, comme articulateurs de l'énoncé, sont assez susceptibles de se retrouver d'une application particulière à une autre alors que les seconds, de par leur forte spécificité sémantique, seront caractéristiques d'un domaine d'application, ce qui ne leur permettrait pas d'être "recyclés" facilement. Dans l'ensemble, une telle affirmation est assez bien vérifiée pour un système de dialogue, mais elle appelle malgré tout un certain nombre de remarques qui traduisent un minimum de recouvrement entre ces deux notions.

Tout d'abord, dans la perspective de définir un lexique pour un système de dialogue homme-machine, il apparaît difficile d'utiliser la notion de sous-langage ou de langage opératif pour réduire le nombre d'entrées lexicales à considérer dans la catégorie des mots grammaticaux. Il semble en effet, comme cela a pu être observé dans de nombreuses expériences de simulations (par exemple [Mignot, 1995]), que l'on arrive très vite à une couverture quasi complète de ce type de vocabulaire. En particulier, puisque c'est là l'un des points essentiels dans la conception de tels systèmes, il ne semble pas que les expressions référentielles puissent être réduites à un sous-ensemble du système d'expressions utilisable en français par exemple. Ainsi, dans une application telle que Gocad, un logiciel de construction de surfaces géologiques, on rencontrera des indéfinis (*Je voudrais construire une surface*), des descriptions définies plus ou moins complexes (*Le fichier Vtop.pts*), des GN démonstratifs (*ces deux lignes*), des pronoms démonstratifs ou non (*celle-ci, il*), etc. Si on se place dans le contexte particulier d'une application graphique, pour laquelle on observera à la fois des références anaphoriques ou déictiques, nous pouvons être assurés que tout le système de détermination du français doit être pris en compte pour garantir une expression spontanée de la part de l'utilisateur. Ces observations garantissent pour une bonne part l'hypothèse qui nous intéresse ici, à savoir qu'il est possible, pour certains sous-systèmes lexicaux, de s'appuyer sur une description existante (et validée par exemple sur un premier corpus de dialogue), en vue de définir le support lexical d'une nouvelle application. Du point de vue de la légitimation du lexique correspondant, on peut dire qu'il s'agira d'une preuve par l'expérience, en supposant que la qualité des descriptions — associées à une bonne théorie pragmatique sous-jacente, mais c'est une autre histoire — ne fera que s'améliorer de système en système.

A l'inverse, la situation semble plus complexe dans le cas du vocabulaire des mots pleins, si l'on cherchait à reproduire une certaine description d'une application à une autre. De fait, il existe deux niveaux différents de difficulté, l'un lié à la non-superposabilité des lexiques lors du passage d'un domaine d'application à un autre, et l'autre, plus fondamental de la réutilisabilité d'une description lexicale quand bien même le premier point serait résolu. Sur ce premier point tout d'abord, il est clair qu'un mot tel que *diapir*, propre au domaine de la géologie, risque

de pas avoir à être défini dans un système de dialogue portant sur les renseignements administratifs, et inversement il est peu probable que le vocabulaire des pages jaunes soit d'une quelconque utilité pour un géologue. D'un point de vue lexical, il est donc nécessaire de décrire de façon spécifique un certain lexique permettant d'accéder à l'ensemble des concepts propres à la tâche considérée.

3. 2. Quelle légitimation ?

De façon pratique, une telle description est particulièrement difficile à légitimer. En effet, il n'existe jamais de source d'informations fiable pour fonder un tel travail. Une première stratégie — au demeurant couramment employée pour des applications de taille réduite — consiste à partir des concepts (objets et actions) de l'application et à "imaginer" le lexique permettant d'y faire référence. En général, les concepts sont faciles à cerner, mais, à moins d'arriver véritablement à se mettre en situation (ce qui est plus ou moins simple suivant le domaine...), il est très délicat de se fier à une expérience mentale si l'on vise une certaine exhaustivité. Une deuxième possibilité consiste généralement à mettre en œuvre une simulation du système de dialogue visé afin de recueillir un corpus d'énoncés "attestés". Sans nous attarder sur les éventuels biais liés à la mise en œuvre de telles expériences, on comprendra que l'exhaustivité est presque tout aussi difficile à atteindre que pour la première option. Il n'est en effet possible que de s'appuyer sur un échantillon réduit de locuteurs, ainsi que de scénarios exemplaires de la tâche visée, pour des raisons bien compréhensibles de temps (celui qu'il faut consacrer à la transcription des résultats).

Une dernière possibilité consiste à concevoir le système de dialogue de façon itérative, c'est-à-dire, partant d'une base lexicale initiale obtenue à l'aide des deux méthodes précédentes, confronter des versions successives du système à un groupe d'utilisateurs pour lesquels on observera toutes les situations où le système n'aura pas été en mesure de comprendre un énoncé donné du fait de l'absence dans son lexique d'un mot ou d'une expression particulière. Cette stratégie particulièrement lourde à mettre en œuvre dans les faits n'est encore que peu appliquée, notamment parce qu'il est très rare de pouvoir véritablement interagir avec des utilisateurs finaux. Cela implique que les projets de recherches et de développements correspondants puissent avoir une certaine durée, et la chose est encore assez rare. Pourtant, c'est *in fine*, la seule stratégie qui nous paraît fondée pour peu que l'on adhère à l'idée qu'il est vain d'imposer un lexique à un utilisateur, mais qu'au contraire, il faut savoir traiter le plus grand nombre possible de ses interventions langagières.

3. 3. Quels modèles de représentation ?

Sans entrer ici dans une discussion théorique qui confronterait tel ou tel modèle de description lexicale, nous pouvons malgré tout continuer à cerner un certain nombre de contraintes qui découlent de cette volonté d'offrir la plus grande liberté d'expression possible à l'utilisateur. Globalement, on observe vite qu'une simple vision de sémantique lexicale qui associerait à un mot un concept de l'application, un peu comme si on se mettait à coller des étiquettes sur les choses observées, n'est guère viable, même pour des applications relativement simples. Il ne s'agit pas d'une difficulté qui serait liée à une quelconque polysémie — qui pour nous n'est pas une véritable question — mais plus à la dynamicité de l'usage lexical dans un dialogue en langage naturel. Il est bien connu maintenant (cf. par exemple [Garrod, Gwyneth, 1994]) que tout dialogue implique une négociation implicite ou explicite entre les deux interlocuteurs sur les termes à employer pour désigner tel ou tel objet de la tâche courante. Ainsi, dans une situation d'aménagement intérieur, le fait qu'un locuteur ait employé *sofa* à la place de *divan* en première mention a de fortes chances d'entraîner que son allocataire suive la même voie. Un autre effet de localisation du lexique vient de l'usage très souvent contrastif associé aux principales expressions référentielles du français, comme nous avons pu le montrer à d'autres occasions ([Romary, Gaiffe, à paraître]). Ainsi, un locuteur ayant à désigner une cuisinière électrique dans un contexte où une cuisinière à gaz apparaissait en même temps sur l'écran, s'est vu employer le terme de *gazinière*, parce que, de son point de vue, il ne s'agissait que de contraster deux appareils (dont un réfrigérateur) faisant l'objet de son attention à ce stade de son interaction avec le système. De façon absolue, comme cela pourrait se concevoir dans une sémantique purement véri-conditionnelle, le terme employé ne devrait pas pouvoir référer à l'objet visé, et pourtant, la réalisation *contrastive* de l'acte de référence aboutit parfaitement dans ce type de situation.

On le voit, les contraintes issues de la mise en œuvre de systèmes de dialogue homme-machine font penser à des phénomènes qui en définitive sont généralisables pour la plupart à la langue dans sa globalité. Dans ce cadre, il peut sembler intéressant d'élargir la réflexion pour essayer de définir un outil un tant soit peu générique qui puisse fonder des études lexicales aussi bien dans le domaine du traitement automatique de la langue que dans celui de la lexicographie, et, au delà, de la linguistique.

4. Quelle description légitime pour des lexiques génériques ?

4. 1. Cadre général

Comme nous l'avons observé dans la section précédente toute réflexion sur le lexique et, plus précisément, sur la construction d'une base

lexicale, doit pouvoir s'appuyer sur un ensemble de données susceptibles d'initier et d'alimenter le processus de description. C'est typiquement ce qui se passe chaque fois que l'on conçoit un dictionnaire où le concepteur va s'appuyer soit sur des dictionnaires existants (cf. les mises à jour annuelles des dictionnaires grand public *Larousse* ou *Robert*), soit sur un corpus de textes, et l'on pensera immédiatement au *TLF* pour le français ou au dictionnaire *Cobuilt* pour l'anglais. Dans ce dernier cas, et probablement parce que la chose est un tout petit peu plus facile à concevoir, les corpus en question sont maintenant disponibles sous une forme électronique, et des initiatives récentes tendent même vers un certain degré de normalisation [Romary, Pierrel, à paraître]. Ceci est nettement moins vrai dans le domaine des lexiques ou dictionnaires informatisés, où, sans aller jusqu'à dire que le travail se fait encore à l'aide du papier et du crayon, il faut bien constater que la communauté académique (francophone en particulier) ne s'est encore que peu penchée sur la normalisation et la mise en commun de ressources lexicales. Les initiatives actuelles, et dont il faut apprécier l'effort, reposent finalement sur différents types d'intérêt. On peut ainsi distinguer :

— des résultats de travaux effectués dans un contexte supra-national, tels que ceux qui ont abouti à la mise en œuvre des lexiques «Multext» ;

— le travail mené au sein d'institutions telles que l'INaLF avec l'informatisation du *TLF* ;

— des initiatives individuelles, comme l'intéressant travail de transcription du "Basnage"⁷ qu'effectue actuellement Sébastien Pettoello par exemple.

Par ailleurs, on constate que dans de nombreux domaines tels que la muséographie, des besoins se font sentir pour établir et surtout maintenir des listes de termes ainsi que de noms propres permettant par exemple d'indexer ultérieurement des documents. Il semblerait donc intéressant de définir un cadre plus large qui permette de gérer ces différentes sources de connaissances, en y incluant les références à des corpus, de sorte que tout informaticien, lexicographe, historien, etc. puisse légitimer le travail effectué sur un mot particulier sur la base des connaissances accumulées autour de ce mot.

⁷Il s'agit d'un dictionnaire en trois tomes du XVIII^e siècle, qui est une édition revue et augmentée du «Dictionnaire Universel de Furetière» par le Sieur Basnage de Beauval.

4. 2. Vous avez dit normalisation ?

De fait, la définition d'une représentation normalisée qui s'applique tout aussi bien à la transcription des entrées d'un dictionnaire qu'au codage de lexiques informatisés a déjà fait l'objet d'un certain nombre d'études au sein de la communauté internationale. Le résultat le plus notable est sans aucun doute l'ensemble des directives définies dans le

⁸Dans un document SGML, une balise telle que `<entry>` représente le début de l'élément, et `</entry>` en représente la fin. Des couples attribut=valeur, insérés dans la balise ouvrante d'un élément permettent d'apporter des informations supplémentaires au niveau de la représentation considérée (par exemple pour fournir un numéro d'ordre à une suite d'éléments de type `<entry>`). Les caractères accentués sont parfois codés à l'aide d'entités SGML (exemple `é` sera codé `é`). L'exemple est une transcription de l'entrée *remoulade* du «Collins Robert French-English English-French Dictionary» (Beryl T. Atkins et al., London, Collins, 1978, rpt. 1983).

cadre de la TEI (*Text Encoding Initiative*) qui comprennent différentes sections spécifiquement consacrées au lexique. Sans entrer dans le détail technique du format adopté, qui s'appuie sur la norme SGML (Standard Generalized Markup Language) de représentation de documents structurés, nous pouvons en donner les principales caractéristiques. Tout d'abord, il s'agit d'une description principalement arborescente des données reposant sur des *éléments* identifiés emboîtés les uns dans les autres. Une entrée de dictionnaire pourra ainsi avoir la forme suivante (extrait des directives de la TEI)⁸ :

```
<entry>
  <form>
    <orth>r&eacute;moulade</orth>
    <pron>Remulad</pron>
  </form>
  <gramGrp>
    <pos>n</pos>
    <gen>f</gen>
  </gramGrp>
  <trans>
    <tr>remoulade</tr>
    <tr>r&eacute;moulade</tr>
    <def>dressing containing mustard and herbs</def>
  </trans>
</entry>
```

où l'on distinguera des informations morphologiques (sous l'élément `<form>`), des informations grammaticales (`<gramGrp>`) et des informations correspondant à la traduction et à la définition dans la langue cible du mot (`<trans>`).

La description de telles entrées s'appuie sur une syntaxe générique du document correspondant, indiquant quelles peuvent être les possibilités d'occurrence de tel élément dans tel autre. Cette description du type de document ou DTD (Document Type Definition dans la norme SGML), donne une grande souplesse de représentation d'informations structurées, telles que celles contenues dans un dictionnaire.

Pour illustrer ce point, nous pouvons reprendre le problème, mentionné en 2. 2. 2., de l'attachement des descriptions phonétiques à tel ou tel niveau de description lexicale. L'avantage d'adopter une structure de représentation lexicale qui soit modulable, c'est-à-dire qui n'impose pas la présence d'une information spécifique à une place parfaitement spécifiée au départ, permet de traiter les cas particuliers de façon élégante. Ainsi, on peut tout aussi bien attacher une prononciation générale pour un mot donné :

```
<entry>
<form>
```



```

<orth>[LEX1]</orth>
<pron>[PRON_LEX1]<pron>
</form>
<sense n=1>...</sense>
<sense n=2>...</sense>
<entry>

```

ou introduire une variante de prononciation associée pour chaque sens spécifique :

```

<entry>
<form>
<orth>[LEX1]</orth>
</form>
<sense n=1>
<pron>[PRON_LEX1_var1]<pron>
...
</sense>
<sense n=2>
<pron>[PRON_LEX1_var2]<pron>
...
</sense>
<entry>

```

Ainsi, pour peu que la structure abstraite de document (ou DTD) accepte une telle variante, le problème de l'attachement d'une description particulière à différents niveaux de représentation se posera de façon bien moins cruciale.

4. 3. Un méta-lexique générique.

Essayons de rêver quelques instants. Supposons que l'ensemble des dictionnaires ou des bases de données lexicales informatisés adoptent un schéma de représentation aussi proche que possible de celui décrit ci-dessus. Nous n'en sommes pas si loin si nous considérons par exemple les trois cas cités dans la section 4. 1. qui — peu ou prou (nous ne détaillerons pas...) — s'inspirent justement de la norme SGML ainsi que des directives de la TEI. Si donc cette pratique venait à se répandre, il serait tout à fait concevable de définir une base lexicale généralisée qui, suivant le mot ou l'expression recherchés, puisse pointer sur les entrées correspondantes des ressources lexicales normalisées qui auraient été recensées. Une telle base ne serait bien évidemment pas un lexique universel, car il est clair — l'étude effectuée dans cet article sur la base de notre expérience en matière de dialogue homme-machine le montrant bien — que l'on ne peut mélanger n'importe quelle source d'information relative à un mot avec n'importe quelle autre sans prendre un minimum de précautions. Cependant, elle permettrait de positionner chaque mot ou

expression dans un contexte temporel et spatial clairement identifié, suivant les références auxquelles il serait associé. Au passage, il serait non seulement nécessaire de pointer sur des bases lexicales existantes, mais aussi sur des occurrences du mot dans différents corpus textuels, pour peu que l'on ait pris la peine de gérer ceux-ci en parallèle avec ceux-là. L'expérience de dictionnaires tels que le *TLF* montrent bien qu'il est nécessaire d'accéder à une base d'exemples pour illustrer la description d'un mot. Dans d'autres cas, ceci permettrait de décrire des mots qui ne sont que des hapax non référencés dans des dictionnaires, mais parfaitement identifiés dans une œuvre particulière (la littérature de l'Oulipo est truffée d'exemples de ce type). Enfin, cette base lexicale doit posséder une certaine souplesse pour que tout utilisateur puisse lui adjoindre — avant de rendre l'information publique — ses propres références lexicales, telles que des listes de noms propres, de termes techniques, etc.

Au sens de la normalisation décrite précédemment, notre méta-lexique aurait une structure particulièrement simple : chaque entrée ne comprendrait que la forme orthographique du mot ou de l'expression, éventuellement quelques informations morpho-syntaxiques simples, et surtout une suite de références externes (l'élément `<xptr>` de la TEI semble bien adapté à cet usage) sur les entrées de lexiques connus et/ou des lieux d'occurrence du mot ou de l'expression.

5. Conclusion

Nous avons essayé dans cet article d'apporter quelques éléments à la réflexion menée au sein de ce numéro spécial sur la légitimation dans le lexique en nous appuyant sur l'expérience spécifique d'une équipe de recherche qui depuis près de vingt ans travaille sur la définition de systèmes de communication homme-machine en langage naturel. Le constat que nous avons fait n'est *a priori* guère optimiste si l'on considère la grande variété des informations à prendre en compte, ainsi que la difficulté à fonder un lexique informatisé sur des données identifiables et fiables, notamment dans le domaine du dialogue homme-machine. Cependant, il nous est apparu nécessaire d'élargir la réflexion pour essayer de cerner, même si c'est sous la forme d'un schéma général qu'il faudra vite valider, ce que pourrait être une base lexicale de référence qui permette à chaque chercheur, quelle que soit sa discipline et *a fortiori* son cadre théorique, de mieux asseoir son propre travail sur les mots et leur usage. Nous pensons que la réalisation d'une telle base est une chose faisable. Elle n'est en effet que peu limitée par des contraintes techniques dans le contexte d'un accroissement de la puissance des machines que nous utilisons et de l'accès maintenant répandu aux réseaux informatiques.

Ce faisant, il est nécessaire que chaque chercheur prenne conscience que la richesse de cette base viendra de ce qu'il pourra y apporter lui-même, en considérant qu'une ressource lexicale particulière a peu de valeur scientifique si elle n'est pas mise à la disposition de la communauté la plus large possible, dans le respect du travail de chacun.

(⁺ CRIN-CNRS
^o Inria-Lorraine)

Références

CARRÉ (R.), DÉGREMONT (J. F.), GROSS (M.), PIERREL (J. M.), SABAH (G.)

1991, *Langage humain et machine*, Paris, Presses du CNRS.

DELATTRE (P.)

1958, "Les Indices de la parole", *Phonetica*, vol. 11, n° 1 et 2.

DEVILLE (G.)

1989, *Modelization of Task-Oriented Utterances in a Man Machine Dialogue System*, Thesis of Computational Linguistics, University of Antwerpen.

[Dialogue, 1996]

1996, "Dialogue homme-machine à forte composante langagière", p. 249-270, in *Rapport d'activité scientifique 1994-1995*, CRIN, Nancy.

DIK (S.)

1979, *Functional Grammar*, Amsterdam, North Holland.

FALZON (P.)

1984, "Les Langages opératifs", p. 364-383, in *Dialogue Homme-Machine à composante orale*, J. M. Pierrel et al., eds., GRECO-CNRS (Communication Parlée).

GARROD (S.), GWYNETH (D.)

1994, "Conversation, Co-ordination and Convention : An Empirical Investigation of How Groups Establish Linguistic Conventions", *Cognition*, 53, p. 181-215.

GREVISSE (M.)

1975, *Le Bon usage*, 2ème éd. revue, Paris, Duculot.

HATON (J. P.), PIERREL (J. M.), PERENNOU (G.), CAELEN (J.), GAUVAIN (J. L.)

1991, *Reconnaissance automatique de la parole*, Paris, Dunod Informatique.

MIGNOT (C.)

1995, *Usage de la parole et du geste dans les interfaces multimodales : étude expérimentale et modélisation*, doctorat d'Université, Université Henri Poincaré, Nancy I.

PIERREL (J. M.)

1989, "Lexique et compréhension automatique de la parole", *Lexique (P. U. L.)*, 8, p. 137-165.

PIERREL (J. M.), CARBONELL (N.), HATON (J. P.), SMAÏLI (K.)

1990, "Vers une meilleure intégration de la parole dans des systèmes de communication homme-machine", *Traitement du Signal*, vol. 7, n° 4, p. 327-344.

ROMARY (L.), GAIFFE (B.)

(à paraître), "Langage, dialogue finalisé et cognition spatiale", in *Langage et cognition spatiale*, M. Denis, ed., Paris, Masson.

ROMARY (L.), PIERREL (J.-M.)

(à paraître), "Le projet *Silfide* : vers un accès ouvert aux ressources linguistiques francophones", *Revue Française de Linguistique Appliquée*.

SÉGUI (J.)

1980, "Code phonétique et code phonologique dans la perception de la parole : le rôle de la syllabe", *Journée d'étude codage phonétique et phonologique*, Ecole des Hautes Etudes.

1989, "L'Accès au lexique, données expérimentales et modèles", p. 215-230, in *La parole et son traitement automatique*, Paris, Calliope, Masson (CNET-ENST).

SPÉRADIO (J. C.)

1993, *L'Ergonomie dans la conception de projets informatiques*, Toulouse, Octares Editions.

