

Un survol des recherches en génération automatique

Laurence Danlos

Volume 14, Number 2, 1985

Linguistique et informatique

URI: <https://id.erudit.org/iderudit/602539ar>

DOI: <https://doi.org/10.7202/602539ar>

[See table of contents](#)

Publisher(s)

Université du Québec à Montréal

ISSN

0710-0167 (print)

1705-4591 (digital)

[Explore this journal](#)

Cite this article

Danlos, L. (1985). Un survol des recherches en génération automatique. *Revue québécoise de linguistique*, 14(2), 65–102. <https://doi.org/10.7202/602539ar>

UN SURVOL DES RECHERCHES EN GÉNÉRATION AUTOMATIQUE

Laurence Danlos

L'utilisation du langage naturel comme moyen de communication entre l'homme et la machine requiert en premier lieu que l'ordinateur comprenne le message de l'utilisateur, ensuite qu'il détermine comment répondre à ce message, enfin qu'il formule la réponse dans la langue désirée par l'utilisateur. La première tâche est du ressort de l'analyse automatique, la deuxième d'un module de raisonnement, et la troisième de la génération automatique. Ces modules sont récapitulés dans la figure 1. Chacun d'eux représente un ou plusieurs domaines de recherche, dans des formes simplifiées, ce sont de gros programmes. Les tâches effectuées en analyse automatique et en génération automatique ne se limitent pas au cas d'un dialogue. D'une manière plus générale, tout système qui produit une représentation formelle reflétant une analyse d'un texte est du domaine de l'analyse automatique, et tout système qui produit un texte à partir d'une représentation formelle est du domaine de la génération automatique.

Historiquement, la génération automatique est un domaine beaucoup plus récent que l'analyse : depuis les années 1950, de nombreux programmes d'analyse ont vu le jour, tandis que la génération automatique ne s'est développée que depuis moins d'une dizaine d'années. Mais la génération automatique est un domaine en plein essor, le besoin de produire des textes élaborés se faisant de plus en plus pressant : interrogation d'une banque de données, résultats d'un système expert, bureautique, etc. Nous allons présenter un certain nombre de travaux en génération automatique en essayant de les situer les uns par rapport aux autres.

Les premiers travaux de génération automatique avaient pour but de tester une grammaire (Yngve 1961 et Freidman 1969). Ainsi, le programme de Friedman était destiné à servir d'aide aux linguistes pour l'étude des relations transformationnelles de la grammaire générative (Chomsky 1957).

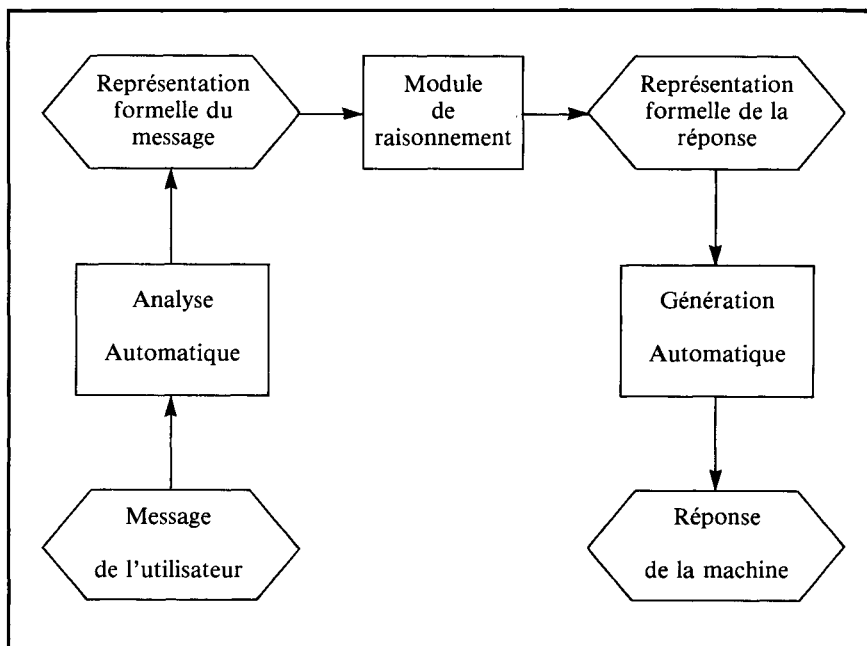


Figure 1. Dialogue entre l'homme et la machine

Il ne s'agissait pas de transmettre un message mais de produire des phrases «au hasard» pour vérifier des règles de grammaire. Ce type de recherche a été délaissé au profit de travaux sur la production de phrases véhiculant un certain contenu informatif. De ce fait, la détermination du contenu informatif du message à produire (la question : **Quoi dire?**) est du domaine de la génération automatique autant que la production d'un texte véhiculant ce contenu informatif (la question : **Comment le dire?**).

1. **Quoi dire? / Comment le dire?**

Examinons en quels termes se pose le problème du contenu informatif du message à produire selon les applications d'un générateur.

Pour un générateur intégré dans un système de traduction automatique, le contenu informatif du texte produit doit être identique à celui du texte source. Le générateur n'a donc pas, en principe, à traiter la question **Quoi dire?** Par contre, lorsqu'il s'agit de produire automatiquement un résumé, on doit extraire du texte source (ou de sa représentation) les informations essentielles; autrement dit, on doit avant tout déterminer le contenu du

résumé. Un sérieux effort dans cette direction a été réalisé par (Lehnert 1981) pour des résumés d'histoires à caractère psychologique, c'est-à-dire d'histoires mettant en jeu les «états d'âme» de certaines personnes. En considérant simplement trois états psychologiques (positif, neutre, négatif), Lehnert construit des graphes dont les noeuds représentent les états psychologiques des protagonistes, et les branches les relations causales entre ces états. Elle a établi des classes d'équivalence entre de tels graphes et réalisé des procédés algorithmiques pour en dégager des noeuds particuliers. Ces noeuds déterminent les pivots de l'histoire, pivots qui constitueront les éléments du résumé.

Un générateur peut aussi être greffé sur un système expert pour exprimer en langue naturelle les résultats du système expert et fournir des explications sur ces résultats. Une méthode de synthèse d'explications consiste à traduire en langue naturelle les règles du programme qui ont été exécutées pour arriver aux résultats. Cette méthode, qui revient à fournir une trace en langue naturelle de l'exécution du programme, est critiquée par Swartout (1981) car elle permet d'apporter des éclaircissements sur le fonctionnement du programme mais elle ne donne aucune justification sur ce fonctionnement. Le problème vient de ce que les connaissances qui sous-tendent la réalisation d'un programme ne font pas obligatoirement partie du programme lui-même. Swartout met donc l'accent sur la nécessité d'intégrer un modèle du domaine — les connaissances d'un expert dans le domaine — aux heuristiques du domaine — l'utilisation de ces connaissances pour réaliser un système expert. Il a utilisé ces deux types de données dans un système expert médical pour déterminer le contenu informatif des explications sur les diagnostics.

Un système de génération peut aussi avoir pour tâche de produire les réponses dans un système d'interrogation de banque de données. La première opération à effectuer consiste alors à sélectionner les éléments de la banque de données qui doivent être exprimés dans la réponse. Comme l'indique McKeown (1982), les informations véhiculées dans la réponse dépendent du type de question posée : les réponses aux questions

Qu'est-ce-que X?

= : Qu'est-ce-qu'une baleine?

Quelle est la différence entre X et Y?

= : Quelle est la différence entre une baleine et un marsouin?

ne contiendront pas les mêmes informations sur X : la réponse à la première question contiendra toutes les caractéristiques de X, tandis que la réponse à la seconde ne contiendra que les caractéristiques de X qui opposent X à Y. De plus, le contenu informatif de la réponse doit être adapté à l'interlocuteur à qui l'on s'adresse. En particulier, la réponse doit être construite en fonction des connaissances de l'interlocuteur : le dialogue

- Qu'est-ce-qu'une sole?
- Une sole est un pleuronectidé...

ne se continuera par la description d'un pleuronectidé que si celle-ci est supposée inconnue de l'interlocuteur.

La situation est analogue lorsqu'un système de génération doit relater des événements. Il faut en premier lieu déterminer les informations pertinentes pour le type d'événement. Ainsi, les conditions météorologiques sont généralement de peu d'importance lorsqu'il s'agit d'un attentat, alors qu'elles sont du plus haut intérêt pour un match de football. De plus, il faut tenir compte de l'interlocuteur auquel on s'adresse. Comme le soulignent Hovy et Schank (1984), un attentat terroriste en Irlande ne sera pas relaté de la même façon dans un journal catholique et dans un journal protestant; les textes suivants, qui relatent le même événement, peuvent apparaître respectivement dans un journal catholique ou protestant :

- (1) Deux occupants britanniques ont été blessés par des membres de l'IRA hier à Belfast.
- (2) Deux soldats britanniques ont été blessés par des terroristes de l'IRA hier à Belfast.

Les récits (1) et (2) diffèrent dans les termes employés pour désigner les auteurs de l'attentat (*membres/terroristes de l'IRA*) et ses victimes (*occupants/soldats Britanniques*). Ces différences sont dans le domaine de la référence. La référence soulève de nombreux problèmes délicats, par exemple le suivant : dans quelle mesure peut-on considérer que les récits (1) et (2) ont la même signification (en admettant qu'il y ait consensus sur la définition de «signification»)? Ce problème a fait l'objet de nombreux débats philosophiques. En génération automatique, il se ramène à décider si la manière de référer à des individus ou objets est dans le champ de la question **Comment le dire?** (si on considère que (1) et (2) ont la même signification), ou dans celui de la question **Quoi dire?** (si on considère que (1) et (2) n'ont pas la même signification). Le fait que ce problème ne reçoit

pas de solution tranchée peut amener à remettre en question l'indépendance des questions **Quoi dire?** et **Comment le dire?** En effet, jusqu'à présent, nous avons supposé que ces questions étaient traitées dans deux modules distincts : un module de raisonnement traite la question du contenu informatif du message à produire, et la sortie de ce module est traduite en langue naturelle par un module indépendant (cf. figure 1). Cette position, admise par la plupart des chercheurs en génération automatique, est critiquée par Appelt (1982). Le travail d'Appelt se situe dans le cadre d'un dialogue entre deux personnes en présence. Appelt s'intéresse à la production de phrases en tant que moyen d'action pour accomplir un but. Cette perspective, inspirée principalement d'Austin (1962), place la formulation d'un énoncé sur le même plan que les actions physiques, par exemple se déplacer, déplacer un objet ou montrer du doigt un objet. Pour illustrer le système d'Appelt, imaginons que deux personnes, Luc et Marie, soient dans la même pièce, qu'il y ait un réveil dans cette pièce et que celui-ci soit caché aux yeux de Luc par Marie; supposons que Luc désire savoir l'heure. Pour satisfaire ce désir, Luc a alors au moins deux solutions : la première consiste à se trouver en position de voir le réveil, la seconde à demander l'heure à Marie. Pour réaliser la première solution, Luc peut se déplacer, déplacer Marie ou demander à Marie de se déplacer. Au niveau des actions que Luc peut entreprendre, il n'y a donc aucune raison de distinguer actes physiques (e.g. déplacer Marie) et actes de parole (e.g. demander à Marie de se déplacer). Si Luc opte pour un acte de parole (en demandant à Marie de se déplacer ou en lui demandant l'heure), cet acte de parole a une «force illocutoire directive» (Searle 1979). Il peut être réalisé de diverses manières : en donnant un ordre, par exemple au moyen d'un impératif (*Déplace-toi, Dis-moi l'heure*); en exprimant un souhait, par exemple en commençant par *Je voudrais* (*Je voudrais que tu te déplaces, Je voudrais que tu me dises l'heure*); en posant une question directe (*Quelle heure est-il?*) ou indirecte (*Peux-tu te déplacer?, As-tu l'heure?*). L'entrée du système d'Appelt est donc la représentation d'un but (e.g. savoir l'heure), la sortie indique un acte physique (e.g. se déplacer) et/ou l'énonciation d'une phrase (e.g. *Quelle heure est-il?*). Le programme examine les différentes solutions permettant d'accomplir le but et les différents moyens permettant de réaliser ces solutions. Il repose sur un modèle de raisonnement logique et sur une base de données décrivant la situation («l'état du monde») et les connaissances des deux personnes en jeu.

À propos du problème de la référence, Appelt prend en compte la possibilité de montrer du doigt des objets. Ainsi, pour informer l'interlocuteur qu'un objet X1 est dans un endroit X2, un locuteur peut énoncer selon les cas :

- (3) Le tournevis est dans le schlumbic.
- (4) Le tournevis est là. **en montrant du doigt X2**
- (5) Le tournevis est dans le schlumbic **en montrant du doigt X2**

La phrase (3) est produite lorsque qu'il n'y a qu'un seul schlumbic et que le locuteur pense que l'interlocuteur sait ce qu'est un schlumbic et connaît l'emplacement du schlumbic. Dans le cas (4), le locuteur considère que l'interlocuteur a simplement besoin de connaître l'emplacement du tournevis sans juger utile de nommer cet emplacement. Par contre, dans (5), le locuteur suppose que l'interlocuteur ne sait pas que X2 est un schlumbic mais que cette information pourra lui être utile par la suite; cet acte de communication a donc deux fonctions : indiquer l'emplacement de X1 et fournir l'information que X2 est un schlumbic. Plus généralement, la manière de désigner un objet est choisie de façon à ce que l'interlocuteur puisse identifier l'objet sans ambiguïté et à ce qu'il en connaisse toute information jugée utile.

La position d'Appelt selon laquelle les actes de parole doivent être mis sur le même plan que les actes physiques s'oppose à juste titre à la position traditionnelle selon laquelle les actes de parole sont la traduction des idées que le locuteur veut exprimer. Néanmoins, lorsque la production de phrases ne se situe pas dans le cadre d'un dialogue oral entre deux personnes en présence, on peut adopter la position traditionnelle, ce qui revient à découper la production de phrases selon les questions **Quoi dire?** et **Comment le dire?** Il n'en reste pas moins que la perspective d'Appelt sur la production de phrases (i.e. moyen d'action pour accomplir un ou plusieurs buts) dépasse le cadre d'un dialogue oral entre deux personnes en présence. Ainsi, dans les interfaces en langue naturelle, les réponses de la machine doivent satisfaire le but d'informer l'utilisateur de la façon la plus «coopérative» possible (Grice 1975). Dans cette optique, il existe un courant de recherche qui tend à formaliser les caractéristiques d'une conversation coopérative. Par exemple, Joshi et alii (1984) avancent le principe suivant : «si vous, locuteur, avez l'intention de dire quelque chose qui peut impliquer

pour l'interlocuteur quelque chose que vous pensez faux, alors fournissez des informations supplémentaires pour empêcher la fausse inférence». Ce principe est destiné à rendre compte, par exemple, que la réponse du dialogue

— Est-ce que Luc est normalien?

— Oui, mais il n'est pas agrégé.

est préférable (plus «coopérative») qu'un simple *Oui*, dans la mesure où il est supposé que les normaliens sont habituellement agrégés. Dans le même ordre d'idée, on souhaiterait obtenir un dialogue homme/machine tel que

— Est-ce qu'il y a un train pour Pau?

— Non, mais il y a un bus.

qui suppose que le système reconnaisse le but de l'utilisateur (aller à Pau), qu'il sache que le moyen envisagé par l'utilisateur (prendre le train) est irréalisable, et qu'en conséquence il propose un autre moyen (prendre le bus).

Comme l'indique Mann (1982), il n'existe guère de systèmes qui traitent sérieusement des différentes questions de la génération automatique. Certains systèmes sont orientés sur la détermination des informations à transmettre à l'interlocuteur, les problèmes syntaxiques et les choix lexicaux étant traités de façon rudimentaire. C'est le cas par exemple des travaux de Swartout (1981) et d'Appelt (1982). C'est aussi le cas du système de Lehnert (1981) mais celui-ci est actuellement prolongé pour constituer un système complet : Cook et alii (1984) présentent une ébauche d'interface entre le système de Lehnert (1981) qui fournit une représentation des éléments devant composer un résumé et le «composant syntaxique» de McDonald (1983). Signalons aussi que des systèmes complets de dialogues homme/machine (cf. figure 1) ont été développés à Carnegie-Mellon et à Berkeley. Le système de Carnegie-Mellon (Carbonell et alii 1983), qui porte sur la sélection des composants d'un environnement informatique, comprend deux générateurs, l'un pour exprimer en anglais les résultats du module de raisonnement (Mauldin 1984), l'autre pour répondre à des questions sur le fonctionnement de ce module (Kurich 1984). Dans le système de Berkeley (Wilensky et alii 1984), logiciel d'aide à l'utilisation d'UNIX, les modules d'analyse et de génération reposent sur le même ensemble de connaissances linguistiques (Jacobs 1983). Les travaux de Joshi et alii (1984) ainsi que ceux de McCoy (1984) et de Hirschberg (1984), sont uniquement consacrés à la formalisation des caractéristiques d'une conversation coopérative. Dans

une autre direction, notre générateur ainsi que ceux que nous allons présenter ne traitent que de la question (**Comment le dire?**) : ils prennent comme entrée une représentation sémantique dont ils expriment certains ou tous les éléments explicitement ou implicitement.

2. Premiers travaux

Les premiers travaux traduisant en langage naturel un contenu informatif donné ont été réalisés par Simmons et Slocum (1972). Simmons a construit un système d'analyse automatique représentant les phrases au moyen de réseaux sémantiques. Il a ensuite réalisé un générateur à partir de ces réseaux sémantiques. Pour l'analyse comme pour la génération, il utilisait un «réseau de transition enrichi» («Augmented Transition Network» ou ATN (Woods 1970))¹. Ses réseaux sémantiques sont constitués de nœuds connectés par des relations sémantiques. À chaque nœud est associé une valeur (TOK) correspondant à un emploi de **mot** d'une langue donnée : la représentation de la phrase française

(6) Max a acheté ce lit à Marie.

serait du type (I) :

C1 = TOK :	acheter	C2 = TOK :	Max
AGENT =	C2	NBRE :	SING
DATIF =	C3	DÉT :	DÉFINI
OBJET =	C4		
C3 = TOK :	Marie	C4 = TOK :	lit
NBRE :	SING	NBRE :	SING
DÉT :	DÉFINI	DÉT :	DÉFINI

Cette représentation est donc calquée sur une langue; elle présente l'avantage de faciliter la production de phrases dans cette langue dans la mesure où le problème des choix lexicaux ne se pose pas. À l'opposé, Schank (1975) a proposé des représentations conceptuelles qualifiées d'indépendantes du langage. Ces représentations s'appuient sur un petit nombre de primitives sémantiques : par exemple, toute action dénotant l'ingestion de quelque chose par quelqu'un est représentée au moyen de la primitive «INGEST» : les phrases

1. Les «réseaux de transitions enrichis» ainsi que d'autres notions techniques utilisées en Intelligence Artificielle sont présentés en détail dans Jayez (1980). L'utilisation des réseaux de transitions enrichis pour la génération à partir de réseaux sémantiques a été développée dans Shapiro (1982).

Luc a mangé une banane.

Luc a bu du lait.

Luc a respiré de l'air.

Luc a fumé une cigarette.

Luc a pris une aspirine.

ont pour représentation :

INGEST

AUTEUR : Luc

OBJET : (une banane + du lait + de l'air + une cigarette
+ une aspirine)

De ce fait, une partie essentielle du générateur de Goldman (1975) programmé à partir de ces représentations conceptuelles est consacrée aux choix lexicaux : il lui faut déterminer si une réalisation de la primitive INGEST doit être exprimée par *manger*, *boire*, *inhaler*, *respirer*, *fumer* ou *prendre*. Pour cela Goldman a construit le questionnaire présenté dans la figure 2. Ce questionnaire est un arbre dont les nœuds sont des tests binaires et dont les branches sont les réponses à ces tests, les branches de droite correspondant à des réponses positives (O), celles de gauche à des réponses

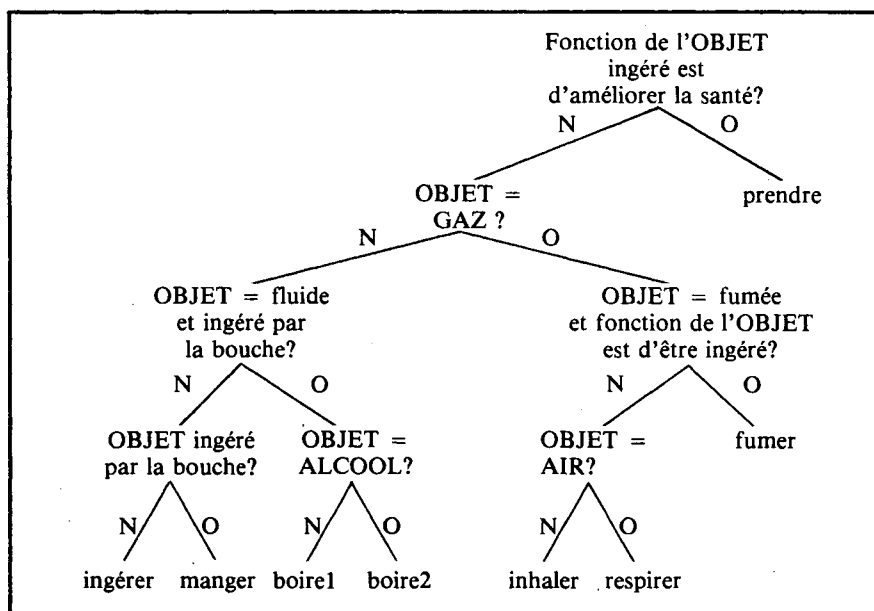


Figure 2. Questionnaire pour la primitive INGEST (Goldman 1975)

négatives (N). Les feuilles de cet arbre correspondent à des emplois de verbes.

Les représentations utilisées par Simmons et Schank illustrent deux attitudes diamétralement opposées par rapport au problème de la décomposition sémantique, problème dont nous allons évoquer les principaux aspects.

3. Représentations sémantiques

Les représentations de Schank (1975) ont été conçues pour faire appel aux concepts traités par la pensée : il est supposé qu'il existe un niveau de représentation du sens indépendant de la langue («language free»), et par là même commun à toute langue. L'hypothèse d'un niveau de pensée indépendant de la langue est l'objet de nombreuses controverses. Rappelons, par exemple, les travaux de Whorf (1956) qui argue que la perception du monde n'est pas la même pour des locuteurs de langues radicalement différentes. Il n'en reste pas moins que l'hypothèse opposée, à savoir que la meilleure représentation de la langue est la langue elle-même, semble exagérée aussi car elle soulève de nombreuses difficultés que nous allons rappeler brièvement.

3.1 *Difficultés avec les représentations calquées sur une langue*

3.1.1 Paraphrases et inférences

L'absence d'analyse sémantique dans les représentations calquées sur une langue masque les éléments de sens communs à plusieurs concepts. De ce fait, ce type de représentation ne permet pas la production de paraphrases sans le recours à des règles de mise en correspondance, e.g. règle reliant

N_0 acheter N_1 à N_2 ² et N_2 vendre N_1 à N_0

Cette contrainte alourdit les systèmes d'interrogation : la réponse à la question *Qui a vendu ce lit à Max?* à partir de la représentation (I) ne peut se faire directement. De plus, l'absence d'analyse sémantique masque les phénomènes d'inférences. Ainsi, la représentation (I), qui n'indique en aucune manière le transfert de propriété du lit de Marie à Max, ne permet pas de répondre à la question *Pourquoi Marie n'a-t-elle plus ce lit?*

2. La forme N_0 acheter N_1 à N_2 est la phrase simple indiquant la complémentation de *acheter*, N_0 désigne le sujet, N_1 le premier complément, N_2 le second.

3.1.2 Passage à une autre langue

La transcription d'une représentation calquée sur une langue en un texte d'une autre langue demande des règles de transfert, e.g. règle reliant

N_0 acheter N_1 à N_2 et N_0 to buy N_1 from N_2

Ces règles de transfert sont faciles à établir et satisfaisantes s'il existe une traduction littérale entre les deux langues; sinon on se heurte à des problèmes largement débattus dans la littérature sur la traduction.

3.1.3 Variations du contenu informatif des textes produits

Les représentations calquées sur une langue ne permettent pas facilement d'engendrer des textes de contenus informatifs adaptés aux lecteurs. Pour illustrer ce point, prenons le cas d'un match de coupe de football dont une représentation calquée sur le français serait du type :

(A) BATTRE

AUTEUR = NOM : Monaco

OBJET = NOM : Valenciennes

SCORE : 3-2

Cette représentation débouche naturellement sur un texte construit autour de la phrase simple N_0 battre N_1 . Or celle-ci ou ses transformées syntaxiques mentionnent obligatoirement l'équipe battue : à l'actif, le complément N_1 , qui réfère à l'équipe battue, est obligatoire : *Monaco a battu (Valenciennes + *E)*³; les transformées syntaxiques de cette phrase simple ont l'équipe battue comme sujet, complément obligatoire :

Valenciennes a été battu (par Monaco + E)

Valenciennes s'est fait battre (par Monaco + E)

La représentation (A) ne permet donc pas la synthèse de formulations telles que

(7) Monaco a remporté la victoire par 3 à 2.

où seule l'équipe victorieuse est mentionnée. La formulation (7) peut être synthétisée à partir d'une représentation plus abstraite comme :

3. Le symbole *E* désigne la séquence vide. La notation *Monaco a battu (Valenciennes + *E)* réfère aux deux phrases : *Monaco a battu Valenciennes*, et **Monaco a battu*. Les symboles «?», «?*» et «*» sont accolés à des phrases ou à des textes qui sont respectivement peu acceptables, guère acceptables ou inacceptables.

(B) *MATCH-COUPÉ*

ÉQUIPE 1 = NOM : Monaco

BUTS : 3

ÉQUIPE 2 = NOM : Valenciennes

BUTS: 2

Cette représentation peut aussi déboucher sur

(8) Valenciennes a été (éliminé + battu) par 3 à 2.

où seule l'équipe battue est mentionnée. Les formulations (7) (resp. (8)) peuvent être souhaitées si elles sont destinées à des supporters de l'équipe victorieuse (resp. battue). En conséquence, on préférera des représentations plus abstraites comme (B) lorsque l'on veut avoir la possibilité de varier les contenus informatifs des textes selon les lecteurs.

3.2 Degré d'imprécision dans les représentations abstraites

Ce sont les difficultés que nous venons d'évoquer qui ont amené Schank, entre autres, à abandonner les représentations calquées sur une langue au profit de représentations plus abstraites. Mais si les décompositions élaborées de Schank (1975) mettent l'accent sur les éléments de sens partagés par plusieurs concepts, elles masquent les propriétés qui les distinguent. De ce fait, elles débouchent sur des textes d'un vocabulaire pauvre. Ainsi il serait difficile de raffiner le questionnaire de la figure 2 destiné à trouver un verbe exprimant la primitive INGEST, afin d'y inclure des variantes de *manger* telles que *paître, croquer, dévorer, avaler, brouter, ronger, baffrer, mâchonner, grignoter, picorer, gober, engloutir*, etc. Avec Saint-Dizier (1984), nous considérons que le problème peut se poser dans les termes suivants : «sachant que les langues naturelles sont d'une finesse d'expression élevée, quel est le degré d'imprécision que l'on peut tolérer dans les représentations sémantiques». Bien entendu, nous n'avons pas la prétention de répondre à cette question. Toutefois, il nous semble important de faire une distinction entre la représentation des «éléments prédicatifs» et celle des «éléments non prédicatifs». La notion d'élément prédicatif (Harris 1968) couvre entre autres les verbes (*respecter*), les adjectifs (*respectueux*) et certains noms (*respect*), en gros les noms qui n'ont pas qu'un emploi concret. Cette notion rejoint celle de phrase simple car on peut associer une phrase simple à un élément prédicatif

N_0 respect N_1 (les lois)

N_0 être respectueux de N_1

N_0 avoir un certain respect de N_1

la phrase simple étant construite autour de *être* pour les adjectifs et d'un «verbe support» (e.g. *avoir*) pour les noms prédicatifs. La notion de verbe support (Gross 1981) réfère à des verbes qui ne sont que les supports des marques de temps et de personne, la fonction prédicative étant portée par un substantif. On notera que le mot *table* fait partie d'un prédicat dans *Luc est à table*; par contre, le mot *chaise* n'est pas un nom prédicatif : il désigne toujours un objet concret.

Pour éviter les difficultés évoquées en 3.1, il peut être nécessaire de décomposer certains éléments prédicatifs. Ainsi la décomposition du prédicat *tuer* de la phrase simple N_0 *tuer* N_1 suivant la relation causale

(N_0 faire quelque chose) = = = > (N_1 être mort)

permet, entre autres, de répondre directement à la question *Qui est mort?* Par contre, il semble sans grand intérêt de représenter un objet concret (e.g. une chaise) par une décomposition reflétant une définition de dictionnaire (e.g. *meuble destiné à s'asseoir et comportant un dossier*). En effet, considérons par exemple le problème de la traduction. S'il est souvent impossible de traduire littéralement une phrase simple ou une construction syntaxique (e.g. les phrases simples anglaises N_0 *to shoot and kill* N_1 ou N_0 *to shoot* N_1 *dead* n'ont pas d'équivalent français, et il en est de même pour la construction passive anglaise *Mary was told that story by Max*), il est généralement possible de traduire littéralement un nom (simple ou composé) désignant un objet concret. Certes, il existe des exceptions à ce principe. Ainsi, le nom composé américain *doggy bag* n'a pas de traduction littérale française. Mais une représentation de *doggy bag* reflétant une définition de dictionnaire débouche sur des textes français comme :

- (9) Hier, au restaurant, John a demandé un sac en plastique pour emporter la nourriture qu'il n'avait pas mangée.

Or, de tels textes ne sont pas la règle : si un roman américain comporte la phrase

- (10) Yesterday, in a restaurant, John asked for a doggy bag.

sa traduction française ne sera généralement pas (9) mais le texte suivant :

Hier, au restaurant, John a demandé un *doggy bag**

N.D.T. : Un *doggy bag* est un sac en plastique qu'un client demande dans un restaurant pour emporter la nourriture qu'il n'a pas mangée. Il est courant aux États-Unis de demander un *doggy bag*.

où le mot *doggy bag* n'est pas traduit dans le texte mais expliqué dans une note du traducteur (N.D.T.). En plus de la définition de *doggy bag*, cette note comporte une information essentielle, à savoir qu'il est courant aux États-Unis de demander un *doggy bag*. Cette information permet d'éviter de mauvaises inférences. En effet, un lecteur français lisant (9) en déduira que *John* est grossier et radin, ce qui n'est pas dans les intentions de l'auteur américain écrivant (10). En conclusion, la représentation d'un *doggy bag* devrait être associée à ce mot et non à sa définition⁴. Une telle représentation pourrait être complétée par une banque de données intégrant définitions de dictionnaires et variations culturelles. Cette position sur la représentation des objets concrets n'est pas sans rapport avec celle soutenue par Appelt (1982) pour les objets dont le nom est inconnu de l'interlocuteur. Appelt argue qu'il est alors préférable de montrer du doigt de tels objets plutôt que d'en donner une description non évocatrice ou qui risque d'induire l'interlocuteur en erreur.

Les rapports entre la représentation des connaissances et la langue étant loin d'être simples, on observe dans la pratique des spécialistes une grande variété de représentations, le choix d'une représentation se faisant en fonction des objectifs visés.

3.3 *Choix d'une représentation en fonction des objectifs visés*

Le choix d'une représentation dépend en premier lieu de la tâche à accomplir. Prenons le cas bien connu de la traduction automatique. Selon les projets, il y a une ou deux représentations intermédiaires, comme indiqué dans la figure 3. Le nombre de représentations intermédiaires et leur nature sont choisis en fonction de la finalité du projet : s'agit-il de traduction d'une langue vers une autre, ou d'une langue vers plusieurs langues? Quelle est la typologie des langues concernées? Quelle est la qualité de traduction demandée? En effet, un système ne comportant qu'une

4. Une représentation de *doggy bag* associée à ce mot présente aussi l'avantage d'aboutir facilement au texte anglais (10) et non à la traduction littérale de (9) qui est peu naturelle :
 ? *Yesterday, in a restaurant, John asked for a plastic bag to carry away the food he did not eat.*

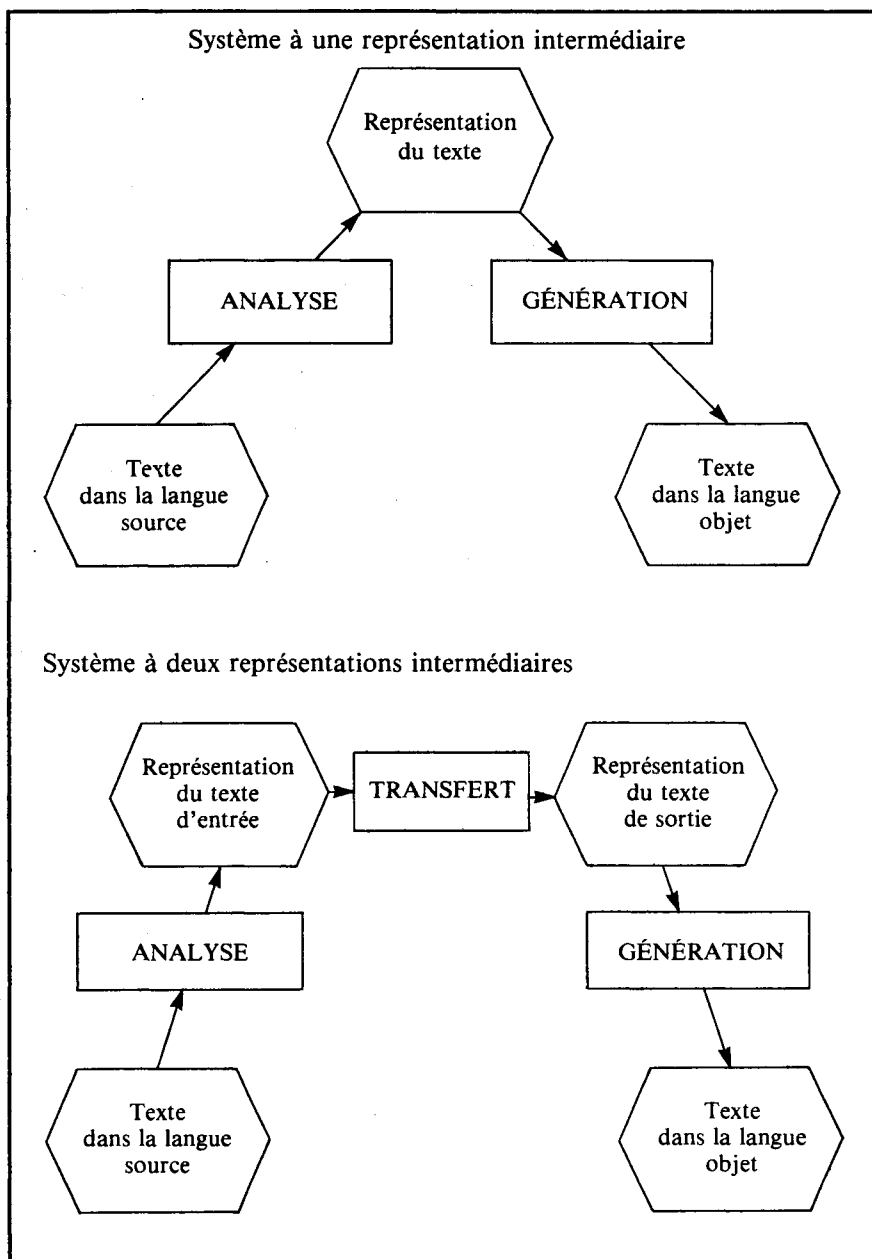


Figure 3. Systèmes de traduction automatique

représentation intermédiaire est plus économique quand il s'agit de traduction en plusieurs langues ; avec deux représentations intermédiaires, il faut adopter pour chaque langue objet deux modules, celui de transfert et celui de génération. Par contre, un système à deux représentations intermédiaires permet de garder des informations sur la structure et le vocabulaire du texte source, ce qui est souhaitable lorsqu'il y a un fort degré de parenté entre les langues source et objet. Étant donné ces diverses finalités, il existe toute une gamme de systèmes de traduction. Certains systèmes ont deux représentations intermédiaires comportant des informations syntaxiques telles que les structures en arbres syntagmatiques des phrases sources et des phrases objets (Boitet et Nedobejkine 1981); les textes produits sont alors des traductions relativement fidèles des textes de départ. D'autres systèmes n'ont qu'une représentation intermédiaire avec aucune information sur la structure et le vocabulaire du texte source (Lytinen et Schank 1982); les textes produits sont alors des «reformulations» des textes de départ.

La variété des représentations sémantiques provient aussi de la diversité des domaines concernés (textes scientifiques, histoires psychologiques ou factuelles, etc.). Ainsi, une représentation utilisée pour interroger une banque de données relationnelle doit mettre l'accent sur les problèmes de quantification afin d'être en mesure de répondre à des questions comme la suivante, extraite de Pique et Sabatier (1982) :

Combien y-a-t-il de colonels qui dirigent un département et qui ne dépendant pas du Général Pierre?

Ces questions, qui comportent des empilements de déterminants, demandent un formalisme à caractère logique, tel que celui développé dans Colmerauer (1979) et Colmerauer et Pique (1980). Mais un formalisme à caractère logique ne s'impose pas pour des drames passionnels par exemple, pour lesquels les problèmes de quantification cèdent le pas devant les connaissances pragmatiques relatives aux causes et manifestations de l'amour et de la jalousie, connaissances qui demandent des représentations analogues à celles développées dans Lehnert (1981). Mais comme Lehnert (1983) le souligne, son approche ne pourrait pas être appliquée à d'autres types de textes que ceux à caractère psychologique.

Pour des textes scientifiques, Harris et alii (1985) ont développé un mode de représentation basé sur les connaissances propres au domaine. Ainsi, l'immunologie met en jeu certaines classes d'éléments : A = les

anticorps, G = les antigènes, C = les cellules. Les relations entre ces classes sont fortement contraintes. Par exemple, toute relation entre anticorps et cellules concerne l'existence d'anticorps dans les cellules, quelles que soient les modalités de cette relation notée T. En conséquence, les phrases

Il y a des anticorps dans ces cellules.

On trouve beaucoup d'anticorps dans ces cellules.

Ces cellules contiennent peu d'anticorps.

Ces cellules produisent rapidement des anticorps.

Ces cellules ne produisent plus d'anticorps.

peuvent toutes recevoir le même type de représentation (i.e. A T C), en annotant cette représentation d'indications reflétant les variantes modales ou aspectuelles. Un élément de type A T C constitue une unité d'information. Les contraintes sur les relations entre classes permettent de structurer les unités d'information en une «grammaire». Ainsi, le fait qu'il ne peut exister d'antigènes dans les cellules amène à exclure la relation G T C de la grammaire⁵. Celle-ci régit une représentation précise du contenu informatif d'un article d'immunologie.

Les représentations de Lehnert et de Harris s'appuient sur les connaissances spécifiques relatives à un domaine. Il en est de même pour les représentations basées sur la notion de «schéma»⁶ (cf. entre autres Minsky 1975, Schank et Abelson 1977, Charniak 1977, Winograd 1977). Les schémas reposent sur la nécessité de raisonner par analogie en imposant des stéréotypes à la réalité. Ce mode de raisonnement, qui permet d'arriver à des conclusions sur la base de similitudes partielles, s'oppose à la déduction logique où tout est vrai ou faux (ou incertain). Il a été étayé par de nombreux travaux sur la mémoire, anciens (Bartlett 1932) ou récents (Schank 1982). Un schéma structure les données en deux niveaux : d'une part, celles qui sont toujours réalisées dans la situation en jeu, e.g. le schéma KIDNAPPING indique qu'un kidnapping commence par un enlèvement suivi d'une demande de rançon, celle-ci étant éventuellement versée, etc.; d'autre part, celles qui sont spécifiques à la situation, e.g. le schéma KIDNAPPING intègre un «formulaire» avec une rubrique sur la

5. Il est admis que des phrases négatives comme

Il n'existe pas d'antigènes dans les cellules.

n'apparaissent jamais dans les articles scientifiques car elles constituent le «b-a-ba» de l'immunologie. Ceci permet d'exclure la relation G T C de la grammaire.

6. Nous utilisons le mot «schéma» comme traduction du mot anglais «frame».

personne enlevée, une rubrique sur l'enlèvement, une sur la demande de rançon, etc. Chaque rubrique comporte une série de questions : nom, prénom, âge, nationalité, fonction pour la rubrique concernant la personne kidnappée; lieu, date, heure, mode de transport, destination pour celle concernant l'enlèvement.

De nombreux systèmes d'analyse automatique reposent sur la notion de schéma. Il s'agit alors d'identifier le schéma correspondant au stéréotype de la situation décrite dans le texte et de remplir le formulaire du schéma suivant les données du texte. Ainsi, lors de l'analyse de la phrase

Un industriel, Pierre Fouin, a été kidnappé hier matin à Paris.

Lytinen et Schank (1982) font appel au schéma KIDNAPPING (dans le jargon des spécialistes, le schéma KIDNAPPING est «instancié»); ce schéma guide l'analyse de la suite du récit : si celui-ci se poursuit par

Un livreur de lait a vu deux hommes l'embarquer dans une voiture rouge. Sa femme a reçu un coup de téléphone lui réclamant 100.000F.

il sera reconnu que l'adjectif possessif *sa* du groupe nominal *sa femme* réfère à Pierre Fouin, la personne kidnappée, et non au livreur de lait, et que le coup de téléphone concerne la demande de rançon, et non une demande de remboursement ou d'emprunt.

Dans l'exemple précédent, le verbe *kidnapper* est directement associé au schéma KIDNAPPING. D'une manière plus générale, les représentations s'appuyant sur la notion de schéma devraient être plus proches du langage que les représentations conceptuelles de Schank (1975). Ainsi, il semblerait que les verbes *manger*, *gober* et *dévorer* doivent être associés au même schéma MANGER, le verbe *fumer* étant associé à un autre schéma. En effet, l'élément de sens commun à *manger* et *fumer* (i.e. ces deux verbes désignent une ingestion) relèvent d'un calcul, ce qui n'est pas le cas des éléments de sens communs à *manger*, *gober* et *dévorer* : ces verbes se construisent avec le même type d'objet, ils sont formulés dans des circonstances analogues et ils éveillent le même ensemble de connaissances (i.e. on (mange + gobe + dévore) un comestible qui est ensuite digéré dans l'estomac puis...). À l'appui de ces observations, signalons des phénomènes linguistiques telles que des relations de fusion (Gross 1975)

Luc a gobé son œuf = Luc a mangé son œuf en le gobant

Luc a dévoré sa choucroute = Luc a mangé sa choucroute en la dévorant

qui indiquent que *dévoré* et *gober* sont sémantiquement proches de *manger*. Il n'existe aucune relation de fusion entre *manger* et *fumer* :

*Luc a mangé sa soupe en la fumant.

*Luc a fumé sa cigarette en la mangeant.

3.4 Entrées des générateurs

Ce bref tour d'horizon des différentes représentations sémantiques indique que les systèmes de génération automatique semblent condamnés à avoir des entrées de natures variées. Apportons cependant la précision suivante. Lorsque la production de textes est réalisée à partir de représentations comportant des informations sur la structure et le vocabulaire du texte à produire, comme c'est le cas dans le système de traduction automatique de Boitet et Nedobejkine (1981), il n'est pas considéré que cette tâche est du domaine de la génération automatique : l'objet de la génération automatique est la production de textes à partir de représentations sémantiques «abstraites», c'est-à-dire des représentations qui peuvent rendre compte des phénomènes de paraphrases et d'inférences et qui permettent la production de textes dans plusieurs langues, ces textes étant éventuellement adaptés aux lecteurs.

En partant de «représentations abstraites», et quelles que soient les particularités propres de ces entrées, les systèmes de génération automatique sont tous confrontés aux problèmes suivants :

— **problèmes conceptuels** : quelles sont les informations qui doivent apparaître dans le texte? Dans quel ordre doivent-elles apparaître? Quelles sont celles qui doivent être exprimées explicitement par rapport à celles qui peuvent être laissées implicites?

— **problèmes linguistiques** : comment découper le texte en phrases? Quelles constructions syntaxiques choisir? Quels mots choisir?

— **problèmes syntaxiques** : former des phrases qui obéissent aux règles de grammaire de la langue objet (bonne formation de la phrase et de ses constituants, règles d'accord, placement de la négation, forme des pronoms, etc.).

Nous allons maintenant examiner comment différents chercheurs traitent ces questions.

4. Travaux récents

4.1 Générateur de McKeown

4.1.1 Présentation

McKeown (1982) a construit un système de génération qui engendre des réponses à des questions sur une banque de données. Son domaine d'application est une classification d'armes militaires. Nous présenterons son générateur en l'illustrant par une classification de mammifères, domaine plus ordinairement appréhendable que celui des armes militaires. Sa banque de données est organisée en un réseau sémantique analogue à celui présenté dans la figure 4. Le travail de McKeown s'appuie sur le fait que les réponses à certaines questions (e.g. *Qu'est-ce-que X?*) suivent un nombre limité de schémas stéréotypés. Ce fait est particulièrement apparent dans les définitions de dictionnaire : celles-ci sont dans un style abrégé, mais elles suivent des normes de présentation, même si ces normes ne sont pas respectées systématiquement. Ainsi, dans le dictionnaire encyclopédique Hachette, on trouve les définitions suivantes :

- **mysticète** : sous-ordre des cétacés comprenant les espèces pourvues de fanons (baleines).
- **odontocète** : sous-ordre des cétacés pourvus de dents (dauphins, cachalots, narvals).
- **baleine** : mammifère marin mysticète (muni de fanons) comptant parmi les plus gros animaux (14m à 24m).
- **narval** : mammifère cétacé odontocète long de 4 à 5m, vivant en bandes dans l'Arctique.

Dans ces définitions, on peut identifier des «prédicats rhétoriques» : les entrées de mysticète et odontocète commencent par une «identification», prédicat rhétorique composé d'une référence à la classe supérieure (les cétacés) suivie de la propriété distinctive (comprenant les espèces pourvues de fanons ou pourvues de dents). Ces entrées se terminent par des «exemples», c'est-à-dire une liste de certains membres de la classe. McKeown a ainsi identifié un certain nombre de prédicats rhétoriques (e.g. identification, exemples, analogie, attributs, illustration, amplification). Elle les a ensuite organisés dans des «moules», schémas stéréotypés pour répondre à un certain type de questions. Un moule contient une liste ordonnée de combinaisons de prédicats rhétoriques, une combinaison de

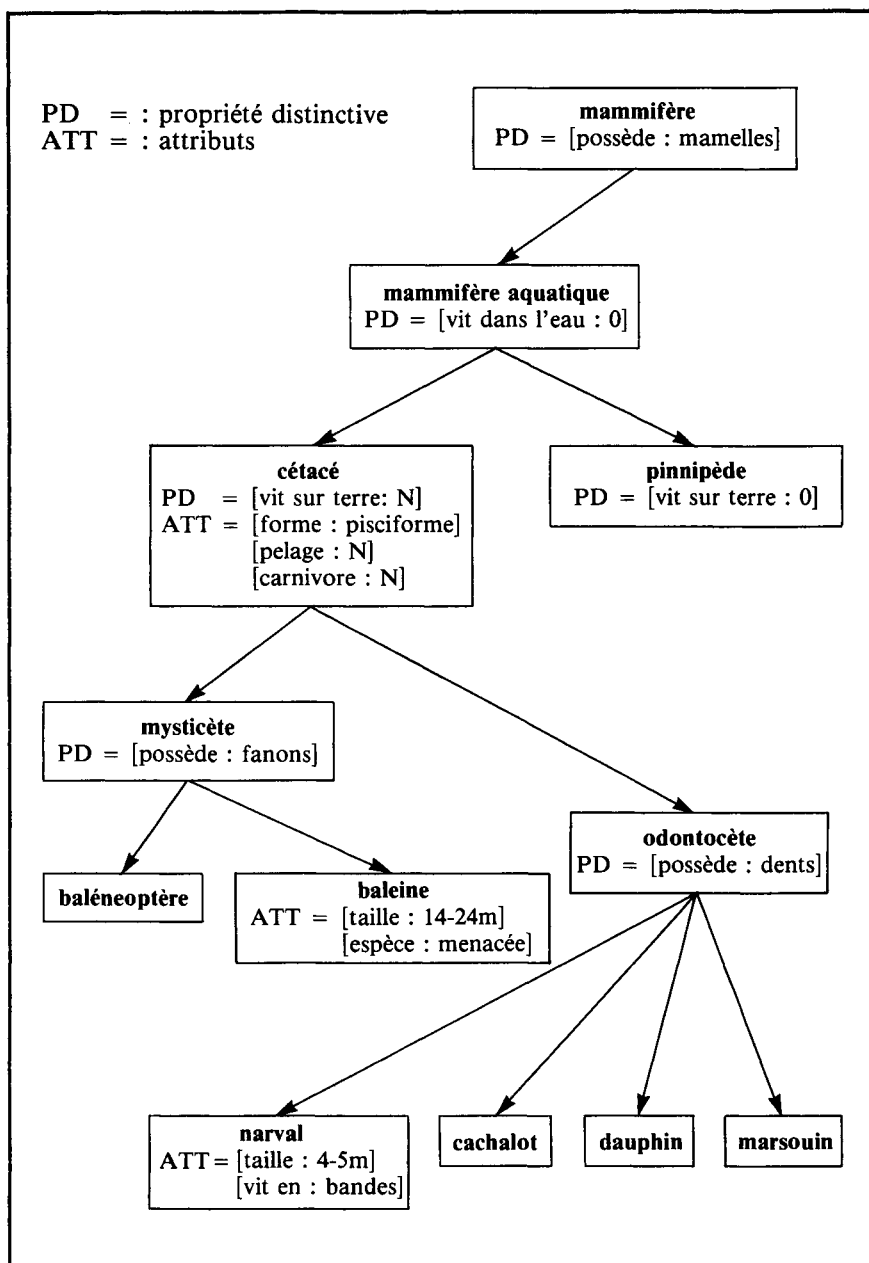


Figure 4. Échantillon d'une classification des animaux

prédicats pouvant être un prédicat obligatoire, un prédicat facultatif, une alternative entre prédicats, ou une répétition du même prédicat. Par exemple, le moule «définition» est composé de la liste suivante⁷ :

Identification

(Analogie / Exemples / Attributs / Autres-noms)*

Illustration / Évidence

(Amplification / Analogie / Attributs)

(Illustration / Évidence)

Ce moule ne constitue qu'un guide pour construire une définition : il n'indique que partiellement la façon d'organiser les informations dans la mesure où il comporte des options. Ces options permettent d'adapter le moule selon les données présentes dans la classification. Différents moules peuvent être associés à un même type de question. Ainsi, la question *Qu'est-ce-que X?* peut recevoir une réponse qui se conforme au moule «définition» ou au moule «constituants». Le choix entre ces deux moules est fait en fonction des données de la classification : si un sous-ordre des mammifères n'a que peu de propriétés caractéristiques mais qu'il contient une riche description de ses constituants, le moule «constituants» est choisi, sinon c'est le moule «définition» qui est choisi.

Un moule ayant été sélectionné, l'étape suivante consiste à le remplir en fonction de la banque de données. Pour cela, des correspondances sont établies entre les prédicats rhétoriques et les relations sémantiques de la classification. Ainsi, au prédicat rhétorique «attribut» correspond le prédicat sémantique

(attribut <entité> <nom-attribut> <valeur-attribut>)

dont une réalisation est :

(P) (attribut <CÉTACÉ> <FORME-CORPS> <PISCIFORME>)

Cette étape aboutit à un ensemble de propositions analogues à (P) qui traduisent les éléments de la classification en correspondance avec les prédicats rhétoriques du moule retenu. Les propositions associées aux options du moule sont facultatives ou constituent des alternatives à d'autres

7. Les signes (et) délimitent une séquence facultative, le signe indique / une alternative et le signe * étoile indique que la séquence peut être répétée.

propositions. Pour sélectionner celles qui seront développées en phrases, McKeown fait appel à la notion de «focus». Le focus d'une phrase correspond, grosso modo, à l'individu ou objet dont il est question dans la phrase. Ainsi, dans la forme *Qu'est-ce-qu'une baleine?* le focus est *la baleine*. Si la réponse est *La baleine est un mysticète* le focus est toujours *la baleine*, mais la notion de *mysticète* ayant été introduite, elle devient un «focus potentiel». Un focus potentiel est un élément qui peut devenir sujet de discussion dans les phrases suivantes. Par exemple, on peut continuer la réponse en définissant un *mysticète* :

La baleine est un mysticète. Le mysticète est un cétacé muni de fanons.

Dans ce texte, le focus de la deuxième phrase est devenu *le mysticète*, focus potentiel de la première phrase; le focus de la première phrase (*la baleine*) est devenu «focus de réserve», c'est-à-dire élément dont il a déjà été question mais sur lequel on peut revenir. McKeown a établi un certain nombre de règles permettant de déterminer si le focus d'une phrase doit être le focus de la phrase précédente, un élément de la liste des focus potentiels ou un élément de la liste des focus de réserve. Par exemple, une de ces règles est la suivante :

Règle 1 : s'il y a le choix entre garder le focus de la phrase précédente ou passer à un élément de la liste des focus potentiels, alors opter pour la seconde solution.

Pour illustrer l'application de cette règle, supposons qu'après avoir formulé

(11) La baleine est un mysticète. [focus = la baleine]
[focus potentiel = le mysticète]

l'on veuille apporter les deux informations suivantes :

(12) La baleine est un des plus gros animaux. [focus = la baleine]
(13) Le mysticète est un cétacé muni
de fanons. [focus = le mysticète]

Il faut alors déterminer si le texte généré sera composé de la phrase (11) suivie de (12) suivie de (13), ou de la phrase (11) suivie de (13) suivie de (12). Avec l'ordre (11)-(12)-(13), le focus de la seconde phrase est le même que celui de la première phrase (i.e. la baleine); par contre avec l'ordre (11)-(13)-(12), le focus de la seconde phrase est le mysticète, focus potentiel de la première phrase. On se trouve donc dans le champ d'application de la

Règle 1. Celle-ci indique qu'il faut opter pour le passage à un focus potentiel, et donc pour l'ordre (11)-(13)-(12); elle permet donc de déboucher sur

(T1) La baleine est un mysticète. Un mysticète est un cétacé muni de fanons. La baleine est un des plus gros animaux. [ordre (11)-(13)-(12)]

et non sur

(T2) La baleine est un mysticète. La baleine est un des plus gros animaux. Un mysticète est un cétacé muni de fanons. [ordre (11)-(12)-(13)]

résultat satisfaisant puisque (T1) est indéniablement préférable à (T2).

Ces règles sur le focus sélectionnent les propositions analogues à (P) qui apparaîtront dans la réponse et déterminent leur ordre d'apparition. Cette liste ordonnée de propositions constitue un «message». Chaque proposition du «message» est transmise à un «dictionnaire» qui a pour rôle de sélectionner un verbe, une «structure profonde» et les mots désignant les arguments du verbe. Ainsi, à partir de la proposition (P) que nous rappelons

(P) (attribut <CÉTACÉ><FORME-CORPS>
<PISCIFORME> <FOCUS = CÉTACÉ>)

la consultation du dictionnaire fournit la «structure profonde» suivante :

[(verbe ((v = avoir)))
(AGENT ((n = cétacé)))
(OBJET ((adj = pisciforme)
(article = indéfini)
(n = corps)))
(FOCUS = AGENT)]

Celle-ci n'est pas basée sur des fonctions syntaxiques telle que sujet ou objet direct, mais sur des fonctions thématiques comme AGENT ou OBJET⁸, ce qui permet l'utilisation d'une «grammaire d'unification». Une description détaillée des grammaires d'unification se trouve dans Kay (1979). Indiquons simplement qu'une grammaire d'unification se distingue d'une grammaire

8. Pour distinguer les fonctions syntaxiques (e.g. sujet, objet direct) des fonctions thématiques (e.g. AGENT, OBJET), ces dernières sont écrites en majuscules.

générative (Chomsky 1957) entre autres par le fait que le même statut est attribué aux fonctions thématiques et aux fonctions ou catégories syntaxiques. De plus, une grammaire d'unification permet d'indiquer le focus d'une phrase. Cette information est utilisée pour certains choix syntaxiques (McKeown 1983). Ainsi, le choix entre la construction active ou passive d'un verbe repose sur la règle suivante :

Règle 2 : si le focus est sur l'AGENT, construire la phrase à l'actif; si le focus est sur l'OBJET, construire la phrase au passif.

Cette règle s'appuie sur le fait qu'il est généralement admis que le focus d'une phrase doit être en position sujet :

(14) Max a abandonné la maison. [focus sur Max]

(15) La maison a été abandonnée
par Max. [focus sur la maison]

Elle suppose de plus une correspondance biunivoque entre fonctions thématiques et fonctions syntaxiques, à savoir : à la voix active, le sujet de la phrase désigne l'AGENT. La **Règle 2** permet donc de dériver (15) de la «structure profonde» suivante :

[(verbe ((v = abandonner)))
(AGENT ((n = Max)))
(OBJET ((article = défini)
(n = maison)))
(FOCUS = OBJET)]

Chaque proposition du message est traduite en une phrase par l'intermédiaire de la grammaire d'unification et le texte final est constitué de la juxtaposition de ces phrases dans l'ordre indiqué par le message. Le fait que la juxtaposition soit le seul procédé de linéarisation employé débouche sur des textes d'une syntaxe simple, en particulier sans relatives ou subordonnées. Mais McKeown et Derr (1984) se sont penchées sur le problème de la linéarisation en phrases. Là aussi, elles font appel à la notion de focus. Ainsi, pour déterminer si les phrases

(16) Max a donné à Marie un livre.

(17) Marie avait besoin du livre.

doivent être juxtaposées

Max a donné à Marie un livre. Marie avait besoin du livre.

ou enchaînées par une relative

Max a donné à Marie un livre dont elle avait besoin.

elles examinent le focus de la phrase qui va suivre. Si celui-ci est *Max*, focus de (16), alors la relativation est choisie

Max a donné à Marie un livre dont elle avait besoin. Il l'a acheté à la librairie de l'Université.

sinon c'est la juxtaposition qui est choisie :

Max a donné à Marie un livre. Mais avait besoin du livre et avait l'intention de l'acheter.

4.1.2 Commentaires

L'idée de départ de McKeown, à savoir que les réponses à certaines questions suivent des schémas stéréotypés, nous paraît totalement justifiée. D'une manière plus générale, il est clair que les discours sont structurés et que l'on doit chercher à dégager leur structure (leur moule)⁹. De plus, même si quelques notions restent floues¹⁰, il doit être possible de répertorier l'ensemble des structures possibles pour un type de discours donné, comme l'a fait McKeown pour les réponses à certaines questions.

Par contre, nous n'adhérons pas aux processus utilisant la notion de focus. Commençons par examiner le rôle du focus pour l'ordre des informations, par exemple pour obtenir (T1) et non (T2), textes que nous rappelons :

(T1) La baleine est un mysticète. Un mysticète est un cétacé muni de fanons. La baleine est un des plus gros animaux.

(T2) La baleine est un mysticète. La baleine est un des plus gros animaux. Un mysticète est un cétacé muni de fanons.

La mise en œuvre de règles générales comme la **Règle 1** pour engendrer (T1) et non (T2) nous paraît précipitée. En effet, cette règle peut aboutir à des

9. Mann (1984) met aussi en avant la nécessité de développer une «théorie de structure rhétorique» et décrit l'ébauche d'une telle théorie. Celle-ci est basée sur des «schémas» comportant un noyau et des satellites. Ainsi, le schéma «requête» a pour noyau la requête elle-même et pour satellites des motivations et des informations permettant à l'interlocuteur de satisfaire la requête. Un texte est organisé autour d'un ou plusieurs schémas et les éléments d'un schéma sont éventuellement décomposés en d'autres schémas.

10. McKeown souligne qu'il n'est pas toujours simple d'identifier le statut rhétorique d'une proposition (e.g. «illustration»/«exemples»). Néanmoins, ces notions ont un niveau opératoire convenable.

résultats malencontreux lorsque, par exemple, des relations causales et/ou temporelles sont en jeu. Pour illustrer cette affirmation, supposons qu'après avoir formulé

- (11') Max était marié à Marie. [focus = Max]
[focus potentiel = Marie]

l'on veuille apporter les deux informations suivantes :

- (12') Max était infidèle à Marie. [focus = Max]
(13') Marie s'est suicidée. [focus = Marie]

Ces deux informations, qui mettent en jeu une relation causale, (12') est la cause de (13'), peuvent être enchaînées dans n'importe quel ordre :

- Max était infidèle à Marie. Elle s'est suicidée. (12') avant (13')
Marie s'est suicidée. Max lui était infidèle. (13') avant (12')

Après (11'), on a donc le choix entre garder le même focus, (11') suivi de (12'), ou passer à un focus potentiel, (11') suivi de (13'). Si on fait appel à la **Règle 1**, celle-ci indique que (11') doit être suivi de (13') et donc débouche sur le texte

- (T1') ?*Max était marié avec Marie. Elle s'est suicidée.
Il lui était infidèle. [ordre (11')-(13')-(12')]

qui est maladroit contrairement au texte suivant où (11') est suivi de (12') :

- (T2') Max était marié avec Marie. Il lui était infidèle.
Elle s'est suicidée. [ordre (11')-(12')-(13')]

Cette situation n'est pas spécifique à la **Règle 1**. En effet, les règles sur le focus sont supposées refléter des principes abstraits de portée générale. De ce fait, elles ne tiennent compte ni du contenu ni de la forme des phrases concernées. Or le langage naturel semble échapper à des principes généraux abstraits : que ce soit en analyse ou en génération, toute décision demande la prise en compte de nombreux facteurs spécifiques, tant pragmatiques que sémantiques, syntaxiques, lexicaux ou stylistiques. Ainsi, nous venons de voir que l'on ne pouvait pas appliquer la **Règle 1** sans tenir compte d'éléments de sens (i.e. relations causales et temporelles entre les événements décrits dans (11'), (12') et (13')). De même, la **Règle 2** sur le choix entre une construction active ou passive ne peut pas s'appliquer sans tenir compte des éléments lexicaux en jeu car il existe des formes verbales qui ne permettent pas le passif :

Max a quitté la maison.

*La maison a été quittée par Max.

Qui plus est, s'il est vrai que les formes verbales référant à une «action» peuvent avoir comme sujet à l'actif l'auteur de cette action, il n'existe pas de correspondance analogue pour les formes verbales référant à un «état» : leur sujet n'est pas associé à un rôle sémantique, comme le montrent les paires suivantes :

(18) a. Max aime Marie.

(19) a. Marie plaît à Max.

(18) b. Max me manque.

(19) b. I miss Max.

Si on admet que *Max* est l'AGENT dans ces différentes phrases et si on veut que le focus soit sur *Max*, la **Règle 2** indique que ces formes verbales doivent être à la voix active. Ceci donne les résultats escomptés pour les phrases (18) mais pas pour les phrases (19).

On retiendra donc qu'il est hors de question d'appliquer les règles sur le focus mécaniquement, c'est-à-dire sans tenir compte du contenu et de la forme des phrases concernées. Comme la notion du focus reflète une intuition indiscutable, on peut chercher à établir les conditions d'application des règles de focus en examinant les interactions entre focus, facteurs pragmatiques, sémantiques, syntaxiques et lexicaux. Néanmoins, l'entreprise est colossale, si tant est qu'elle soit réalisable car on a du mal à imaginer une méthode permettant de la mener à bien. Il est donc raisonnable d'éviter les règles de focus. Examinons comment traiter le problème de l'ordre des informations sans avoir recours à la notion de focus, par exemple comment aboutir à (T1) et non à (T2). Pour cela, il semble qu'il suffise de s'appuyer sur les structures de discours¹¹. Par exemple, on peut poser que si une définition commence par une référence à la classe supérieure (e.g. *La baleine est un mysticète*) et que cette classe supérieure doit être identifiée (e.g. *Un mysticète est un cétacé muni de fanons*), alors cette identification doit suivre la référence à la classe supérieure en précédant tout autre information (e.g. *La baleine est un mysticète. Un mysticète est un cétacé muni de fanons. La baleine est un des*

11. À ce propos, McKeown signale qu'elle ne tient pas compte de l'interaction entre structures de discours (moules) et contraintes sur le focus mais qu'indéniablement ces deux phénomènes interfèrent.

plus gros animaux). Autrement dit, il est envisageable de raffiner les «moules» afin d'y inclure les procédures de choix lorsqu'il y a plusieurs alternatives. D'une manière plus générale, le problème de l'ordre des informations peut être traité au moyen d'une «grammaire de discours» (cf. 4.3) et il en est de même pour les décisions sur la linéarisation en phrases et les constructions syntaxiques. À l'opposé des règles de focus, la notion de grammaire de discours reflète le fait que le traitement automatique du langage naturel nécessite un ensemble de connaissances spécifiques.

4.2 Modèles *«psychologiquement plausibles»*

Certains travaux de génération automatique tendent à simuler un locuteur parlant et cherchent à construire des générateurs «psychologiquement plausibles». Ces travaux s'appuient sur des remarques comme : lorsque l'on parle, on n'a qu'une vague idée de ce que l'on va dire, ou bien, lorsque l'on commence une phrase, on ne sait pas exactement comment on va la terminer. Une stricte observation de ces remarques amènerait à considérer que l'on produit les discours pas à pas de «gauche à droite» : pour les phrases affirmatives de langues à ordre de base sujet—verbe—complément, la procédure consisterait à déterminer l'élément qui va constituer le sujet, ensuite à rechercher un terme pour le verbe et ainsi de suite. Néanmoins cette position n'est généralement pas retenue, vu que l'absence totale d'anticipation déboucherait sur des magmas informes et incohérents, ce que ne sont pas les discours oraux même si ceux-ci sont moins soignés que les textes écrits. Les générateurs psychologiquement plausibles ne fonctionnent donc pas complètement de «gauche à droite», mais l'utilisation de procédures de gauche à droite est quand même leur principe de base.

L'architecture du système de McDonald (1983) est semblable au système de McKeown : l'entrée du générateur est examinée par un «conceptualiseur» qui prend les décisions conceptuelles et fournit un «message», liste ordonnée de «propositions». Chaque proposition est transmise à un dictionnaire qui détermine les mots à employer et fournit des arbres syntagmatiques. Ceux-ci sont développés en phrases par un «composant syntaxique». Dans ce système, seul le composant syntaxique fonctionne «de gauche à droite». McDonald affirme même que les «messages» doivent être hautement planifiés¹², néanmoins il ne s'occupe

12. Nous ne comprenons pas bien comment McDonald accorde cette affirmation avec son hypothèse de base, qui est qu'un locuteur parlant n'a qu'une vague idée de ce qu'il va dire.

pas de leur formation. Son composant syntaxique repose sur l'hypothèse que seules les dépendances syntaxiques séquentielles peuvent être appréciées lors d'un discours oral. De ce fait, les dépendances de gauche à droite sont prises en compte, mais ce n'est pas le cas des dépendances de droite à gauche, ni des phénomènes qui demandent une vue d'ensemble de la phrase. En conséquence, son générateur peut produire (20a) mais pas (21a) :

- (20) a. Quand Max chante, il agace
Marie. [dépendance de gauche à droite]
- (21) a. Quand il chante, Max agace
Marie. [dépendance de droite à gauche]

En effet, (20a) met en jeu une dépendance de gauche à droite (le pronom *il* est placé à droite du nom auquel il réfère *Max*), tandis que (21a) met en jeu une dépendance de droite à gauche (le pronom est placé à gauche du nom auquel il réfère). De même, son générateur peut produire (20b) mais pas (21b) :

- *Max aime qu'il chante.¹³ [dépendance de gauche à droite]
- (20) b. Max aime chanter.
- *Qu'il chante plaît à Max. [dépendance de droite à gauche]
- (21) b. Chanter plaît à Max.

On ne peut donc juger la valeur de ce générateur qu'en apportant une réponse à la question suivante : est-il plus difficile de formuler les phrases (21) que les phrases (20)?

Le générateur de Hovy et Schank (1984) cherche à rendre compte du fait qu'un locuteur peut changer sa planification au cours de l'élocution : par exemple, un locuteur ayant planifié d'exprimer les informations $I_1, I_2, \dots, I_n$ et ce, dans cet ordre, peut formuler P_1 exprimant I_1 , puis décider d'exprimer I_3 tout de suite en raccrochant P_3 à P_1 . Le fait que l'enchaînement des phrases se fasse par des procédures improvisées expliquerait que les discours oraux ne soient pas toujours cohérents et bien construits.

Comme McDonald, Kempen (1982) suppose l'existence d'un «conceptualiseur» fournissant un «message» qui est transmis à un dictionnaire. Les arbres syntagmatiques fournis par le dictionnaire ne sont

13. Cette phrase est inacceptable dans l'interprétation où *il* réfère à *Max*, interprétation retenue ici.

pas développés complètement par des procédures de «gauche à droite» : il existe des procédures en parallèle qui développent simultanément certains constituants de l'arbre. De ce fait, Kempen peut rendre compte des phénomènes de lapsus comme *Max a laissé son manteau dans son portefeuille* en posant que *son manteau* et *son portefeuille* sont construits en même temps et qu'ils sont mal placés dans la phrase finale.

L'étude des lapsus faite dans Garrett (1980) ainsi que d'autres travaux de psychologues (cf. entre autres Miller et alii 1960, Goldman-Eisler 1980 et Butterworth 1980) montrent que le locuteur planifie à l'avance ce qu'il va dire et que ce qu'il est en train d'exprimer, à un moment donné, dépend de la séquence qui précède comme de celle qui va suivre. Il faudrait donc raffiner l'hypothèse que l'énonciation d'un discours repose sur des procédures de gauche à droite, en excluant toute appréciation globale et en particulier toute structure de discours. À l'instar de McKeown (1982), on peut justifier les structures de discours en considérant qu'un locuteur a des «idées préconçues» sur les procédés utilisables pour réaliser un acte de communication (e.g. répondre à un certain type de question) et sur la façon d'organiser ces procédés pour former un discours.

Ajoutons une remarque sur les aspects oraux d'un discours, tels que les variations mélodiques de la phrase. De nombreux travaux (cf. Martin 1982, Danlos et Émerard 1985) montrent que la structure prosodique de la phrase est basée sur des alternances de montées et de descentes de la fréquence fondamentale, ces alternances étant déterminées **de droite à gauche**. Ainsi, dans une phrase affirmative comme *Luc vient*, le contour mélodique final a une pente descendante. Par le principe d'alternance de montées et de descentes, *Luc* a une pente montante. Par contre, dans une phrase interrogative comme *Luc vient?*, le contour mélodique final a une pente montante, ce qui implique une pente descendante pour *Luc*. Il est donc exclu qu'un locuteur prononce le mot *Luc* sans savoir ce que va être la suite de la phrase.

4.3 Notre générateur

Dans les différents systèmes que nous venons de présenter, les décisions conceptuelles (Dans quel ordre doivent apparaître les informations? Quelles sont les informations qui doivent être exprimées explicitement par rapport à celles qui ne sont exprimées qu'implicitement?) sont prises **avant** les décisions linguistiques (Comment découper le texte en phrases? Quelles

constructions syntaxiques choisir? Quels mots choisir?). De tels systèmes supposent que les décisions conceptuelles peuvent être prises en dehors de toutes considérations linguistiques, et que les décisions linguistiques doivent être prises de façon à respecter les décisions conceptuelles. Dans Danlos (1985), nous montrons, d'une part, que les décisions conceptuelles et linguistiques ne doivent pas être prises indépendamment les unes des autres, d'autre part, que toutes ces décisions, hormis le choix des mots, requièrent l'utilisation d'une «grammaire de discours». L'interdépendance entre les décisions provient principalement du caractère non isomorphe de la relation entre sens et forme. Cette caractéristique du langage naturel se manifeste, entre autres, par des contraintes lexicales qui ne s'expliquent pas par des éléments de sens. De telles contraintes ont été mises en relief dans les travaux sur la phrase simple réalisés au Laboratoire d'Automatique Documentaire et Linguistique (LADL) de Maurice Gross. Ainsi, la transformation actif-passif ne peut pas être appliquée sans tenir compte des items lexicaux en jeu, du verbe d'une part

Cette affaire (concerne + regarde) Luc¹⁴

Luc est (concerné + *regardé) par cette affaire (Gross 1975)

de ses compléments d'autre part :

(De riches politiciens + Luc) habite(nt) Manhattan

Manhattan est habité par (de riches

politiciens + *Luc)

(Gross 1979)

Dans le prolongement des travaux du LADL, nous avons mis en évidence des contraintes lexicales au niveau du discours. Par exemple, un discours enchaînant deux phrases, l'une décrivant un acte, l'autre le résultat de cet acte, n'est acceptable qu'avec certains verbes exprimant le résultat lorsque celui-ci apparaît après l'acte :

Max a (tué + assassiné) Marie. Il lui a tiré une balle dans la tête.

[ordre RÉSULTAT-ACTE]

Max a tiré une balle dans la tête de Marie. Il l'a (tuée + *assassinée).

[ordre ACTE-RÉSULTAT]

14. Nous utilisons des mises en facteur pour alléger le texte : la notation *Cette affaire (concerne + regarde) Luc* réfère aux deux phrases *Cette affaire concerne Luc* et *Cette affaire regarde Luc*. Le symbole *E* désigne la séquence vide. Les symboles «?», «?*» et «*» sont accolés à des phrases ou à des textes qui sont respectivement peu acceptables, guère acceptables ou inacceptables.

Ces diverses contraintes lexicales, tant au niveau de la phrase simple que du discours, font que les formes verbales ne peuvent être choisies indépendamment des constructions syntaxiques et/ou de l'ordre des informations.

Comme nous l'avons déjà signalé, la formation d'un texte bien construit ne se ramène pas à une simple juxtaposition des phrases exprimant les différentes informations de la représentation sémantique : il existe d'autres procédés de linéarisation en phrases que la juxtaposition, par exemple la relativation, la subordination ou la coordination. Lorsqu'un de ces procédés est réalisable formellement, il n'est pas garanti qu'il aboutisse à un texte dégageant la sémantique voulue. Ainsi, considérons la relation causale entre l'acte décrit dans *Un policier a matraqué Marie* et le résultat décrit dans *Un policier a blessé Marie*, et examinons comment articuler ces deux phrases en faisant varier le procédé de linéarisation, leur ordre d'apparition et leurs constructions syntaxiques respectives. Les textes

- (22) a. Un policier a matraqué Marie qui a été blessée.
- (23) a. Marie a été blessée. Elle a été matraquée par un policier.
- (24) a. Un policier a matraqué Marie et il l'a blessée.

peuvent évoquer une relation causale entre l'acte et le résultat, mais ceci n'est pas le cas des textes obtenus respectivement à partir de (22a), (23a) et (24a) en intervertissant *matraquer* et *blessé* :

- (22) b. Un policier a blessé Marie qui a été matraquée.
- (23) b. Marie a été matraquée. Elle a été blessée par un policier.
- (24) b. Un policier a blessé Marie et il l'a matraquée.

Ces textes présentent l'acte et le résultat comme deux événements indépendants. Les différences d'interprétation entre (22a)-(24a) et (22b)-(24b) ne sont pas spécifiques aux phrases *Un policier a matraqué Marie* et *Un policier a blessé Marie*. On les retrouve avec l'acte décrit dans *Max a renversé ce verre* et le résultat décrit dans *Max a cassé ce verre* : les textes

- (22')a. Max a renversé ce verre qui a été cassé.
- (23')a. Ce verre a été cassé. Il a été renversé par Max.
- (24')a. Max a renversé ce verre et il l'a cassé.

dont les structures sont respectivement identiques à celles de (22a), (23a) et (24a), peuvent évoquer une relation causale contrairement aux textes

- (22')b. Max a cassé ce verre qui a été renversé.

(23')b. Ce verre a été renversé. Il a été cassé par Max.

(24')b. Max a cassé ce verre et il l'a renversé.

dont les structures sont respectivement identiques à celle de (22b), (23b) et (24b). Pour une relation causale directe (i.e. un acte qui provoque directement un état), nous avons mis en évidence que les options sur la linéarisation en phrases, l'ordre des informations et les constructions syntaxiques offrent une combinatoire de 36 cas sur lesquels seuls 15 sont réalisables en français, et ce avec la sémantique voulue. Ces 15 structures de discours, qui ne sont prédictibles ni sémantiquement ni syntaxiquement, constituent des données linguistiques d'un type nouveau : une «grammaire de discours». La notion de grammaire de discours que nous introduisons a l'originalité d'intégrer des informations conceptuelles et syntaxiques. La grammaire de discours que nous avons établie ne s'applique qu'à des relations causales directes, mais des grammaires de discours analogues pourraient être établies pour d'autres types de relations sémantiques : relations causales non directes, implications logiques, définitions de dictionnaire, etc. L'utilisation de grammaires de discours semble l'unique garantie de formation de textes corrects quelle que soit la situation en jeu.

L'interdépendance entre les différentes décisions et la nécessité d'utiliser une grammaire de discours conduisent à réaliser un générateur composé de deux modules : le premier, le composant stratégique, prend les décisions conceptuelles et linguistiques mentionnées. Il construit des schémas de phrases qui indiquent la linéarisation en phrases du texte et les formes verbales exprimant les principaux concepts avec leur construction syntaxique. Des exemples de schémas de phrases simplifiés sont :

T1 AUTEUR matraquer OBJET qui être blessé.

T2 CIBLE être assassiné DATE GÉO. AUTEUR tirer sur CIBLE
alors que CIBLE défiler dans VOITURE.

Le deuxième module, le composant syntaxique, développe les schémas de phrases en phrases, par exemple pour T1 et T2 :

Un policier a matraqué Marie qui a été blessé.

John Kennedy a été assassiné hier à Dallas. Un individu a tiré sur le Président des États-Unis alors qu'il défilait dans une voiture décapotable.

Les schémas de phrases sont annotés d'informations syntaxiques non indiquées dans les exemples ci-dessus. Les informations syntaxiques qui doivent être présentes dans les schémas de phrases ainsi que le fonctionnement du composant syntaxique s'appuient sur les travaux du LADL qui a développé une grammaire-lexique du français couvrant 50.000 entrées utilisées dans 600 types de constructions. De ce fait, notre composant syntaxique est basé sur une grammaire à large extension et non sur des fragments de grammaire inspirés par un corpus étroit, comme c'est trop souvent le cas bien que cet état de fait soit déploré par certains chercheurs (Mann 1982).

*Laurence Danlos
LADL, CNRS, France*

Références

- APPELT, D.E. (1982) *Planning Natural-Language Utterances to satisfy Multiple Goals*, Technical Note 259, SRI International, Menlo Park, Californie.
- AUSTIN, J. (1962) *How to do things with words*, J.O. Urmson éd., Oxford University Press.
- BARTLETT, F. (1932) *Remembering, a study in experimental and social psychology*, Cambridge University Press, England.
- BOITET, C. et N. Nedobejkine (1981) «Recent developments in Russian-French machine translation at Grenoble», *Linguistics*, n° 19, pp. 199-271.
- BUTTERWORTH, B. (1980) «Evidence from Pauses in Speech» dans *Language Production*, B. Butterworth éd., Academic Press, London.
- CARBONELL, J.W., W. Boggs et M. Mauldin (1983) «The XCALIBUR project : A Natural Language Interface to Expert Systems» dans *Proceedings of the Eight International Joint Conference on Artificial Intelligence*, Karlsruhe, West Germany.
- CHARNIAK, E. (1977) «A framed PAINTING : The representation of a common sense knowledge fragment», *Cognitive Science*, vol. 1, n° 4.
- CHOMSKY, N. (1957) *Syntactic Structures*, Mouton, La Haye.
- COLMERAUER, A. (1979) «Un sous-ensemble intéressant du français», *R.A.I.R.O.*, vol. 13, n° 4.
- COLMERAUER, A. et J.F. Pique, (1980) «About natural logic», dans *Advances in Data Base Theory*, vol. 1, Gallaire and Minker éd., Plenum Press.
- COOK, M., W. Lehnert et D. McDonald (1984) «Conveying implicit content in narrative summaries» dans *Proceedings of COLING84*, Stanford University, California.
- DANLOS, L. (1985) *Génération automatique de textes en langues naturelles*, Manon, Paris.
- DANLOS, L. et F. Emerard (1985) «Une structure prosodique du français», *Actes du GALF, journées d'études sur la parole*, Paris.
- FRIEDMAN, J. (1969) «A computer system for transformational grammar», *Communications of the Association for Computing Machinery*, vol. 1, n° 4.
- GARRETT, M.F. (1980) «Levels of Processing in Sentence Production» dans *Language Production*, B. Butterworth éd., Academic Press, London.
- GOLDMAN, N. (1975) «Conceptual Generation» dans *Conceptual Information Processing*, Schank éd., North Holland, Amsterdam.
- GOLDMAN — Eisler, F. (1980) «Psychological Mechanisms of Speech Production as studied through the Analysis of Simultaneous Production» dans *Language Production*, B. Butterworth éd., Academic Press, London.
- GRICE, H. (1975) «Logic and Conversation» dans *Syntax and Semantics III : Speech Acts*, P. Cole et J.L. Morgan éd., Academic Press, New York.
- GROSS, M. (1975) *Méthodes en syntaxe*, Hermann, Paris.
- GROSS, M. (1979) «On the failure of generative grammar», *Language*, n° 55, pp. 859-885.
- GROSS, M. (1981) «Les bases empiriques de la notion de prédicat sémantique», *Langages* 63, Larousse, Paris.
- HARRIS, Z., T. Riekman, M. Gottfried et A. Dalladier (1985) *Structure of information in science*, Columbia University Press, New York.
- HIRSCHERBG, J (1984) «Toward a redefinition of Yes/No questions» dans *Proceedings of COLING84*, Stanford University, California.

- HOVY, E. et R.C. Schank (1984) «Language Generation by Computer» dans *Natural Language Processing*, B.G. Bara and G. Guida éd., North Holland, Amsterdam.
- JACOBS, P. (1983) «Generation in a Natural Language Interface» dans *Proceedings of the Eight International Joint Conference on Artificial Intelligence*, Karlsruhe, West Germany.
- JAYEZ, J.H. (1980) «Un survol de recherches sur le traitement automatique du langage naturel», *Linguisticae Investigationes*, IV : 1, pp. 39-101.
- JOSHI, A., B. Webber et R. Weischbedel (1984) «Preventing False Inferences» dans *Proceedings of COLING84*, Stanford University, California.
- KAY, M (1979) «Functionnal Grammar» dans *Proceedings of the Annual Meeting of the Linguistic Society of America*.
- KEMPEN, G. et E. Hoenkamp (1982) *An incremental procedural grammar for sentence formulation*, Internal Report 82 FU 14, Katholieke Universiteit Nijmegen, The Netherlands.
- KURICH, K. (1984) «Fluency in Natural Language reports» dans *Natural Language Generation Systems*, L. Bolc éd., Springer Verlag, Berlin.
- LEHNERT, W. (1981) «Plot Units and Narrative Summarization», *Cognitive Science*, vol 5, n° 4.
- LEHNERT, W. (1983) «Narrative complexity based on Summarization algorithms» dans *Proceedings of the Eight International Joint Conference on Artificial Intelligence*, Karlsruhe, West Germany.
- LYTINEN, S. et R. Schank (1982) *Representation and Translation*, Technical Report 234, Yale University.
- MANN, W. (1982) «Text Generation», *American Journal of Computational Linguistics*, vol. 8, n° 2.
- MANN, W. (1984) «Discourse Structures for Text Generation» dans *Proceedings of COLING84*, Stanford University, California.
- MARTIN, P. (1982) «Phonetic realisations of prosodic contours in french» dans *Speech Communication* 1, North Holland.
- MAULDIN, M. (1984) «Semantic rule based text generation» dans *Proceedings of COLING84*, Stanford University, California.
- McCOY, K. (1984) «Correcting Object-Related Misconceptions : How should the system respond» dans *Proceedings of COLING84*, Stanford University, California.
- McDONALD, D. (1983) «Natural Language Generation as a Computational Problem : an introduction» dans *Computational Models of Discourse*, Brady et Berwick éd., MIT Press, Cambridge, Massachussets.
- McKEOWN, K. (1982) *Generating Natural Language Text in response to Questions about database structure*, PhD Dissertation, University of Pennsylvania.
- McKEOWN, K (1983) «Focus constraints on language generation» dans *Proceedings of the Eight International Joint Conference on Artificial Intelligence*, Karlsruhe, West Germany.
- McKEOWN, K et M. DERR (1984) «Using focus to generate complex and simple sentences» dans *Proceedings of COLING84*, Stanford University, California.
- MILLER, G.A., E. Galanter et K.H. Pribam (1960) *Plans and the Structure of Behavior*, Holt, Rinehart and Winston Inc., New York.

- MINSKY, M (1975) «A framework for representing knowledge» dans *The psychology of computer vision*, P.H. Winston éd., McGraw-Hill, New York.
- PIQUE, J.F. et P. Sabatier (1982) «An informative, adaptable and efficient natural language consultable database system» dans *Proceedings of the 1982 European Conference on Artificial Intelligence*, Orsay, France.
- SAINT-DIZIER, P., P. Bosc et P. Sebilot (1984) «Compréhension du langage naturel et programmation en logique, Application à l'interrogation des bases de données» dans *Actes du Colloque Traitement Automatique du Langage Naturel*, Université de Nantes, Nantes.
- SCHANK, R. (1975) *Conceptual Information Processing*, North Holland, Amsterdam.
- SCHANK, R. et R.P. Abelson (1977) *Scripts, plans, goals and understanding*, Lawrence Erlbaum Associates, Hillsdale, New Jersey.
- SCHANK, R. (1982) *Dynamic memory : A theory of learning in computers and people*, Cambridge University Press, Cambridge, Massachussets.
- SEARLE, J. (1979) «A taxonomy of Illocutionary Acts» dans *Expression and Meaning : Studies in the Theory of Speech Acts*, Cambridge University Press, England.
- SHAPIRO, S. (1982) «Generalized Augmented Transition Network Grammars for Generation from Semantic Networks», *American Journal of Computational Linguistics*, vol. 8, n° 1.
- SIMMONS, R. et J. SLOCUM (1972) «Generating English discourse from semantic networks», *Communications of the Association for Computing Machinery*, vol. 15, n° 10.
- SWARTOUT, W. (1981) *Producing Explanations and Justifications of Expert Consulting Programs*, Technical Report MIT/LCS/TR-251, MIT.
- WHORF, B. (1956) *Language, Thought and Reality*, MIT Press, Cambridge, Mass.
- WILENSKY, R., Y Arens et D. Chin (1984) «Talking to UNIX in English : an overview of UC», *Communication of the ACM*, vol. 27, n° 6.
- WINOGRAD, T. (1977) «A framework for understanding discourse» dans *Cognitive processes in comprehension*, M. Just and P. Carpenter éd., Lawrence Erlbaum Associates, Hillsdale, New Jersey.
- WOODS, W. (1970) «Transition network grammar for natural language analysis», *Communications of the Association for Computing Machinery*, n° 13, pp. 591-606.
- YNGVE, V. (1961) «Random generation of English sentences» dans *Proceedings of the 1961 International Conference on Machine Translation of Languages and Applied Language Analysis*, Her Majesty's Stationery Office, London.