

Chapitre 6

Modélisation du discours : visées et structures du discours

6.1. Introduction

Pour comprendre notre environnement, nous recherchons la meilleure explication pour les faits que nous observons. De même, pour interpréter les textes, nous cherchons la meilleure explication pour les faits « observables » présentés dans le texte. Nous obtenons la version computationnelle de cette approche en traitant l'interprétation d'un texte comme la preuve abductive la plus économique de la forme logique des phrases du texte, « abductif » signifiant que des suppositions sont admises, qui ont un coût, et « économique » signifie que ces coûts sont minimisés [HOB 93].

6.2. Structurer la contiguïté

Parmi les traits observables de notre environnement que nous cherchons à expliquer, il y a la contiguïté ou la proximité des objets ; nous n'y prenons en général pas garde, sauf lorsque les choses sortent de l'ordinaire, par exemple quand nous voyons une chaise sur une table, ou un chien dans une salle de classe. En principe, l'explication donnée pour la présence d'une chose en telle place, est qu'il s'agit là de sa place normale.

Une situation analogue se produit pour le langage. Un texte est une suite de mots, et un des traits d'un texte qui demande explication est la contiguïté de paires de mots, ou de segments plus grands du texte.

L'exemple le plus simple est fourni par les noms composés. Quand nous rencontrons le syntagme « flacon de térébenthine » dans un texte, le problème d'interprétation que nous rencontrons consiste à trouver la relation la plus plausible entre térébenthine et flacon, en utilisant ce nous savons entre ces objets. Dans de nombreux cas, la relation provient d'un des noms lui-même, comme dans « échantillon d'huile », où la relation entre l'huile et l'échantillon est simplement la relation « échantillon de ». « Huile » fournit l'un des arguments du prédicat « échantillon de ».

La syntaxe et la sémantique compositionnelle peuvent être vues comme émergeant du même besoin d'expliquer la contiguïté. Quand nous voyons la suite de mots « les hommes travaillent », nous devons trouver quelque relation entre eux. L'hypothèse selon laquelle les phrases ont une structure syntaxique revient à accepter un ensemble de contraintes quant aux relations que l'on peut obtenir entre deux mots ou dans des segments plus larges d'un texte. Dans ce cas, la contrainte est que le dernier mot fournisse la relation. Les hommes doit être l'agent du travail, tandis que dans le cas de « échantillon d'huile », « échantillon » fournit une relation possible, dans le cas de « les hommes travaillent », « travaille » fournit la relation exigée. (Ce n'est pas tout à fait vrai ; une métonymie est possible, de telle sorte que le dernier mot doive seulement fournir une relation entre l'événement et quelque chose qui soit fonctionnellement relié aux mots qui précèdent)¹.

La structure en arbre des phrases provient du fait que la relation de contiguïté peut exister entre des segments de texte plus larges que simplement des mots seuls, et où les segments ont leur structure arborescente interne propre résultant des relations de contiguïté. Ainsi, dans :

John croit que les hommes travaillent,

nous ne cherchons pas à expliquer la contiguïté entre « croit » et « hommes ». Nous préférons expliquer d'abord la contiguïté entre « hommes » et « travaillent », et seulement ensuite la contiguïté entre « croit » et « les hommes travaillent » (ou la contiguïté entre « John », « croit » et « les hommes travaillent », selon votre vision de la syntaxe de la proposition). Ce regroupement apparaît même en l'absence de contraintes syntaxiques. Considérons les deux noms composés « *Stanford Research*

1. Cette approche de la syntaxe m'est apparue à la suite d'une conversation avec Mark Johnson.

Institute » et « *Cancer Research Institute* ». Dans le dernier cas, nous devons d'abord trouver la relation entre le cancer et la recherche, et ensuite trouver la relation entre la recherche sur le cancer et l'institut, alors que pour le premier cas, nous regroupons « Recherche » avec « Institut » et ensuite « Stanford » avec « Institut de Recherche ».

Afin d'expliquer les contiguïtés entre des segments de texte de taille plus importante que le mot, nous avons besoin d'avoir une idée de l'information principale transmise par les segments. Par exemple, dans « institut de recherche », la référence à l'institut est, en un sens, plus basique que celle à la recherche ; un institut de recherche est un type d'institut, plutôt qu'un type de recherche. De la même façon, dans « les hommes travaillent », l'information portant sur certaines entités en train de travailler est plus basique que l'information portant sur ces entités comme étant des hommes. En conséquence, quand nous composons des segments de plus en plus importants dans un texte, nous devons nécessairement avoir quelque idée de l'information principale transmise par ces segments composés.

Les règles de la syntaxe et de la sémantique compositionnelle sont un ensemble de contraintes sur la façon dont les segments d'un texte peuvent être groupés ensemble et sur ce qu'est l'information basique transmise par le segment composé. Au niveau de la proposition principale, l'information de base est souvent celle qui est transmise par le verbe principal et/ou les adverbiaux, bien que ceci puisse être dépassé par d'autres facteurs comme l'intonation, la nouveauté, la topicalisation et autre. Nous pouvons appeler cette information de base, quel que soit son mode de détermination, l'*assertion* de la proposition. C'est, en un sens, un résumé de ce que la phrase transmet.

L'origine de la structure du discours est précisément la même que l'origine de la structure intraphrastique. Deux expressions, deux propositions ou deux phrases, ou de plus larges segments de discours sont contigus, et ceci requiert une explication. Une relation entre eux doit être trouvée pour expliquer cette contiguïté. Tandis que dans les structures syntaxiques des phrases, les relations entre éléments adjacents sont le plus souvent des relations prédicat-arguments, le cas des segments de discours plus importants de textes est plus proche du cas des noms composés, car la relation qui explique la contiguïté peut être en principe n'importe quelle relation plausible entre les situations décrites par les segments du texte.

En fait, quoique je milite pour une transition continue de la syntaxe au discours, si on cherchait à définir l'unité de discours minimal [POL 88], un bon choix serait l'unité ci-dessus dans laquelle les relations prédicat-argument ne sont plus les interprétations dominantes de la contiguïté. Dans les textes écrits, cela tend à être des propositions ou des phrases ; à l'oral, il s'agit souvent de syntagmes, ou d'éléments plus petits.

Comme dans le cas de la syntaxe, tandis que nous composons des segments de texte, nous devons avoir quelque idée de l'information principale véhiculée par le segment complexe. Nous devons être capables de spécifier le « contenu », ou « résumé » d'un segment de texte combinant plusieurs propositions. C'est à la fois plus facile et plus difficile que pour les noms composés. Plus difficile parce que s'il est vrai que c'est presque toujours, en anglais, le second nom d'une paire nom-nom qui porte l'information principale, pour une paire de propositions, l'information principale peut être portée par la première, par la seconde, ou peut émerger de l'une et de l'autre. Plus facile parce que le nombre de relations typiques qui peut s'établir entre deux situations ou éventualités décrites dans des segments supra-propositionnels est inférieur à celui qui peut s'établir entre deux noms. Pour les noms composés, [DOW 77] et d'autres, ont établi de manière convaincante que les relations entre deux noms peuvent être quasiment quelconques, pourvu qu'on se donne le contexte voulu. Mais [LEV 78] a établi de manière convaincante que la grande majorité de ces relations peuvent être vues comme des exemples d'une petite douzaine environ de relations abstraites telles que prédicat-argument, fonction, appartenance, etc.

De même, il est possible que la relation entre deux segments supra-propositionnels d'un texte soit quelconque dans un contexte approprié. Le récepteur peut simplement imaginer la relation la plus plausible entre les situations décrites et le contenu ou résumé le plus plausible pour le segment complexe. Une telle vue de la cohérence discursive, malheureusement, ne donne aucune aide pour déterminer ce qu'est le contenu, ou résumé d'un segment complexe.

Fondamentalement, cependant, la relation entre des segments supra-propositionnels peut être vue comme réalisation de l'une des trois relations abstraites suivantes : causalité, figure-fond, et similarité. J'appellerai ces relations des *relations de cohérence*. Une théorie de la cohérence discursive qui admet cela doit développer des caractérisations de chacune de ces relations, expliquer les diverses variétés ou exemples de chaque relation et pour chacune d'elles, définir le contenu ou résumé du segment complexe. C'est ce que j'ai essayé de faire dans des travaux antérieurs (par exemple, [HOB 78, HOB 85]) et c'est ce que je vais tenter de reformuler ici dans un cadre abductif.

Je vais m'intéresser ici surtout à la façon dont ces trois relations s'exercent entre les segments suprapropositionnels du texte, ces liens existent aussi entre des éléments internes aux propositions indépendantes. Ainsi, les éléments d'une liste illustrent la relation de similarité requise dans le parallélisme [POL 88]. Dans la phrase :

Une voiture a heurté un jogger à Palo Alto hier soir,

il y a une relation causale implicite entre le jogger et l'automobiliste. Ces relations vont au-delà de ce qui nous est donné par la sémantique compositionnelle.

J'ai mis l'accent ici sur le problème de l'*interprétation* et je vais continuer à le faire, mais il est important de garder à l'esprit que l'interprétation et la génération sont indissociablement inter-reliées. Un locuteur cherche à émettre des énoncés destinés à être compris. Lorsque deux segments de discours cohérent sont émis à la suite, c'est que le locuteur espère que l'auditeur va trouver la relation qui est intentionnellement produite par la contiguïté. A l'inverse, l'auditeur doit souvent réfléchir pour trouver ce que le locuteur essaie de réaliser et aux autres façons qu'il aurait pu utiliser pour le faire, de façon à obtenir la meilleure interprétation pour un fragment de discours.

Dans ce papier, j'adopte une perspective informationnelle sur le discours, dans laquelle l'information portée par le discours est pour la majeure partie prise au premier degré. En fait, les interprétations dérivées comme je les décris ici sont *considérées* plutôt que *crues* ; la révision de la croyance peut être considérée comme un processus distinct, venant logiquement après l'interprétation. En outre, le cheminement de l'information dans le discours est un acte intentionnel. Dans [HOB 95a], je montre comment la perspective informationnelle est incluse dans une perspective d'abduction intentionnelle et comment les deux perspectives interagissent dans l'interprétation.

6.3. Le cadre : l'interprétation comme phénomène d'abduction

6.3.1. La résolution des questions pragmatiques locales par l'abduction

L'abduction est l'inférence qui donne la meilleure explication. Le processus d'interprétation des phrases en discours peut être considéré comme le processus qui fournit la meilleure explication de l'information portée par les phrases, c'est-à-dire la meilleure justification à ce que les phrases soient prises pour vraies. Ceci peut être matérialisé de façon procédurale comme suit.

Pour interpréter une phrase :

- (1) Prouver la forme logique de la phrase,
 en tenant compte des contraintes que les prédicats imposent à leurs arguments,
 et permettant les coercions
 en réduisant les redondances là où c'est possible,
 en faisant des hypothèses là où c'est nécessaire.

A la première ligne, prouver signifie « dériver au sens logique, à partir des axiomes du calcul des prédicats dans la base de connaissance, la forme logique qui a été produite par l'analyse syntaxique et sa traduction sémantique dans la phrase ».

Dans une situation de discours, le locuteur et l'auditeur ont chacun leurs propres ensembles de croyances privées, et il existe un large ensemble commun de croyances partagées. Un énoncé vit à la frontière entre la croyance partagée et les croyances privées du locuteur. C'est une offre d'extension du terrain des croyances partagées par inclusion de quelque croyance privée du locuteur. L'énoncé est ancré référentiellement dans la croyance partagée et quand nous réussissons à prouver sa forme logique avec les contraintes, nous reconnaissons cet ancrage référentiel. C'est l'information donnée, celle qui est définie, présupposée. Quand il est nécessaire de faire des suppositions, l'information vient des croyances privées du locuteur, et donc il s'agit d'une information nouvelle, celle qui est indéfinie, assertée. Réduire les redondances est une façon d'obtenir une interprétation minimale et, par là, meilleure.

Considérons un seul exemple :

Le bureau de Boston a appelé.

Cette phrase pose au moins trois problèmes d'interprétation au niveau pragmatique local, les questions de la référence de « le bureau de Boston » en élargissant la métonymie à « la personne du bureau de Boston » et en déterminant la relation implicite entre Boston et le bureau en question. Laissons de côté ces questions pour le moment et considérons la phrase en fonction de la procédure (1). Nous devons prouver la forme logique de la phrase en tenant compte des contraintes que « appeler » impose à l'actant, et qui autorise la coercion. C'est-à-dire que nous devons prouver par abduction l'expression (indépendamment du temps verbal et autres complications) :

$$(3) \exists x, y, z \text{ a-appelé } (e, x) \ \& \ \text{personne}(x) \ \& \ \text{rel}(x, y) \ \& \ \text{bureau}(y) \ \& \ \text{Boston}(z) \\ \& \ nn(z)$$

Soit : il y a un événement d'appel e , effectué par une personne x , x pouvant ou non être identique au sujet explicite de la phrase, mais étant au moins lié à lui, ou coercible à partir de lui, ce que nous représentons par $\text{rel}(x, y)$. y est un bureau, et il entretient quelque relation non spécifiée nn à z , Boston. $\text{personne}(x)$ est la contrainte que *a appelé* impose sur son agent x .

La phrase peut être interprétée par rapport à une base de connaissances partagées² qui contient les faits suivants :

Boston (B_1)	(B_1 est la ville de Boston)
bureau (O_1) & à (O_1, B_1)	(O_1 est un bureau sis à Boston)
personne (J_1)	(J_1 est une personne)
travaille-pour (J_1, O_1)	(John travaille pour O_1)
$\forall y, z$ dans(yz) $\rightarrow$ $nn(z, y)$	(si y est dans z , z peut entretenir une relation à y)
$\forall x, y$ travaille-pour (x, y) $\rightarrow$ rel (x, y)	(si x travaille pour y , alors y peut être coercé en x)

La preuve de la totalité de (3) est directe, sauf pour *a-appelé* (e, x). Nous supposons donc que c'est l'information nouvelle apportée par la phrase.

Cette interprétation est illustrée dans le graphe de preuve de la figure 6.1, où un rectangle est tracé autour du littéral admis *a-appelé* (e, x). De tels graphes de preuves jouent le même rôle que les arborescences des analyses syntaxiques. Ce sont des figurations de l'interprétation, et nous en verrons d'autres dans ce chapitre.

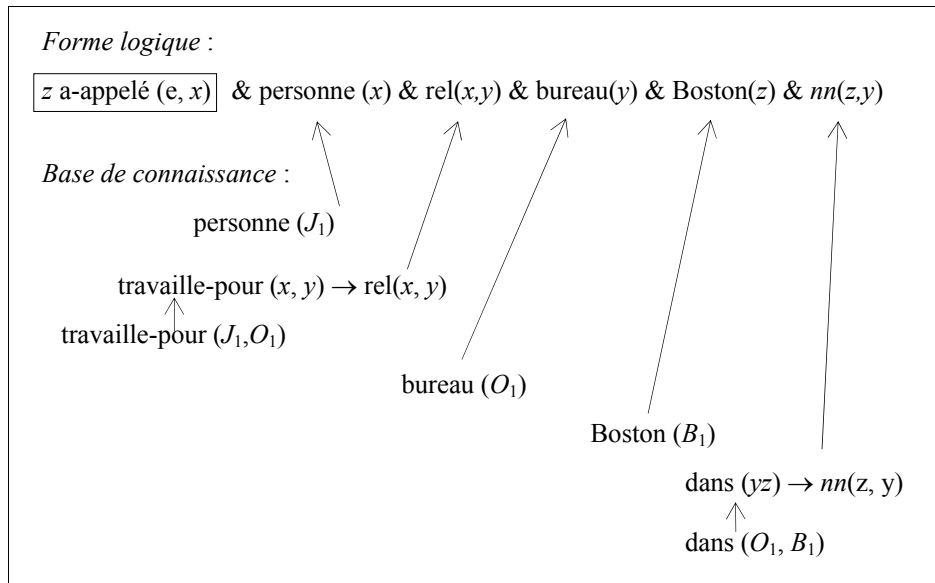


Figure 6.1. Interprétation de « Le bureau de Boston a appelé »

2. Dans ce chapitre, nous admettrons que tous les axiomes sont des connaissances partagées par le locuteur et le récepteur, qu'elles font partie de leur arrière-plan culturel commun.

Notons maintenant que les trois questions pragmatiques locales ont été résolues comme des sous-produits. Nous avons résolu « bureau de Boston » en O_1 . Nous avons déterminé la relation implicite contenue dans le nom composé comme étant « dans ». Et nous avons élargi la métonymie à « John, qui travaille pour le bureau de Boston, a appelé ».

D'autres preuves sont possibles pour (3). Par exemple, nous pourrions avoir asserté « personne(x) » plutôt que de le prouver par l'unification de « personne (J_1) ». Cette preuve correspond à l'interprétation la moins spécifique, à savoir que *quelqu'un* qui travaille au bureau de Boston a appelé.

En général, il devrait y avoir une méthode pour choisir entre plusieurs interprétations possibles. Un schéma des coûts est présenté dans [HOB 93], dans lequel un coût est attribué à chaque preuve, la meilleure preuve étant celle qui coûte le moins. Rien d'autre ne sera dit à ce sujet dans cet article. Nous nous contenterons de montrer que les interprétations attendues sont dans l'espace des interprétations possibles, sans fournir des arguments pour dire qu'elles sont les meilleures interprétations dans cet espace.

6.3.2. Intégration de la syntaxe, de la sémantique computationnelle et de la pragmatique locale

Si on combine l'idée de l'interprétation comme abduction avec l'idée plus ancienne du *parsing* comme déduction ([KOW 80], p. 52-53), [PER 83], il devient possible d'intégrer la syntaxe, la sémantique, et la pragmatique d'une manière très complète et élégante³.

Nous présenterons cela dans les termes de l'exemple (2), repris ici par commodité :

(2) Le bureau de Boston a appelé.

Rappelons que pour interpréter cela nous devons prouver l'expression :

(3a) $\exists x, y, z$ a-appelé (e, x) & personne (x) & rel(x, y)

(3b) & bureau (y) & Boston (z) & nn(z)

Considérons maintenant une grammaire simple capable d'analyser cette phrase et écrite dans le style de Prolog :

$\forall(w_1, w_2) \text{ GN}(w_1) \ \& \ \text{V}(w_2) \rightarrow \text{P}(w_1, w_2)$

$\forall(w_1, w_2) \text{ Det}(\text{le}) \ \& \ \text{N}(w_1) \ \& \ \text{N}(w_2) \rightarrow \text{GN}(\text{le } w_1 \ w_2)$

3. Cette idée est due à Stuart Shieber.

Soit : si la suite w_1 est un groupe nominal, et que la suite w_2 soit un verbe, alors la concaténation $w_1 w_2$ est une phrase. La seconde règle est interprétée de la même manière. Analyser une phrase W consiste à prouver $s(W)$.

Nous pouvons intégrer la syntaxe, la sémantique compositionnelle et la pragmatique locale en augmentant les axiomes de cette grammaire par des éléments de la forme logique aux places appropriées, de la manière suivante :

$$(4) \forall (w_1, w_2, y, p, e, x) \text{GN}(w_1, y) \ \& \ \text{V}(w_2, p) \ \& \ p'(e, x) \ \& \ \text{rel}(x, y) \ \& \ \text{Req}(p, x) \rightarrow \text{P}(w_1 w_2, e)$$

$$(5) \forall (w_1, w_2, q, r, y, z) \text{Det}(le) \ \& \ \text{N}(w_1, r) \ \& \ \text{N}(w_2, q) \ \& \ r(z) \ \& \ q(y) \ \& \ nn(z, y) \rightarrow \text{GN}(le \ w_1 \ de \ w_2, y)$$

Le second argument du prédicat lexical N et V dénote le prédicat correspondant au mot, comme *Boston*, *bureau*, ou *a appelé*. La formule atomique GN (w_1, y), signifie que la suite w_1 est un groupe nominal référant à y . La formule atomique Req (p, x) signifie que le prédicat p doit avoir x pour argument. Cette contrainte peut ensuite être renforcée s'il y a un axiome :

$$\forall (x) \text{personne}(x), \rightarrow \text{Req}(\text{appeler}, x)$$

Cet axiome dit qu'une manière de satisfaire la contrainte est de prendre pour x une personne. L'axiome (4) peut donc être paraphrasé comme suit: « si w_1 est un GN référant à y , et que w_2 soit un verbe dénotant le prédicat p , que p' soit vrai d'une éventualité e et d'une entité x , que x soit relié à y (ou coercible à partir de y), et que x satisfasse la contrainte qu'impose p' sur son second argument, alors la concaténation w_1, w_2 , est une phrase décrivant l'éventualité e ". L'axiome (5) peut être paraphrasé comme suit : « Si *le* est un déterminant, que w_1 soit un N dénotant le prédicat r , que w_2 soit un nom dénotant le prédicat q , et que le prédicat r soit vrai de quelque entité z , et le prédicat q de quelque entité y , qu'il y ait une relation implicite nn entre z et y , alors la concaténation *le* w_1 *de* w_2 est un GN référant à l'entité y ». Notons que la conjonction de (3a) a été incorporée dans l'axiome (4), et la conjonction de (3b) dans l'axiome (5)⁴.

Quand nous avons prouvé que $p(W)$, nous avons prouvé que W était une phrase. Maintenant, si nous prouvons $(\exists e) p(W, e)$, nous prouvons qu'il existe une expression interprétable W et que l'événement e est son interprétation.

4. Tels qu'ils sont présentés, ces axiomes sont de second ordre, mais ils ne le sont pas en fait, puis que les variables-prédicat sont à instancier comme des constantes de prédicat, jamais comme des expressions lambda. Il est donc facile de les transformer en axiome du premier ordre au moyen d'une constante individuelle correspondant à chaque constante de prédicat.

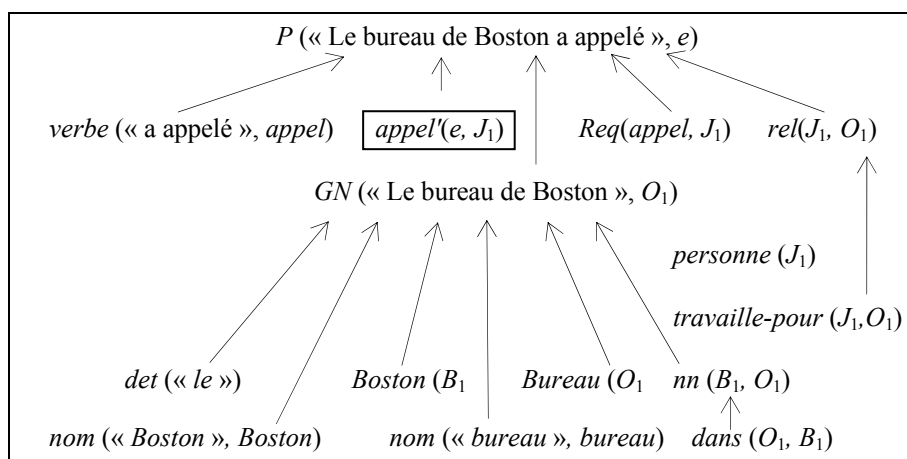


Figure 6.2. Analyse et interprétation de « Le bureau de Boston a appelé »

Chaque axiome de la « grammaire » a donc un composant « syntaxique » – les conjonctions comme $GN(w_1, y)$ et $Verbe(w_2, p)$ – qui spécifie la structure syntaxique, et un composant « pragmatique » – les conjonctions comme $p'(e, x)$ et $rel(x, y)$ – qui pilote l'interprétation. Autrement dit, la pragmatique locale est capturée en vertu du fait que pour pouvoir prouver $(\exists e) p(W, e)$, on doit dériver la forme logique de la phrase en tenant compte en même temps des contraintes que les prédicats imposent à leurs arguments, autorisant la métonymie, voir la description (1). La sémantique compositionnelle de la phrase est spécifiée par la façon dont les dénominations données par la part syntaxique sont utilisées pour construire la part pragmatique.

A noter que la structure de chacun des axiomes (4) et (5) est en gros la suivante : pour prouver qu'un tout est une entité signifiante, il faut prouver que chacune de ses parties est une entité signifiante et qu'il existe une relation signifiante entre elles. En (4), la relation est la coercion de la relation prédicat-argument $p'(e, x)$ & $rel(x, y)$. En (5), c'est la relation nn .

Dans ce cadre, interprétation et génération ne diffèrent que par les conditions initiale et finale du processus de preuve. Interpréter une phrase W , c'est prouver $(\exists e) p(W, e)$. L'événement résultant e est l'information nouvelle qui est donnée. En génération, l'éventualité E est connue, et la phrase doit être dérivée, c'est-à-dire que l'on doit prouver $(\exists w) p(w, E)$. Cette idée est plus largement développée dans [HOB 90, HOB 93].

Une prise en compte beaucoup plus étendue et quelque peu différente de la syntaxe de l'anglais dans la perspective de « L'interprétation comme Abduction » est donnée dans [HOB 95b] ; elle suit la grammaire HPSG [POL 94] d'assez près.

6.3.3. Conventions de notation

Avant de poursuivre, je voudrais introduire quelques conventions de notation utilisées ici.

$p(x)$ signifie que p est vrai de x , et $p'(e, x)$ signifie que e est l'éventualité, ou situation possible dans laquelle p est vrai de x . Cette éventualité peut exister ou non dans le monde réel. Les prédicats primes et les prédicats sont reliés par l'axiome suivant :

$$\forall(x) p(x) \equiv \exists e p'(e, x) \ \& \ \text{Rexiste}(e)$$

où $\text{Rexiste}(e)$ signifie que l'éventualité e existe réellement. Cette notation, en réifiant les événements et les conditions fournit une manière de spécifier des propriétés de second ordre en logique du premier ordre. Cette réification davidsonnienne est une stratégie courante en intelligence artificielle. Voir Hobbs (1985) pour des explications supplémentaires et l'ontologie sous-jacente.

Nous écrirons souvent les axiomes qui intuitivement devraient être écrits :

$$\forall(x) p(x) \rightarrow q(x)$$

sous la forme :

$$\forall(e_1, x) p'(e_1, x) \rightarrow \exists(e_2) q'(e_2, x)$$

C'est-à dire : si e_1 est l'éventualité selon laquelle p est vrai de x , alors il existe une éventualité e_2 selon laquelle q est vrai de x . Il sera parfois utile de noter cela sous une forme plus forte. Ce n'est pas seulement que si e_1 existe, alors e_2 existe également, c'est que l'éventualité e_2 existe *du fait que* e_1 existe. Exprimons cette connexion étroite par le prédicat *gen*, pour « génère ». L'axiome précédent peut alors être renforcé en :

$$\forall(e_1, x) p'(e_1, x) \rightarrow \exists(e_2) q'(e_2, x) \ \& \ \text{gen}(e_1, e_2)$$

En outre, il pourrait sembler, que puisque dans l'abduction nous faisons uniquement usage du chaînage arrière pour trouver une preuve et un ensemble de propositions, nous ne pouvons pas utiliser d'information relative aux ensembles supérieurs. Cependant, le fait que nous puissions faire des suppositions nous permet de retourner les axiomes. En général, les axiomes de forme :

Espèce $\rightarrow$ genre

peuvent être convertis en axiomes de forme

espèce $\&$ différences spécifiques $\equiv$ genre

Souvent, nous ne serons pas en mesure de prouver les différences et, en de nombreux cas, nous ne pouvons même pas les formuler. Mais dans notre schéma abductif, cela n'a pas d'importance ; nous pouvons seulement les supposer. En fait, nous n'avons pas à les formuler explicitement. Nous pouvons seulement introduire un prédicat, un par axiome, qui représente toutes les propriétés restantes. Il ne sera jamais prouvable, mais nous pouvons le supposer. Ainsi, outre des axiomes comme :

$$\forall(y) \text{ éléphant}(y) \rightarrow \text{maladroit}(y)$$

nous pouvons avoir des axiomes comme

$$\forall(y) \text{ maladroit}(y) \ \& \ \text{etc}(y) \rightarrow \text{éléphant}(y)$$

Ainsi, même si nous faisons du chaînage arrière strict à la recherche d'une explication, le fait que quelque chose soit maladroit peut toujours être utilisé comme un indice (fragile, peut-être) qu'il s'agit d'un éléphant, puisque nous admettons la prédication et cœtera *etc(x)*, au prix d'un certain coût.

Cette stratégie peut sembler *ad hoc*, au premier abord. Mais je considère que cette stratégie fournit une solution technique assez générale au problème de la non monotonie dans le raisonnement de sens commun, et à la question du vague dans la signification des langues naturelles, une stratégie très semblable à l'usage des prédicats d'anormalité dans la logique de la *circonscription* [MCC 87]. Alors que dans le cadre de la *circonscription* on spécifie typiquement un ordre partiel des prédicats d'anormalité en vertu duquel ils sont minimisés, dans le cadre de l'abduction pondérée, on utilise un système de coût un peu plus flexible.

Il n'y a pas de difficulté particulière à spécifier la sémantique des prédicats « *et cœtera* ». Formellement, *etc_i* dans l'axiome ci-dessus peut être utilisé pour dénoter l'ensemble de tous les individus qui sont soit non maladroits, soit des éléphants maladroits. Intuitivement, *etc* contient toute l'information qu'on aurait besoin de connaître au-delà de « maladroit » pour conclure à ce qu'un individu soit un éléphant. Comme dans le cas de presque tous les prédicats dans une axiomatisation des connaissances de sens commun, il est sans espoir de chercher à formuler les conditions nécessaires et suffisantes pour un prédicat « *et cœtera* ». En fait, l'usage de tels prédicats est dû en général largement à la reconnaissance de ce phénomène à propos des connaissances de sens commun.

Les prédicats « *et cœtera* » pourraient être utilisés comme les prédicats d'anormalité le sont en logique de la *circonscription*, avec des axiomes séparés indiquant les conditions dans lesquelles ils sont valides. Cependant, dans l'optique adoptée ici, des conditions plus détaillées seraient formulées par l'extension des axiomes de la forme :

$$\forall(x) p_1(x) \ \& \ \text{etc}_1(x) \rightarrow q(x)$$

en axiomes de la forme :

$$\forall (x) p_1(x) \& p_2(x) \text{ etc}_1(x) \rightarrow q(x)$$

Un prédicat « *et cætera* » n'apparaîtrait que dans l'antécédent d'un axiome simple et jamais dans le conséquent. Ainsi, les prédications « *et cætera* » signalent seulement qu'il y a des suppositions ayant un coût. Elles ne sont jamais prouvées. Ce sont de simples suppositions.

Ces prédicats constituent un des principaux moyens « d'arrondir les angles » de notre logique. Nous aimerions qu'ils se répandent dans la base de connaissance. Virtuellement, à chaque fois qu'il y a un axiome reliant une espèce à un genre, il devrait y avoir un axiome correspondant, incorporant une prédication « *et cætera* » qui exprime la relation inverse.

Résumons, en ce point, la forme la plus élaborée des axiomes de notre base de connaissance. Si nous voulons exprimer une relation d'implication entre les concepts p et q , la manière la plus naturelle de le faire est un axiome de type :

$$\forall (x, z) p(x, z) \rightarrow \exists y q(x, y)$$

où z et y sont des arguments qui apparaissent dans une prédication, mais non dans l'autre. Quand nous introduisons des éventualités, les axiomes deviennent :

$$\forall (e_1, x, z) p'(e_1, x, z) \rightarrow \exists (e_2, y) q'(e_2, x, y)$$

En utilisant la relation *gen* pour exprimer la connexion étroite entre les deux éventualités, l'axiome devient :

$$\forall (e_1, x, z) p'(e_1, x, z) \rightarrow \exists (e_2, y) q'(e_2, x, y) \& \text{gen}(e_1, e_2)$$

Ensuite, nous introduisons une proposition « *et cætera* » dans l'antécédent pour tenir compte de l'imprécision de notre savoir concernant l'implication :

$$\forall (e_1, x, z) p'(e_1, x, z) \& \text{etc}_1(x, z) \rightarrow \exists (e_2, y) q'(e_2, x, y) \& \text{gen}(e_1, e_2)$$

Finalement, nous transformons la relation entre p et q en biconditionnelle en écrivant également la relation converse :

$$\forall (e_1, x, z) p'(e_1, x, z) \& \text{etc}_1(x, z) \rightarrow \exists (e_2, y) q'(e_2, x, y) \& \text{gen}(e_1, e_2)$$

$$\forall (e_1, x, y) p'(e_2, x, y) \& \text{etc}_2(x, y) \rightarrow \exists (e_1, z) p'(e_1, x, z) \& \text{gen}(e_2, e_1)$$

Telle est donc la forme d'expression formelle la plus générale dans notre logique abductive de ce que nous percevons intuitivement comme une association entre les concepts p et q .

Dans cet article, par commodité de notation, nous utilisons la forme la plus simple possible des axiomes. Le lecteur doit cependant se souvenir que ce ne sont que des abréviations pour la forme complète, biconditionnelle des axiomes⁵.

6.4. Axiomatisation des structures arborescentes du discours

Les structures arborescentes du discours peuvent être captées par deux axiomes :

$$\begin{aligned} &\forall(w, e) S(w, e) \rightarrow \text{Segment}(w, e) \\ &\forall(w_1, e_1, w_2, e_2, e) \text{Segment}(w_1, e_1) \& \text{Segment}(w_2, e_2) \& \\ &\quad \text{Relation_de_cohérence}(e_1, e_2, e) \rightarrow \text{Segment}(w_1, w_2, e) \end{aligned}$$

Le premier axiome stipule qu'une phrase w , décrivant une éventualité e est un segment de discours cohérent décrivant e . Le second dit que si deux segments w_1 et w_2 décrivent les éventualités e_1 et e_2 respectivement, que e_1 et e_2 soient reliés par une relation de discours, alors la concaténation $w_1 w_2$ est un segment de discours cohérent.

La variable e dans le second axiome est l'assertion du contenu du segment complexe $w_1 w_2$. Il est déterminé par l'assertion des segments consituants e_1 et e_2 et de la relation qui s'établit entre eux du fait que $w_1 w_2$ lui-même est un segment de discours cohérent.

Prouver $\text{Relation_de_cohérence}(e_1, e_2, e)$ revient à expliquer la contiguïté de w_1 et w_2 .

Interpréter un texte W est donc prouver l'expression

$$\exists(e) \text{Segment}(W, e)$$

C'est-à-dire que W est un segment de discours cohérent qui véhicule, ou décrit la situation ou éventualité e , e étant l'assertion ou contenu de W .

Développer une théorie de la cohérence discursive et de la structure du discours dans cette perspective, consiste à spécifier les diverses manières dont :

$\text{Relation_de_cohérence}(e_1, e_2, e)$
peut être établie, et à spécifier e .

5. Les axiomes complets ne sont pas sous forme de clauses de Horn, mais ce n'est pas essentiel. Ils peuvent être skolemisés et éclatés en deux axiomes ayant les mêmes fonctions de Skolem. Cette remarque vaut aussi pour les autres axiomes de cet article qui ont des conjonctions dans le conséquent.

La distinction courante entre des relations de cohérence hypotaxiques, avec un segment dominant et un segment dominé, des relations de cohérence parataxiques est facile à capter dans ce cadre. Si la relation est hypotaxique, alors e est soit e_1 soit e_2 , selon que w_1 ou w_2 est dominant. Si la relation de cohérence est parataxique, alors e doit être calculé à partir de e_1 et e_2 pris ensemble. C'est en fait ce que signifient hypotaxe et parataxe.

Cette approche évoque les grammaires de discours. Ce qui a toujours fait problème avec les grammaires de discours, c'est que leurs symboles terminaux (par exemple *Introduction*) et parfois leur composition ne sont pas calculables effectivement. Puisque dans notre cadre abductif nous sommes en mesure de raisonner sur le contenu des énoncés du discours, ce problème disparaît.

6.5. Les relations de cohérence

6.5.1. Les sources de la cohérence

A un niveau suffisamment abstrait, quand les relations entre propositions ou phrases successives d'un texte ne sont pas explicitement indiquées, il y a trois relations qui s'établissent prioritairement : causalité, figure-fond, et similarité. On ne devrait pas être surpris que ces relations soient si saillantes. On voit bien pourquoi la causalité intéresse des créatures comme nous qui devons nous frayer un chemin parmi des événements qui échappent à notre contrôle ; les prédictions sont un facteur de survie. Notre intérêt pour les relations de figure-fond et de similarité peut aussi se réduire à la causalité. Une entité, la figure, est causalement influencée par l'environnement (le fond) dans lequel elle se situe, et des entités similaires se comportent causalement d'une manière analogue (et quand ce n'est pas le cas, cela mérite d'être noté). Par conséquent, connaître ces relations est une aide à la prédiction.

On peut soutenir que cette observation trouve son origine dans Hume. Dans son *Inquiry Concerning Human Understanding* (section III), il soutient qu'il y a des principes généraux à la base des discours cohérents qui reposent sur des principes généraux relatifs à l'association des idées. « Transcrivons une conversation quelconque, la plus futile et la plus libre qu'on veut, nous observerons qu'il y a une connexion pour toutes ses transitions. Ou lorsque cela fait défaut, la personne qui rompt le fil du discours, vous informera qu'il s'est produit dans son esprit une succession de pensées qui l'ont peu à peu éloignée du sujet de la conversation. » En outre, les trois principes proposés par Hume sont très proches de nos propres relations de causalité, de figure-fond, et de similarité ». « Il n'y a pour moi que trois

principes de connexion entre les idées, c'est-à-dire la *ressemblance*, la *contiguïté* dans le temps ou l'espace, et la *cause*, ou *effet*. »

Dans la suite de cette section, je décris les relations de cohérence qui sont requises par le fragment de conversation analysé dans la section 6.7 : Causalité, Parallélisme, Elaboration et Contraste. Je donne des axiomes définissant chaque relation de cohérence et des exemples, de même que les schémas d'inférence pertinents.

Dans les exemples, nous rencontrerons beaucoup d'autres problèmes d'interprétation, tels que la résolution des pronoms et l'interprétation d'une métaphore, qui seront résolus comme simple conséquence annexe du processus d'identification des relations de cohérence.

6.5.2 Causalité

La causalité peut relier deux segments de discours de l'une des manières suivantes. Nous pouvons décrire un événement ou une situation et ensuite décrire une conséquence logique : dans ce cas, la cause est énoncée en premier et l'effet en second lieu, et l'un ou l'autre peut être le segment dominant, bien que généralement l'effet soit le segment dominant. Ou alors, nous pouvons décrire une situation, et ensuite décrire, en guise d'explication, la situation qui lui a donné naissance : dans ce cas, l'effet vient en premier et la cause en second lieu, et l'effet est le membre dominant de cette paire. Je donnerai un exemple détaillé de la relation d'explication.

Pour qu'un fragment de discours S_2 fonctionne comme explication d'un autre fragment de discours S_1 , S_2 doit décrire un événement ou une situation qui entraîne ou pourrait entraîner la situation ou l'événement décrit en S_1 . C'est ce qu'exprime l'axiome suivant (du moins comme une première approximation) :

$$\forall (e_1, e_2) \text{ cause } (e_2, e_1) \rightarrow \text{Explication } (e_1, e_2)$$

Soit : si ce qui est asserté par le second segment peut causer ce qui est asserté par le premier, alors il y a une relation d'explication entre les deux segments. L'explication est généralement hypotaxique. Le « fait à expliquer » (*explanandum*) est dominant et contribue ainsi à l'assertion ou résumé. L'expliquant est subordonné. L'axiome suivant capte ce fonctionnement :

$$\forall (e_1, e_2) \text{ Explication}(e_1, e_2 \rightarrow \text{Relation de cohérence}(e_1, e_2, e_1))$$

Comme exemple de relation d'explication, considérons une variante de l'exemple classique de [WIN 72] :

Les policiers ont interdit aux jeunes de manifester.

Ils redoutaient des actes de violence.

Interpréter le texte consiste à prouver par abduction l'expression suivante :

Segment (« Les policiers... violence », e)

Cela implique de prouver que chaque phrase est un segment, en prouvant qu'elles sont des phrases, et qu'il y a une relation de cohérence entre elles. Pour prouver que ce sont des phrases, nous recourons à une version étendue de la grammaire de phrases (paragraphe 6.3.2). Cela nécessiterait de prouver par abduction la forme logique des phrases.

Une façon de prouver qu'il y a une relation de cohérence entre ces phrases est de prouver qu'il y a une relation d'explication entre elles, et une façon de le démontrer est de prouver qu'il y a une relation de cause entre les deux affirmations.

Après avoir fait ce chaînage arrière, nous devons donc prouver l'expression :

$$\exists(e_1, p, d, w, e_2, y, v, z) \text{ interdire}'(e_1, p, d) \ \& \ \text{manifester}'(d, w) \ \& \ \text{cause}(e_2, e_1) \ \& \ \text{redouter}'(e_2, y, v) \ \& \ \text{actes de violence}'(v, z)$$

A savoir : il y a un événement e_1 , l'agent est la police, pour un événement « manifestation » d , pour l'agent *les jeunes*. Il y a un événement à craindre e_2 craint par quelqu'un y , et violent : v , produit par quelqu'un w . L'événement à craindre e_2 cause l'interdiction de l'événement e_1 . Cette expression est simplement la forme logique de chacune des deux phrases, plus la relation supposée de causalité qui les relie.

Supposons, ce qui est assez plausible, que nous ayons les axiomes suivants :

$$\forall(e_2, y, v) \text{ redouter}'(e_2, y, v) \rightarrow \exists(d_2) \text{ non_voulu}(d_2, y, v) \ \& \ \text{cause}(e_2, d_2)$$

Soit : si e_2 est la crainte par y de v , cela causera l'état d_2 dans lequel y ne souhaite pas v .

$$\forall(d, w) \text{ manifeste}'(d, w) \rightarrow \exists(v, z) \text{ cause}(d, v) \ \& \ \text{violent}'(v, z)$$

c'est-à-dire manifester est cause de violence.

$$\forall(d, v, d_2, y) \text{ cause}(d, v) \ \& \ \text{non_voulu}(d_2, y, v) \rightarrow \exists(d_1) \text{ non_voulu}'(d_1, y, d) \ \& \ \text{cause}(d_2, d_1)$$

C'est-à-dire : si quelqu'un ne veut pas v , et que v soit causé par d , alors ce sera la cause que cette personne ne voudra pas d non plus. Si on ne veut pas l'effet, on ne veut pas la cause.

$$\forall(d_1, p, d) \text{ non_voulu}(d, p, d) \ \& \ \text{autorité}(p) \rightarrow \exists(e_1) \text{ interdit}'(e_1, p, d) \ \& \ \text{cause}(d_1, e_1)$$

Si ceux qui ont autorité n'aiment pas quelque chose, cela les conduit à l'interdire.

$$\forall(e_1, e_2, e_3) \text{ cause}(e_1, e_2) \ \& \ \text{cause}(e_2, e_3) \rightarrow \text{cause}(e_1, e_3)$$

Cause est transitif.

$$\forall(p) \text{ les policiers}(p) \rightarrow \text{autorité}(p)$$

La police détient l'autorité.

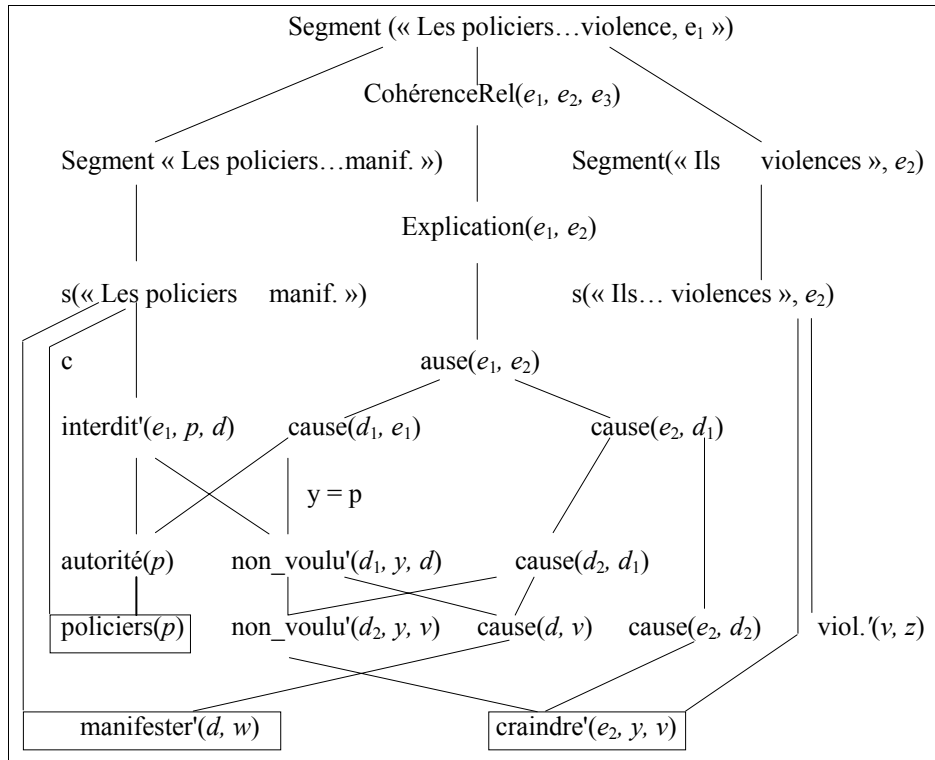


Figure 6.3. Interprétation de : « Les policiers ont interdit aux jeunes de manifester. Ils redoutaient des violences »

Au moyen de ces axiomes, nous pouvons prouver toutes les formes logiques ci-dessus, sauf les propositions *police* (p), *manifeste* (d, w), et *craint*'(f, y, v), que nous tenons pour vraies. Cela est illustré par la figure 6.3. Notez qu'au cours de la preuve, nous unifions y avec p , donnant au pronom *ils* l'antécédent *les policiers*, ce qui revient à choisir une solution pour le cas problématique de résolution anaphorique qui motivait l'exemple original de Winograd.

On peut imaginer un grand nombre de variations sur cet exemple. Si nous n'avions pas inclus l'axiome selon lequel les manifestations sont causes de violence, nous aurions dû supposer la violence et la relation causale entre manifestation et violence. En outre, d'autres relations de cohérence auraient pu être imaginées ici en construisant le contexte de manière appropriée. Nous pourrions avoir ensuite la phrase « Mais puisqu'ils n'avaient jamais manifesté auparavant ils ne savaient pas que des violences pouvaient en résulter ». Dans ce cas, la seconde phrase serait subordonnée à la troisième, obligeant à résoudre le *ils* comme renvoyant aux jeunes. Chaque exemple, bien sûr, devrait être analysé pour lui-même, et changer l'exemple changerait l'analyse. Dans l'adaptation de l'exemple original de Vinograd :

La police a interdit aux jeunes de manifester parce qu'ils redoutaient des violences,

la causalité était explicite, éliminant de ce fait les relations de cohérence comme une source d'ambiguïté. L'expression littérale *cause*(e_1, e_2) ferait partie de la forme logique.

6.5.3 Les relations de similarité

Une autre classe de relations de cohérence est basée sur la similarité. Ces relations, en un sens, développent le discours « sur place » au lieu de le faire progresser, ou de compléter son arrière-plan. Elles peuvent être classifiées en termes de passage entre assertions générales et spécifiques et d'interaction entre ces passages et la négation comme illustré figure 6.4.

	Spécifique à spécifique	Spécifique à général	Général à spécifique
Positive	Parallèle	Généralisation	Exemplification
Négative	Contraste	—	—

Figure 6.4. Relations de cohérence fondées sur la similarité

J'ai laissé deux cases vides dans la ligne *Négative* car de telles relations consisteraient une contradiction. Elles pourraient être remplies par une relation

d'« exception ». On donne une vérité générale, puis une exception à cette règle, ou *vice versa*. Mais j'ai choisi, de manière assez arbitraire, de considérer ces exemples comme des cas de Contraste.

Il y a deux cas limites intéressants. La relation d'Elaboration est un cas limite de la relation de Parallélisme; la relation d'Attente déçue est un cas-limite de Contraste.

Dans le fragment de conversation analysé dans la section 6.7, seules les relations de Parallélisme, d'Elaboration et de Contraste apparaissent ; nous ne considérons donc que celles-ci pour l'instant.

6.5.4. La relation de Parallélisme

La définition de la relation de Parallélisme implique la reconnaissance d'une similarité entre des entités. Les segments de texte doivent affirmer la même propriété d'entités similaires. Plus précisément, la relation Parallélisme est satisfaite quand, de l'assertion du premier segment S_1 , nous pouvons inférer $p(x_1, x_2, \dots)$ et de l'assertion du second segment S_2 , nous pouvons inférer $p(y_1, y_2, \dots)$, x_i et y_j étant similaires pour tous les i . Deux entités sont similaires si elles partagent quelque propriété (suffisamment spécifique). Les déterminations de similarité sont sujettes au même flou et aux mêmes évaluations de « qualité » que les relations de cohérence en général.

Une première approximation d'axiome captant cette caractérisation est la suivante :

$$\forall(p, q, x, y, e_1, e_2, e_3, e_4) p'(e_3, x) \& p'(e_4, y) \& \text{gen}(e_1, e_3) \& \text{gen}(e_2, e_4) \& q(x) \& q(y) \rightarrow \text{Parallel}(e_1, e_2, e)$$

Soit : il y a une relation de Parallélisme entre e_1 et e_2 s'il a un p tel que e_3 est le fait que p' soit vrai d'un x , où e_1 génère e_3 , et e_4 est le fait que p' soit vrai d'un y , où e_2 génère e_4 , et il y a un q vrai de x et y .

La relation de parallèle est une relation de cohérence :

$$\forall(e_1, e_2) \text{Parallèle}(e_1, e_2, e) \rightarrow \text{Relation de cohérence}(e_1, e_2, e)$$

La relation est parataxique. Je n'ai pas spécifié dans les axiomes ce que serait l'assertion d'un segment composite, mais il s'agirait de la généralisation dont l'assertion de chaque segment est une instance. La généralisation peut subsumer seulement les x et les y , ou toutes les entités pour lesquelles q est vrai.

Un exemple simple de la relation de Parallélisme est cet énoncé extrait d'une description d'algorithme :

Mettre la pile A à zéro, et mettre le pointeur de la variable P à T

De chacune des propositions, on peut inférer (de manière triviale) qu'une structure de données est associée à une valeur. Le prédicat p est par conséquent *mettre* ; la pile A et le pointeur de la variable P sont similaires dans la mesure où ce sont des structures de données, la mise à zéro de la pile, et l'association de P à T sont toutes deux des conditions initiales.

L'exemple qui suit est un peu moins simple. Il est emprunté à un problème extrait d'un manuel de physique.

L'échelle pèse 50 kg et son centre de gravité est à 6 mètres de son pied. Un homme de 75 kg est à 3 mètres du sommet.

Etant donné la nature de la tâche, le lecteur doit tirer des inférences de cette phrase à propos des forces pertinentes. Nous pouvons représenter ces inférences comme suit :

$force(50 \text{ kg}, E, \text{Bas}, x_1) \ \& \ distance(P, x_1, 6 \text{ m}) \ \& \ pied(P, E)$

$force(75 \text{ kg}, x, \text{Bas}, x_2) \ \& \ distance(H, x_2, 3 \text{ m}) \ \& \ haut(H, y)$

Ici, le prédicat p est *force*, les premiers arguments sont similaires (ce sont des poids), les second et troisième sont identiques (une fois x identifié avec E), donc similaires, et les quatrième sont similaires (ce sont des points sur l'échelle situés à une certaine distance d'une de ses extrémités (étant admis que y est E)). L'assertion du segment complexe est quelque chose comme : « Il y a des forces descendantes qui agissent sur l'échelle à certaines distances de ces extrémités », une information qui à elle seule ne permettrait pas de résoudre le problème.

L'exemple suivant est extrait d'un livre médical sur l'hépatite :

(§) Le sang est sang doute le milieu qui contient la plus forte concentration du virus de l'hépatite B après le foie.

Le liquide séminal, les sécrétions vaginales et le sang menstruel contiennent cet agent et sont contagieux.

La salive a une concentration plus faible que le sang, l'antigène superficiel de l'hépatite B ne peut d'ailleurs être détecté que pour la moitié des sujets infectés.

L'urine contient de faibles concentrations quel que soit le moment considéré.

Le prédicat p est *contient* ; le diagramme de la figure 6.5 indique les arguments similaires correspondants et les propriétés partagées (les titres de colonne) qui fondent leur similarité.

Notez également que les phrases sont données dans l'ordre décroissant de concentration ; Il est très fréquent que des genres particuliers, ou des « microgenres » soient caractérisés par des contraintes supplémentaires ajoutées à ces structures de cohérence universelles.

MILIEU	CONTIENT	CONCENTRATION	AGENT
sang	contient	la plus forte	VHB
Liquide séminal sécrétions vaginales sang menstruel	contient	plus faibles	agent
salive	a	plus faible	antigène superficiel de VHB
salive des sujets infectés	dans	détectable pour la moitié des cas	
urine	contient	faibles	

Figure 6.5. La relation de Parallélisme dans l'exemple (6)

6.5.5. Elaboration

L'élaboration est un cas extrême de la relation de Parallélisme, dans laquelle les entités ne sont pas seulement similaires, mais encore identiques. Cela revient à dire que la même proposition peut être inférée de l'assertion de chacun des segments. Les deux segments disent la même chose à un certain niveau. Dans notre notation, cela peut être capté par la relation *gen* :

$$\forall(e_1, e_2, e) \text{ Elaboration } (e_1, e_2, e) \rightarrow \text{relation de Cohérence } (e_1, e_2, e)$$

$$\forall(e_1, e_2, e) \text{ gen } (e_1, e) \ \& \ \text{gen } (e_2, e) \rightarrow \text{Elaboration } (e_1, e_2, e)$$

Soit : s'il existe une éventualité e qui est « générée » par e_1 d'une part et e_2 de l'autre, alors il y a une relation d'Elaboration entre e_1 et e_2 , et l'assertion du segment composé sera e .

Fréquemment le second segment ajoute une information cruciale, et c'est ce que j'ai appelé la relation d'Elaboration, mais cela n'est pas spécifié dans la définition

puisque'il est nécessaire d'inclure de pures répétitions sous ce chef. Un exemple simple de la relation d'Elaboration est le cas suivant :

Descendre la première rue

Suivez cette rue quelques mètres jusqu'à la rue A.

Remarquons qu'il est important de voir ici une relation d'Elaboration plutôt que deux instructions se succédant dans le temps.

De la première phrase, nous inférons :

aller(Agent : vous, But : x , Parcours : Première Rue, Mesure : y)

pour x et y . De la seconde phrase, nous inférons : *aller* (Agent : vous, But : Rue A, Parcours : Première Rue, Mesure : quelques mètres).

Si nous assertons que x est la rue A et que y est quelques mètres, alors les deux phrases sont identiques et représentent la proposition P de la définition. Ceci montre de façon relativement simple que nous pouvons ignorer les détails de la dérivation abductive de la relation (en ignorant les « quelques mètres ») ; pour interpréter le discours, nous devons prouver abductivement l'expression :

Segment (« Descendez...Rue A. », e)

Pour prouver que le texte est un segment, nous devons prouver que chaque phrase est un segment, en prouvant que c'est une phrase. Ceci nous ramène à une version étendue de la grammaire de phrase (section 6.3.2), ce qui exige que nous prouvions la forme logique de chaque phrase. Ainsi nous devons prouver (en simplifiant un peu) :

$\exists(g, u, x, y, f, f_1) \text{ descendre}'(g, u, x, y) \ \& \ \text{dans}(g, 1^{\text{ère}} \text{ Rue}) \ \& \ \text{relCohérence}(g, f, f_1) \ \& \ \text{suivre}'(f, u, 1^{\text{ère}} \text{ Rue}, \text{ Rue A})$

Soit : il y a une éventualité d'aller pour un agent u de x vers y , et le trajet est « en descendant » la 1^{ère} Rue. Il y a aussi une éventualité de suivre f pour u , la 1^{ère} Rue en direction de la Rue A. Enfin, il y a une relation de cohérence entre « aller » g et « suivre » f , avec une assertion complexe f_1 .

Supposons que nous ayons les axiomes suivants dans notre base de connaissances :

$\forall(f) \text{ gen}(f, f)$. La relation *gen* est réflexive.

$\forall(g, u, x, y, z) \text{ descendre}'(g, u, x, y) \ \& \ \text{le long de}(g, z) \rightarrow \exists(f) \text{ suivre}'(f, u, z, y) \ \& \ \text{gen}(g, f)$

Soit : si g est l'éventualité de descendre faite par u de x vers y , et que cela se fasse le long de z , alors g génère une éventualité de suivre f par u en parcourant z vers y .

$\forall(g, z) \text{ « en descendant »}(g, z) \rightarrow \text{« le long de »}(g, z)$

Donc une relation *descendre* est une forme de relation *parcourir*.

Si nous admettons $aller'(g, u, x, y)$ et $descendre(g, SR)$, alors la preuve logique du texte est directe. Elle est illustrée dans la figure 6.6.

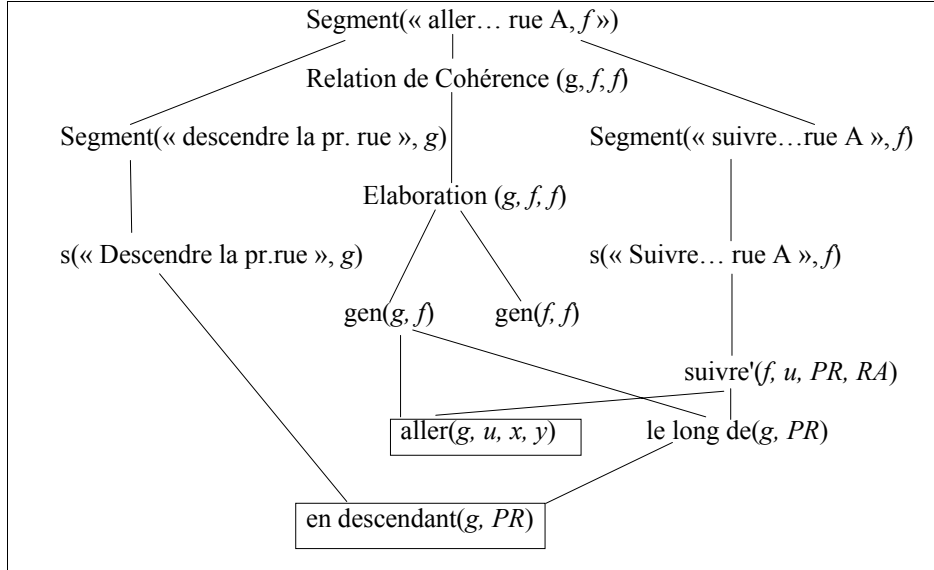


Figure 6.6. Interprétation de « Descendre la première rue.
Suivre la première rue jusqu'à la rue A »

6.5.6. Contraste

Il y a au moins deux cas de Contraste. Le premier est le suivante : nous pouvons inférer $p(x)$ de l'assertion d'un segment S_1 , et nous pouvons inférer $\neg p(y)$ de l'assertion du segment suivant S_2 , x et y étant des entités similaires.

L'axiome qui suit est une première approximation de cette caractérisation.

$$(7) \quad \forall(p, q, x, y, e_1, e_2, e_3, e_4) p'(e_1, x) \& \text{non}'(e_3, e_4) \& p'(e_4, y) \& \text{gen}(e_3, e_2) \& q(x) \& q(y) \rightarrow \text{Contraste}(e_1, e_2)$$

Soit : il y a une relation de contraste entre e_1 et e_2 s'il y a un p tel que e_1 est le fait que p est vrai d'un x , et s'il y a un e_3 qui engendre e_2 qui soit la négation d'un e_4 tel que p' soit vrai d'un y , et s'il y a un q vrai de x et de y .

La relation de contraste est une relation de cohérence :

$$\forall(e_1, e_2, e) \text{ Contraste}(e_1, e_2) \rightarrow \text{Relation de Cohérence}(e_1, e_2, e_2)$$

La relation est en général hypotaxique. En principe, le second segment est le segment dominant, mais ce n'est pas essentiel ici.

Le second cas de la relation de Contraste vaut quand nous pouvons inférer une proposition $p(x)$ de l'assertion d'un segment S_1 et que nous pouvons inférer $p(y)$ de l'assertion du segment suivant S_2 , s'il y a quelque propriété q telle que $q(x) \& \neg q(y)$. Le second segment est en général dominant.

Le premier cas est illustré par un exemple détaillé impliquant l'interprétation d'une métaphore.

John est un éléphant.

Pris isolément, cet énoncé peut signifier différentes choses : que John est pesant, qu'il a la peau dure, ou une bonne mémoire, etc. Dans des contextes particuliers, l'énoncé prendra des significations spécifiques. Dans le contexte suivant :

(8) Marie est gracieuse. John est un éléphant.

L'interprétation la plus raisonnable est que John est maladroit⁶.

Pour interpréter ce texte abductivement, nous utilisons d'abord un axiome liant la maladresse aux éléphants :

$$(9) \forall(y) \text{maladroit}(y) \& \text{etc}_2(y) \rightarrow \text{éléphant}(y)$$

Cela peut se lire ainsi : si un individu est maladroit et que certaines autres propriétés non spécifiées sont vraies de lui, alors cet individu est un éléphant, ou, en d'autres termes, être un éléphant est une des manières d'être maladroit.

Nous devons maintenant introduire une complication supplémentaire, car nous allons devoir nous référer explicitement aux propriétés d'être éléphant, gracieux, maladroit, etc. L'axiome (9) doit être réécrit ainsi :

$$(10) \forall(e_3, y) \text{maladroit}(e_3, y) \& \text{etc}_2(e_3, y) \rightarrow \exists(e_2) \text{éléphant}'(e_2, y) \& \text{gen}(e_3, e_2)$$

Soit : si e_3 signifie que y est maladroit, et que quelques autres choses non spécifiées soient vraies de e_3 et de y , alors, il existe un fait e_2 selon lequel y est un

6. Les éléphants, naturellement, ne sont pas maladroits ; mais selon nos stéréotypes, ils le sont. Cette propriété est donc dans notre base de connaissances et par conséquent accessible pour l'interprétation du discours. [SEA 79] l'a montré pour les gorilles réputés « fiers, méchants, violents, etc. ».

éléphant. En outre, il y a une relation très étroite entre e_3 et e_2 : y est un éléphant du fait qu'il est maladroit et que les autres choses sont vraies.

Ensuite, nous avons besoin d'un axiome pour relier la maladresse et la grâce.

$$(11) \forall(e_3, e_4, y) \text{ non}'(e_3, e_4) \& \text{gracieux}(e_4, y) \rightarrow \text{maladroit}'(e_3, y)$$

Soit : si e_3 est la condition selon laquelle e_4 n'est pas vrai, où e_4 est la condition y est gracieux, alors e_3 est la condition selon laquelle y est maladroit.

Supposons que nous sachions aussi que Marie et Jean sont des personnes :

personne (M), personne (J)

Maintenant, nous pouvons interpréter le texte (8). Sa forme logique est :

$$\exists(e_1, e_2, e) \text{ gracieux}'(e_1, M) \& \text{éléphant}'(e_2, J) \& \text{Relation de Cohérence}(e_1, e_2, e)$$

Nous pouvons inférer en chaînage arrière de « éléphant » à « maladroit », assumer $\text{etc}_2(e_3, J)$. Nous pouvons également admettre $\text{gracieux}(e_1, M)$. Nous avons alors une preuve de *Contraste* (e_1, e_2) en utilisant l'axiome (7), avec p instancié par gracieux , et q instancié par personne . Cela établit *Relation de Cohérence* (e_1, e_2, e).

Comme toujours, d'autres relations de cohérence sont théoriquement possibles, dans cet exemple. La phrase suivante pourrait être « Marie pourrait danser sur son dos », auquel cas la seconde phrase ne serait pas en contraste avec la première, mais en arrière-plan pour la troisième, et John serait plutôt un vrai éléphant qu'un éléphant métaphorique (figure 6.7).

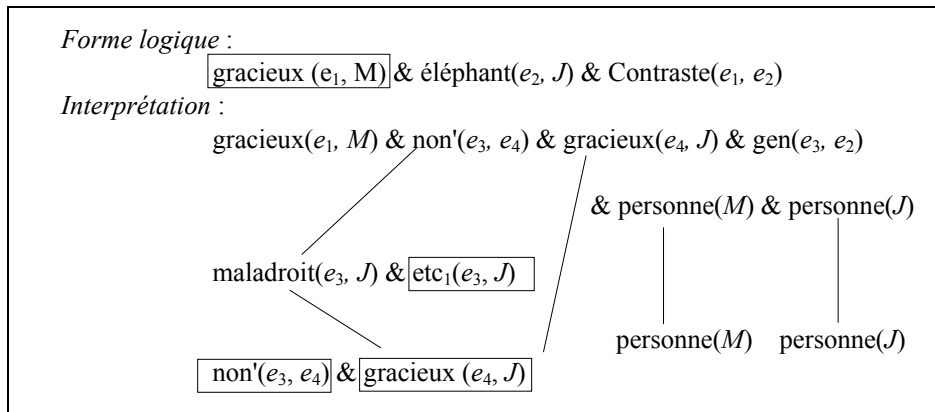


Figure 6.7. Interprétation de l'exemple de contraste

6.6. Une méthode pour analyser le discours

Ce traitement de la structure du discours suggère une méthode d'analyse du discours. La méthode consiste en quatre étapes, classées dans un ordre croissant de difficulté. J'illustrerai cela sur le texte suivant :

- (12.a) J'aimerais maintenant examiner l'hypothèse dite de l'« innéisme »
- (12.b) pour en identifier les éléments discutables, ou qui devraient être discutés
- (12.c) et aborder quelques-uns des problèmes qui surgissent lorsqu'on tente de trancher ce débat.
- (13) Nous pourrions essayer de voir ensuite ce qu'on peut dire de la nature et l'exercice de la compétence linguistique, ainsi que de quelques questions connexes.

Chomsky (Réflexions sur le langage, p. 23)⁷

La méthode est la suivante :

1. On identifie la ou les coupures majeures du texte, et l'on coupe à cet endroit. En substance, on choisit la manière la plus naturelle de diviser le texte en deux ou trois segments. Cela peut être fait sur une base totalement intuitive par toute personne ayant compris le texte, et parmi ceux-ci, on trouvera un large accord. Ce processus est ensuite répété pour chacun des segments, en opérant la segmentation la plus naturelle. On continue le processus jusqu'au niveau des propositions simples. Cela produit une structure en arbre pour le texte entier.

Dans le texte (12), par exemple, la coupure majeure intervient entre les phrases (12) et (13). Dans la phrase (12) il y a une coupure entre la première phrase et les deux suivantes, et naturellement, une coupure finale entre la seconde et la troisième proposition de la première phrase. Cela produit l'arbre de la figure 6.8.

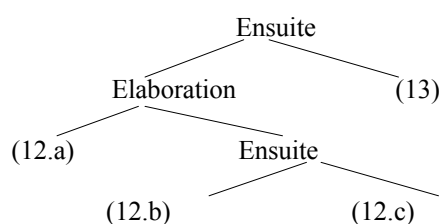


Figure 6.8. Structure des phrases (12)-(13)

7. Traduction et références de la traduction française (Editions Maspéro, 1977).

2. On légende les nœuds non terminaux de l'arbre avec des relations de cohérence. Procédant de bas en haut, on donne une esquisse de ce qui est asserté par chaque segment complexe. Ainsi, dans l'exemple de Chomsky, nous légendons le nœud qui relie (12.b) et (12.c) au moyen de la relation « Ensuite » (ou *Occasion*). Nous légendons le nœud qui lie le segment résultant et (12.a) au moyen de la relation *Elaboration*. Finalement, nous étiquetons le nœud reliant (12) et (13) avec la relation *Ensuite*. Ce processus produit l'étiquetage des nœuds représentés dans la figure 6.8.

A ce stade, la méthode devient dépendante des théories particulières, car on doit savoir de quelles relations on dispose, et avoir au moins une caractérisation grossière de chacune d'elles. On peut s'aider ici en déterminant quelles conjonctions ou adverbes de phrase il serait justifié d'insérer. Si nous pouvons insérer « Ensuite » entre S_0 et S_1 , et que le sens du texte serait altéré par un changement d'ordre des segments, alors la relation *Occasion* est un excellent candidat. Si nous pouvons insérer « parce que », la relation *Explication* devient une très forte probabilité. « C'est-à-dire », suggère *Elaboration*, « de même » suggère *Parallélisme*, « par exemple » suggère *Exemplification*, et « mais » suggère *Contraste*, ou *Attente déçue*. On doit cependant insister sur le fait que ces tests sont informels. Ils ne sont pas des définitions des relations. Les conjonctions et adverbes de phrase imposent des contraintes sur le contenu propositionnel des propositions qu'ils lient ou modifient, et en de nombreux cas, ces contraintes sont presque les mêmes que celles qui sont imposées par quelque relation de cohérence. Dans le meilleur des cas, il y a un lien suffisant pour que la conjonction nous indique quelle est la relation de cohérence.

3. On cherche à préciser (plus ou moins) les connaissances ou croyances qui légitiment ces assignations de relations de cohérence aux nœuds. Chacune de ces relations de cohérence a été définie en fonction des inférences qui doivent être tirées de la base de connaissance du récepteur pour que la relation soit reconnue. Quand nous disons, par exemple, qu'une relation « ensuite » s'établit entre (12.b) et (12.c), nous devons spécifier le changement asserté en (12.b) (en l'occurrence, un changement dans le savoir partagé concernant le point de la controverse, à partir du mot « identifier ») qui est présupposé par l'événement décrit en (12.c) (l'effort fait pour trancher le débat). Nous avons donc besoin d'un savoir concernant le changement effectué par l'action d'identifier, et nous devons connaître les significations de « controverse » et « trancher » qui nous autorise à parler de trancher un débat.

La précision avec laquelle nous spécifions ce savoir est réellement fonction de nos objectifs. Nous pouvons nous satisfaire d'une formulation claire en français, ou nous pouvons exiger une formulation dans un langage logique, inclus dans une théorie formelle plus vaste des représentations du sens commun.

4. On valide les hypothèses de l'étape 3 quant aux savoirs sous-jacents au discours. Mike Agar et moi-même [AGA 82] avons discuté en détail comment on pourrait procéder. En bref, on considère le corpus plus vaste auquel appartient le texte, un corpus produit par le même locuteur ou par la même culture supposant la même réception. On tente de construire une base de connaissances ou un système de croyances partagées qui légitimerait les analyses de tous les textes du corpus. Si l'étape 1 est une affaire de minutes pour un texte d'un paragraphe, l'étape 2 l'affaire d'une heure ou deux et l'étape 3 une affaire de jours, l'étape 4 est une affaire de mois ou d'années.

Dans chacune de ces étapes, des difficultés peuvent surgir, mais ces difficultés d'analyse révèlent en général des aspects réellement problématiques du texte. Dans l'étape 1, nous pouvons trouver difficile de segmenter le texte en certains endroits, mais cela révèle probablement une réelle zone d'incohérence dans le texte lui-même. Nous pouvons trouver la segmentation aisée parce que les segments concernent des sujets différents, mais avoir du mal à trouver des relations de cohérence pour lier les segments. Quand cela arrive, c'est peut-être que nous avons trouvé deux textes consécutifs plutôt qu'un texte unique. Par moments, le savoir sous-jacent à un segment complexe n'est pas évident, mais cela nous conduit souvent à des hypothèses nouvelles très intéressantes quant au système de croyance des participants. Enfin, souvent, nous ne pouvons pas être certains que le savoir que nous avons supposé actif l'est réellement ; la considération de données supplémentaires nous oblige à réviser nos hypothèses.

Il y a un rait structural de quelques discours, particulièrement des conversations orales, qui doit être considéré dans la segmentation, les pivots de discours. Un pivot de discours est un segment S_2 tel que dans $S_1S_2S_3$, S_1 et S_2 sont reliés en vertu d'une partie du contenu de S_2 , et S_2 et S_3 sont reliés l'un à l'autre en vertu d'une autre partie du contenu de S_2 . Les effets de la structure globale du texte est qu'un arbre subsume S_1S_2 , et un autre S_2S_3 . Le segment (14.b) ci-dessous est un exemple de pivot de discours.

6.7. Analyse d'un fragment de conversation

Des textes de tous genres peuvent être éclairés par une analyse de cohérence. Pour l'exemple détaillé de ce chapitre, j'ai choisi une espèce particulièrement difficile, un fragment d'une conversation à trois dans une discussion devant aboutir à une décision (même s'il n'y a que deux locuteurs dans ce fragment). Les participants mettent sur pied un programme pour la visite d'un représentant commercial, en s'assurant que plusieurs contraintes seront satisfaites au mieux.

- (14.a) A : Donc, euh, si je passe d'abord, disons avec, euh, disons moi
 (14.b) Comme j'ai dit, il me faut à peu près une heure et quart
 (14.c) Je peux faire le, mon rapport sur le projet en cours, ah pour cette première heure.
 (14.d) Regardez, si nous faisons le compte de tout le temps qu'il nous faut,
 (14.e) Disons une heure pour Brian
 (14.f) A : Une heure et quinze minutes pour moi
 B : Donc ça fait
 (14.g) A : Trente cinq minutes,
 B: A peu près, exactement...
 (14.h) A : C'est à peu près exactement correct. Trois heures
 (14.i) B : Mais il faut considérer qu'ils sont régulièrement en retard.
 (14.j) B : Oui, donc on va être très serrés
 A : Oui, c'est sûr. Oui, euh,
 (14.k) A : Je crois que ce qu'on devrait faire si on démarre en retard pour le, bon, si je passe en premier, et si on est en retard, je récupérerai le retard sur mon temps d'exposé.

Les énoncés qui se chevauchent sont donnés sous le même numéro.

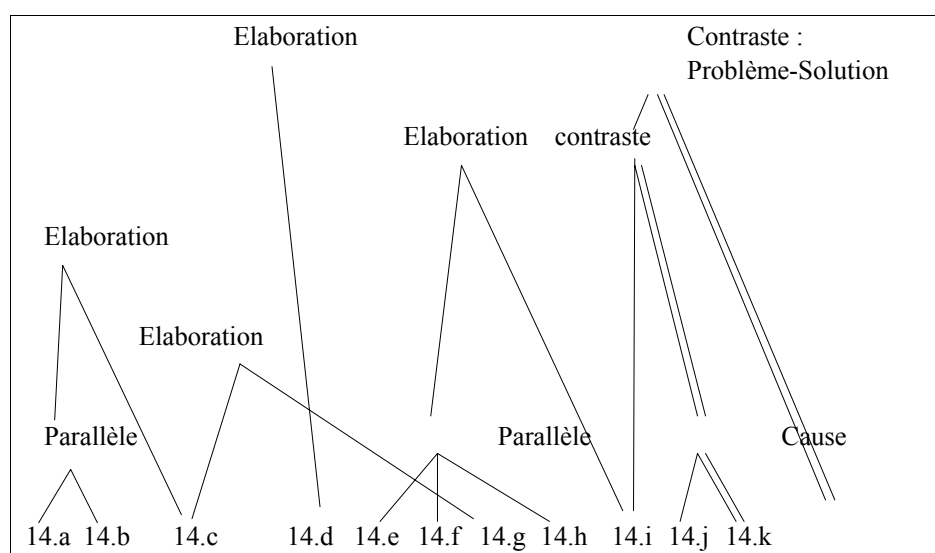


Figure 6.9. Structure de Cohérence d'un fragment de conversation

La structure de cohérence est représentée dans la figure 6.9.

Ce fragment se divise en trois parties. La première inclut les segments (14.a)-(14.c), dans lesquels A dit combien de temps il lui faut, et quand il devra l'avoir. Puis, en (14.d)-(14.h) il passe à une description du programme d'ensemble dans lequel s'insère sa présentation ; une sous-partie, le segment (14.f) répète le contenu du segment (14.c). Le segment (14.h) est un pivot de discours, et le segment (14.h)-(14.k) discute un problème et sa solution. Nous allons examiner successivement la structure détaillée de ces trois segments.

Pour construire un programme, les participants doivent affecter un certain temps à chaque activité et spécifier l'ordre dans lequel elles se succèdent. En (14.a) et (14.b) s'intéresse successivement à ces deux aspects de la conception du programme pour ce qui concerne sa propre présentation. En (14.a) il dit qu'il parlera en premier. En (14.b) il dit qu'il prendra une heure et quinze minutes. En (14.c) il développe cela en décrivant l'activité « pour cette première heure ».

En ce point, il revient en arrière sur un problème plus général, la durée totale requise par l'ensemble des trois participants, et dont son heure et quart sera une partie. Il commence par donner l'ébauche de ce qu'il fera, et totalise le temps total exigé. Cela est suivi par un exposé en parallèle des exigences de chaque participant : une heure pour B en (14.e), une heure et quart pour A en (14.f), et trente-cinq minutes pour C en (14.g). Cela est ensuite résumé, ou plus précisément totalisé en (14.h). Le segment (14.e)-(14.h) constitue donc l'élaboration promise en (14.d) et fournit le temps total exigé.

B a été un participant actif dans ce segment, mais seulement pour renforcer ou contribuer en quelque manière aux développements de A.

En (14.i), B interrompt (14.h) sur un mode différent. A le donne comme une récapitulation du temps exigé. Mais en (14.i) B met en question le fait que trois heures soit la durée totale qu'il leur faut prévoir. Les visiteurs sont parfois en retard. Les retards réduisent le temps disponible, comme le dit B en (14.j). Qu'ils disposent de moins de temps est en contraste avec l'affirmation (14.h) selon laquelle ils disposeraient exactement de trois heures.

Le segment (14h) est par conséquent un pivot de discours ; il fonctionne comme une élaboration du segment précédent, et vis-à-vis du segment suivant, il se comporte comme l'élément subordonné d'une relation de contraste.

Un problème est une situation qui n'est pas en accord avec nos objectifs. Une solution au problème est une modification de cette situation, en général au moyen

d'une action, qui soit en accord avec les objectifs. Par conséquent, l'exposé d'un problème et de sa solution est une variété de la structure *contraste*. Ici, le segment (14.h)-(14.j) décrit le problème, l'énoncé principal du problème se trouvant en (14.j). Le segment (14.k) décrit la solution. Si du temps se trouve perdu, A abrégera d'autant sa propre présentation.

6.8. Conclusion

La théorie de la cohérence locale du discours que j'ai esquissée dans ce chapitre est une partie d'une théorie plus vaste qui cherche à expliciter les liens entre l'interprétation d'un texte et le système de connaissances ou de croyances sous-jacente à ce texte. Les relations de cohérence qui structurent un texte sont une partie de ce qu'est une interprétation ; elles sont reconnues sur la base d'inférences, et par conséquent spécifient un lien qui doit exister entre des interprétations des connaissances-croyances. La méthode indiquée dans ce chapitre peut être utilisée pour exploiter cette connexion de différentes manières.

Si, comme en ethnographie, notre intérêt réside dans le système de croyances, ou la culture, partagé par les participants, la méthode agit comme une fonction de « forçage ». Elle ne nous dit pas ce que sont les croyances sous-jacentes, mais elle nous force à faire l'hypothèse de croyances que nous pourrions sinon laisser échapper, et cela impose des contraintes fortes sur ce que nous devons considérer comme une croyance.

Si notre intérêt réside pour l'essentiel dans l'interprétation d'un texte, comme en critique littéraire, la méthode nous donne une technique pour trouver la structure d'un texte, un aspect important de l'interprétation. En imposant des contraintes sur la structure idéale d'un texte, cette méthode peut diriger notre attention vers des zones problématiques du texte dans lesquelles les idéaux de cohérence proposés ici ne semblent pas atteints. Nous pourrions en fin de compte décider qu'il s'agit bien de cas dans lesquels ces idéaux ne sont pas atteints, mais dans de nombreuses circonstances nous découvrirons que nos tentatives pour satisfaire à ces idéaux nous conduisent à des très intéressantes réinterprétations du texte dans son ensemble.

6.9. Bibliographie

[AGA 82] AGAR M., HOBBS J., « Interpreting Discourse: Coherence and the analysis of Ethnographic Interviews », *Discourse Processes*, Vol. 5, n°1, p. 1-32, 1982.

- [DAV 67] DAVIDSON D., « The Logical Form of Action Sentences », dans N. Rescher (dir.), *The logic of Decision and Action*, p. 91-95, University of Pittsburgh Press, Pittsburgh, Pennsylvania, 1967.
- [DOW 77] DOWNING P., « On the Creation and Use of English Counpound Nouns », *Language*, vol. 53, n°4, p. 810-842, 1977.
- [HOB 78] HOBBS J. R., « Why is Discourse Coherent », SRI Technical Note 176, SRI International, Menlo Park, California, 1978, reproduit dans F. Neubauer (dir.) *Coherence in Natural-Language Texts*, p. 29-69, Helmut Buske Verlag, Hambourg, 1983.
- [HOB 85a] HOBBS J. R., « Ontological Promiscuity », actes, *23rd Annual Meeting of the Association for Computational Linguistics*, p. 61-69, Illinois, Chicago, 1985.
- [HOB 85b] HOBBS J. R., « On the Coherence and Structure of Discourse », *rapport CSLI 85-37*, Center for the Study of Language and Information, Stanford University, 1985.
- [HOB 95a] HOBBS J. R., « The Roles of Intention and Information in the Interpretation of Utterances », D. Scott et E. Hovy (dir.), actes, *NATO Workshop on Burning Issues in Discourse*, Maratea, Italie, 1995.
- [HOB 95b] HOBBS J. R., *The Syntax of English in an Abductive Framework*, manuscrit, 1995.
- [HOB 90] HOBBS J. R., KAMEYAMA, M., « Translation by Abduction », dans H. Karlegren (dir.), actes, *Thirteenth International Conference on Computational Linguistics*, vol. 3, p. 155-161, Finlande, Helsinki, 1990.
- [HOB 93] HOBBS J. R., STICKEL M., APPELT D., MARTIN, P., « Interpretation as Abduction », *Artificial Intelligence*, vol. 63, n°1-2, p. 69-142, 1993.
- [KOW 80] KOWALSKI R., *Logic for Problem Solving*, North Holland, New York, 1980.
- [LEV 78] LEVI J., *The Syntax and Semantics of Complex Nominals*, Academic Press, New York, 1978.
- [MCC 87] MCCARTHY J., « Circumscription: a Form of Nonmonotonic Reasoning » dans M. Ginsberg (dir.) *Readings in Nonmonotonic Reasoning*, p. 145-152, Morgan Kaufmann Publishers, Inc., California, Los Altos, 1987.
- [PER 83] PEREIRA F. C.N., WARREN D., « Parsing as Deduction », *Proceedings of the 21st Annual Meeting*, Association for Computational Linguistics, p. 137-144, Mass, Cambridge, 1983.
- [POL 88] POLANYI L., « A Formal Model of Discourse Structure », *Journal of Pragmatics*, vol. 12, p. 601-638, 1988.
- [POL 94] POLLARD C., SAG I.A., *Head-Driven Phrase Structure Grammar*, University of Chicago Press et CSLI Publications, 1994.
- [SEA 79] SEARLE J., « Metaphor », dans A. Orthony (dir.) *Metaphor and Thought*, p. 92-123, England, Cambridge University Press, Cambridge, 1979.
- [WIN 72] WINOGRAD T., *Understanding Natural Language*, Academic Press, New York, 1972.