

Approches pour la recomposition de fragments sémantiques

Sommaire

11.1 Introduction	132
11.2 Composition d'arbres	132
11.3 Stratégies de décision	134
11.3.1 Méthode de connexion forte	134
11.3.2 Méthode de connexion par classifieur	135
11.4 Conclusion	140

Résumé

Ce chapitre présente les algorithmes de recomposition d'arbres utilisés pour finaliser le processus de compréhension à partir des fragments de structures proposés par le décodage en DBN. Deux approches sont proposées et évaluées dans ce chapitre : la méthode de connexion forte et la méthode de connexion par classifieur.

11.1 Introduction

A l'issue du décodage sémantique réalisé par les modèles DBN présentés au chapitre 8, le système de compréhension dispose de fragments structures composés de frames et de FE. Ces fragments sont des entités sémantiques composées qui représentent des situations élémentaires incluses dans le message à interpréter. Les relations sémantiques liant ces situations élémentaires ne sont pas toujours décodées par les modèles DBN. En effet, ces relations sont souvent portées par des dépendances “longue-distance”, non modélisables par les DBN.

La capture de ces relations nécessite donc l'examen global du message. Dans le contexte de dialogue du corpus MEDIA, le message de l'utilisateur est borné par les interventions de l'opérateur simulant le serveur vocal. Les relations entre les fragments sémantiques produits par les DBN sont donc recherchées en considérant un tour de parole complet de l'utilisateur.

Deux algorithmes de recomposition des arbres sémantiques dont les fragments (i.e. branches, ou sous-arbres) sont produits par les DBN ont été évalués. Ils sont détaillés dans la section 11.2. La première approche est déterministe, la seconde s'appuie sur les décisions de classifieurs SVM.

11.2 Composition d'arbres

Les arbres étiquetés ont été définis en 10.2. Leur utilisation pour représenter les connaissances sémantiques et leur manipulation dans ce travail nécessitent la définition des opérations pouvant être réalisées sur ces arbres.

Comme décrit en 8.5, l'apprentissage des paramètres des modèles DBN nécessite la projection de l'arbre de frames et FE associé à la phrase complète pour relier sous-branches sémantiques et concepts de base. Lors de cette projection, des opérations de deux types sont réalisées. Les opérations de *séparation* rompent des liens entre frames et FE selon les concepts qui leurs sont associés. Les opérations de *duplication* des objets sémantiques (frames ou FE) sont nécessaires lorsque ces objets sont présents dans plusieurs sous-branches distinctes.

L'algorithme de recomposition est développé pour rassembler les fragments sémantiques produits par les DBN et rétablir l'arbre sémantique associé à la globalité du message. Il décide des opérations réciproques de celles effectuées lors de la projection, soit des opérations de *liaison* entre frames et FE et des opérations de *regroupement* entre frames ou FE.

Les liaisons potentielles inter-fragments s'appuient sur l'ontologie en frames et FE développée pour le domaine du corpus MEDIA : deux objets sémantiques ne peuvent être reliés que s'ils le sont dans l'ontologie. Ainsi, deux sous-branches sémantiques ne peuvent être connectées que si elles portent des nœuds dont les labels sont reliés dans l'ontologie. Les opérations de liaison consistent donc en l'ajout d'arêtes entre des

nœuds de sous-branches sémantiques distinctes pour les rassembler sous un arbre sémantique unique.

Les identifications potentielles concernent les objets sémantiques semblables présents au sein de plusieurs fragments associés à un même message. L'algorithme de recomposition considère ces objets et décide de la pertinence de leur présences multiples. Les opérations de regroupement ont ainsi pour rôle de supprimer les objets sémantiques redondants produits par les DBN. Lorsque deux nœuds de sous-branches sont identifiés, un seul est conservé dans l'arbre sémantique global et les nœuds fils du nœud supprimé sont reliés au nœud conservé.

L'exemple donné dans la figure 11.1 illustre ces processus. L'arbre sémantique associé au message "réserver un hôtel à Bourg-en-Bresse" est reproduit sur la gauche de la figure. Les branches de l'arbre décomposé sont reproduites à droite. Lors de la décomposition pour l'apprentissage des paramètres des modèles DBN, la frame `LODGING` est dupliquée et le FE `reserve_theme` est séparé d'une des deux instances de cette frame.

L'algorithme de recomposition doit être à même de recomposer l'arbre complet en disposant du message, des concepts associés et des branches générées par les DBN à partir de ces connaissances. Dans notre exemple, les deux frames `LODGING` doivent être regroupées. La liaison entre le FE `reserve_theme` et la frame `LODGING` résultante est naturellement réalisée dans ce cas. On remarquera que le message "réserver un séjour à Bourg-en-Bresse" aurait généré les branches `RESERVE-reserve_theme` et `LODGING-lodging_location-LOCATION-location_town`. L'algorithme de recomposition aurait eu dans ce cas à décider de la liaison entre le FE `reserve_theme` et la frame `LODGING`.

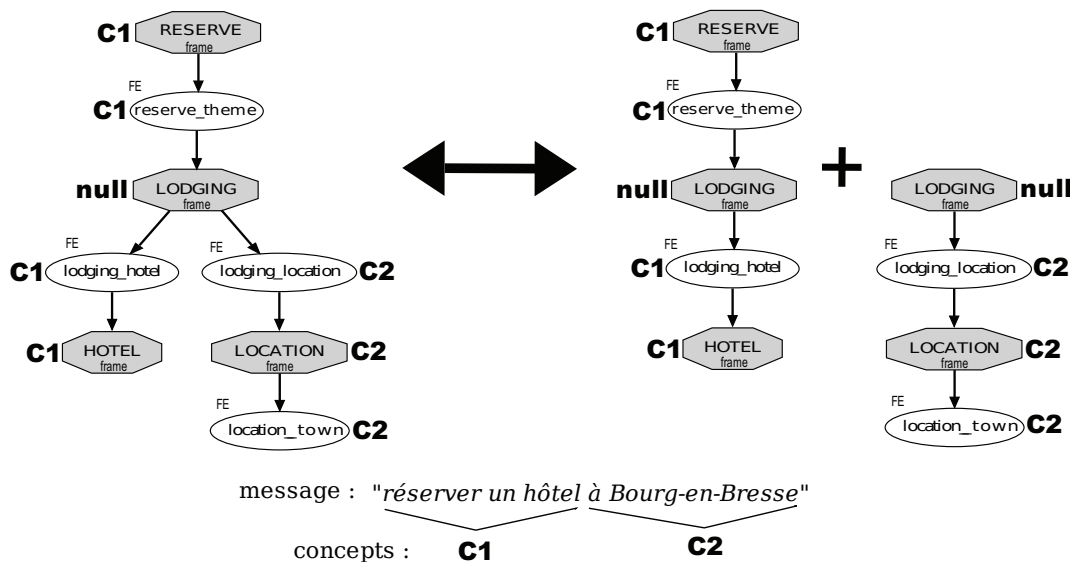


FIGURE 11.1 – Décomposition et recomposition de l'arbre sémantique associé au message "réserver un hôtel à Bourg-en-Bresse"

La pertinence de l'arbre sémantique recomposé est donc directement dépendante de

la pertinence des décisions de liaisons et de regroupement. Deux stratégies de décision sont évaluées dans ce travail. Elles sont présentées ci-après.

11.3 Stratégies de décision

11.3.1 Méthode de connexion forte

La première stratégie évaluée est une heuristique visant à obtenir pour chaque message une représentation sémantique compacte dans le cadre autorisé par l'ontologie. Dans cette méthode de **connexion forte**, toute liaison ou regroupement, possible selon l'ontologie, est réalisée.

Cette approche est a priori efficace pour les messages simples contenant des phrases courtes et peu ambiguës. En revanche, elle ne prend pas en compte les mots et les concepts associés au message. Elle n'est donc pas très adaptée aux messages complexes dont la représentation sémantique peut contenir de nombreux sous-structures non connectées. Le principe de cette heuristique est présenté dans la figure 11.2.

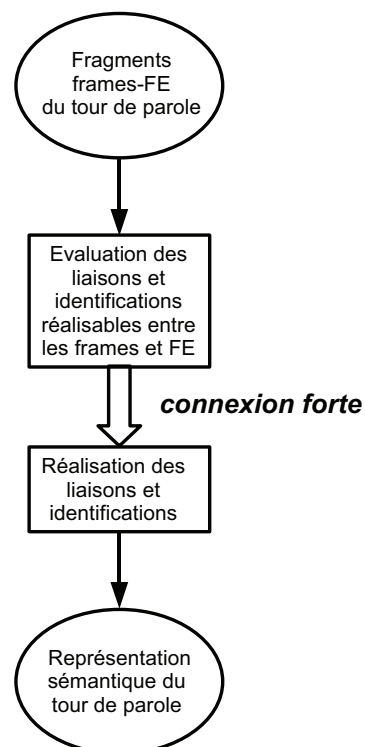


FIGURE 11.2 – Principe d'application de la méthode de **connexion forte**

11.3.2 Méthode de connexion par classifieur

La seconde stratégie évaluée s'appuie sur une méthode de connexion basée sur l'apprentissage de classifieurs SVM, que nous avons nommée **connexion SVM**. Le choix du type de classifieurs linéaires employés est dicté par plusieurs considérations : la quantité de données disponibles, la rapidité de réponse ou encore les performances obtenues sur des données comparables. En raison de leurs propriétés, décrites en 10.3, les classifieurs SVM s'adaptent parfaitement au contexte applicatif de ce travail.

Apprentissage

Les opérations de séparation et de duplication réalisées lors de la projection des arbres nécessaire à l'apprentissage des paramètres des DBN sont recensées. L'algorithme utilisé est donné ci-après (Algorithme 2).

t] Algorithme de projection d'arbre avec extraction des opérations réalisées

Entrée : $\{c_i\}$ séquence de concepts associés à la phrase, $\mathcal{T}$ arbre de frames et FE représentant le message
Sortie : $\mathcal{B}$ ensemble des branches

```

1:  $\mathcal{B} \leftarrow \emptyset$ 
2: pour tout  $c \in \{c_i\}$  faire
3:    $b_c \leftarrow \emptyset$ 
   Génération des branches principales
4:   pour tout  $l \in \text{feuilles}(\mathcal{T}, c)$  faire
5:     si  $l.\text{concept}() == c_i$  alors
6:        $b_c.\text{ajouter}(\text{extraire\_branche}(l))$ 
7:     fin si
8:   fin pour
   Contrôle des branches internes
9:   pour tout  $n \in \text{noeuds}(\mathcal{T}, c)$  faire
10:    si  $n \notin b_c$  ET  $n.\text{concept}() == c_i$  alors
11:       $b_c.\text{ajouter}(\text{extraire\_branche}(n))$ 
12:    fin si
13:   fin pour
14:    $\mathcal{B} \leftarrow b_c$ 
15: fin pour
16: retourner  $\mathcal{B}$ 

Fonction extraire_branche
Entrée :  $c_i$  concept,  $n$  noeud
Sortie : branche  $b \leftarrow \emptyset$ 
17: répéter
18:    $b. = n$ 
19:   si  $\mathcal{B}.\text{contient}(n)$  alors
20:     duplication/regroupement( $n$ )
21:   fin si
22:    $n \leftarrow n.\text{père}()$ 
23: jusqu'à  $!(n \text{ ET } (n.\text{concept} \in \{c_i, \text{null}\}))$ 
24: si  $n$  alors
25:   séparation/liaison( $b.\text{last}(), n$ )
26: fin si

```

27: retourner b

end

A chaque opération est associé l'ensemble des exemples du corpus d'entraînement contenant les objets sémantiques qu'elle fait intervenir. Ces messages sont répartis en deux classes selon qu'ils ont ou non *déclenché* l'opération.

On dispose de $\mathcal{T}$, ensemble des exemples d'apprentissage annotés en arbres sémantiques par le système à base de règles décrit au chapitre 6.

Soit $\mathcal{A}$ l'ensemble construit à partir de $\mathcal{T}$ tel que tout élément de $\mathcal{A}$ est composé des mots, des concepts et de l'arbre sémantique associés à un exemple de $\mathcal{T}$.

Soit $\mathcal{A}^p$ l'ensemble construit à partir de $\mathcal{A}$ tel que tout élément de $\mathcal{A}^p$ est composé :

- des mots,
- des concepts,
- des fragments sémantiques obtenus après projection de l'arbre sémantique,
- des opérations de projection (séparation, duplication) réalisées lors de la projection de l'arbre sémantique

d'un exemple de $\mathcal{A}$.

Soient $\mathcal{O}$ l'ensemble des opérations observées dans $\mathcal{A}^p$. Par souci de simplification, une opération de projection et sa réciproque de recombposition (ou *regroupement*) seront également notées $\mathcal{O}_i$, le contexte d'application levant toute ambiguïté.

Chaque opération de projection $\mathcal{O}_i \in \mathcal{O}$ met en jeu deux objets sémantiques f_{i1} et f_{i2} et on notera $\mathcal{O}_i = f_{i1} \mathcal{R} f_{i2}$.

Pour chaque paire $\{f_{i1}, f_{i2}\}$ associée à une opération de $\mathcal{O}$, on construit l'ensemble $\mathcal{A}_i^p$ des exemples de $\mathcal{A}^p$ contenant f_{i1} et f_{i2} .

Les exemples de $\mathcal{A}_i^p$ pour lesquels l'opération $\mathcal{O}_i$ s'est appliquée lors de la projection sont dits "positifs" pour $\mathcal{O}_i$.

Les exemples $\mathcal{A}_i^p$ contenant s_{i1} et s_{i2} pour lesquels $\mathcal{O}_i$ n'a pas été appliquée sont "négatifs" pour $\mathcal{O}_i$.

On dispose pour chaque opération $\mathcal{O}_i$ de la partition $\{\mathcal{A}_i^{p+}, \mathcal{A}_i^{p-}\}$ de $\mathcal{A}_i^p$ où $\mathcal{A}_i^{p+}$ et $\mathcal{A}_i^{p-}$ sont respectivement les sous-ensembles d'exemples positifs et négatifs de $\mathcal{A}_i^p$.

Pour permettre l'emploi de la méthode de classification SVM, telle que décrite en 10.3, il est nécessaire de plonger les données dans $\mathbb{R}^n$. Un exemple $\mathcal{E}$ est représenté dans $\mathbb{R}^n$ par un point E dont les coordonnées sont les index numériques :

- des mots et trigrammes de mots de l'exemple ;
- de la séquence de concepts associée à l'exemple ;
- des frames et FE présents dans les fragments sémantiques associés à cet exemple.

L'introduction des n-grammes de mots dans le point caractérisant un exemple permet de prendre en compte une information séquentielle.

Les paramètres d'un classifieur linéaire binaire à base de SVM sont appris sur les points représentant les exemples de chaque ensemble $\mathcal{A}_i^p$. A l'issue de cette procédure, on

dispose d'un classifieur S_i par opération $\mathcal{O}_i$ et on a $|\mathcal{O}| = |\mathcal{S}| = I$, avec $\mathcal{S}$ l'ensemble des classifieurs entraînés.

Le tableau 11.1 donne un exemple de l'élaboration des caractéristiques des cas positif et négatif pour l'opération de regroupement de deux occurrences de la frame `HOTEL`.

	Message POSITIF	Message NÉGATIF
<i>Caractéristiques du point représentant le message</i>		
<i>mots</i>	"Ibis Prado"	"Ibis ou Mercure"
<i>concepts</i>	hotel-marque nom-hotel	hotel-marque connectattr hotel-marque
<i>frames et FE</i>	HOTEL hotel_name hotel_marque	HOTEL hotel_marque
<i>Fragments sémantiques associés au message</i>		
<i>avant projection</i>	HOTEL-hotel_name-hotel_marque	HOTEL-hotel_marque
<i>après projection</i>	HOTEL-hotel_name HOTEL-hotel_marque	HOTEL-hotel_marque HOTEL-hotel_marque

TABLE 11.1 – Caractéristiques des messages positif et négatif dans le cas du regroupement de deux occurrences de la frame `HOTEL`.

Les deux exemples ont en commun de posséder deux frames `HOTEL` dans les fragments projetés qui leur sont associés. Avant projection, le fragment sémantique du cas positif possède une unique frame `HOTEL`. En effet, le message mentionne un unique hôtel, l'Ibis Prado. Après projection, la frame `HOTEL` est dupliquée. Le cas est donc positif pour l'opération de regroupement de deux frames `HOTEL`.

Les frames `HOTEL` du cas négatif sont présentes dans les fragments sémantiques *avant* projection. Elles sont distinctes et symbolisent les deux (marques d') hôtels différent(e)s mentionné(e)s par l'utilisateur, Ibis et Mercure. L'exemple appartient donc à l'ensemble des cas négatifs pour l'opération de regroupement de deux frames `HOTEL`.

Application aux exemples de l'ensemble de test

Pour chaque exemple $\mathcal{E}$ de l'ensemble de test, annoté en fragments sémantiques, on construit l'ensemble des opérations pouvant le concerner en fonction des paires d'objets sémantiques contenues dans ses fragments.

Soit $\mathcal{O}^{\mathcal{E}}$ cet ensemble, on a :

$$\mathcal{O}^{\mathcal{E}} = \{\mathcal{O}_{i \in I} = f_{i1} \mathcal{R} f_{i2} \text{ tq } f_{i1} \text{ et } f_{i2} \text{ appartiennent aux fragments sémantiques associés à } \mathcal{E}\}$$

et $\mathcal{O}^{\mathcal{E}} \subset \mathcal{O}$.

Pour toute opération $\mathcal{O}_i \in \mathcal{O}^{\mathcal{E}}$, le point E représentant l'exemple $\mathcal{E}$ est soumis au classifieur S_i . La réponse de S_i quant à la classe de E détermine la pertinence de la réalisation de $\mathcal{O}_i$ sur les objets sémantiques de l'exemple.

A l'issue de ce processus, la phase de composition sémantique est achevée par la réalisation sur les objets sémantiques de $\mathcal{E}$ de toutes les opérations jugées pertinentes par les $S_{i \in I}$.

Le principe général d'application de cette méthode est présenté dans la figure 11.3.

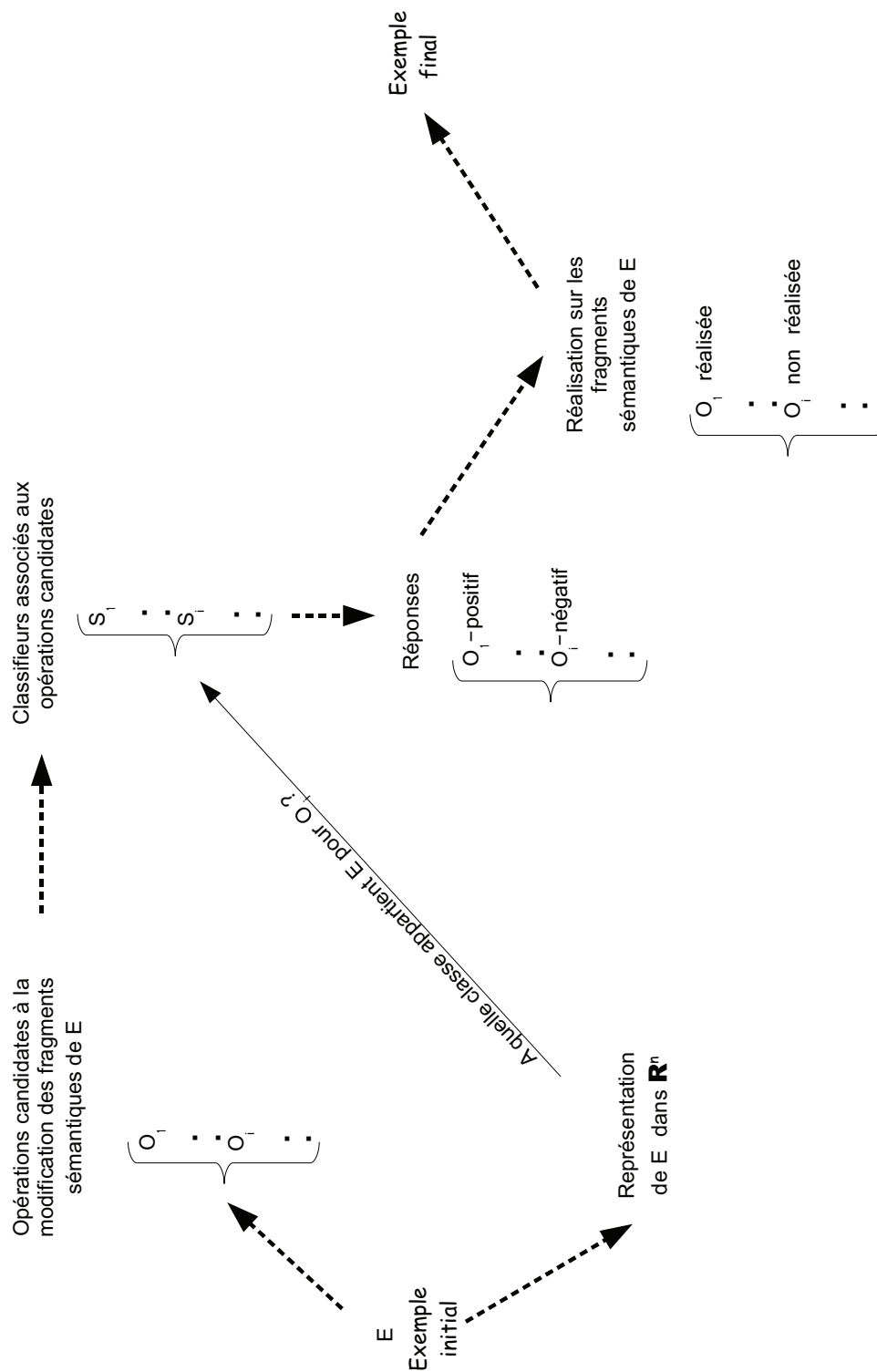


FIGURE 11.3 – Principe d'application de la méthode de connexion par classifieur SVM

11.4 Conclusion

Les méthodes de recomposition sémantiques que nous développons dans ce travail sont basées sur le formalisme des arbres sémantiques. Les arbres utilisés sont non orientés. Deux types d'opérations permettent de les recomposer : le regroupement de deux nœuds de même étiquette sémantique et la liaison de deux nœuds d'étiquettes reliées dans la base de connaissances.

Deux méthodes sont proposées pour décider de la réalisation des opérations de recomposition. La méthode de **connexion forte** est déterministe. Elle considère les objets des fragments sémantiques d'un message et réalise toutes les opérations de recomposition compatibles avec les contraintes relationnelles de la base de connaissances.

La méthode de **connexion par classifieur SVM** considère toutes les informations caractérisant un message (mots, concepts, fragments sémantiques). Un ensemble de classifieurs SVM dédiés lui permet de décider des opérations à réaliser sur les fragments sémantiques de chaque message.

Ces deux méthodes sont évaluées dans le chapitre [12](#) sur les fragments sémantiques produits par le modèle DBN compact.

Chapitre 12

Expériences et résultats

Sommaire

12.1 Introduction	142
12.2 Expériences	142
12.3 Résultats	143
12.4 Conclusion	146

Résumé

Ce chapitre décrit les expériences menées pour évaluer les algorithmes de composition des fragments sémantiques présentés dans le chapitre 11. Les deux algorithmes proposés sont appliqués à la composition des fragments sémantiques produits par le modèle DBN compact sur l'ensemble de test de MEDIA. Les résultats obtenus par chaque algorithme sont détaillés selon la nature manuelle ou automatique des données de test utilisées.

12.1 Introduction

Deux méthodes de composition des fragments sémantiques sont présentées dans le chapitre précédent 11 : la méthode de connexion forte et la méthode de connexion SVM. Ces deux méthodes sont évaluées sur les tours de parole utilisateur annotés en frames et FE de l'ensemble de test MEDIA.

Les messages d'entraînement annotés par le système à base de règles décrit en 6 permettent l'entraînement des classifieurs SVM utilisés par la méthode de connexion SVM.

Les expériences réalisées sont décrites dans la section 12.2 et les résultats obtenus sont donnés en 12.3.

12.2 Expériences

De même que lors de l'évaluation des systèmes DBN (Chapitre 9), les expériences sont menées sur l'ensemble des 3005 tours de parole de test MEDIA dans trois conditions différentes, fonctions de la nature des données utilisées :

- MAN : les tours de parole du locuteur sont manuellement transcrits et annotés en concepts ;
- SLU : les concepts de base sont décodés à partir des transcriptions manuelles des tours de parole locuteur, en utilisant le modèle SLU à base de DBN décrit dans (Lefèvre, 2006) ;
- ASR+SLU : les concepts sont décodés par le modèle de compréhension en utilisant la meilleure hypothèse (1-best) de séquence de mots générée par un système ASR conforme à (Barrault et al., 2008).

Les données SLU et ASR+SLU comportent des erreurs de transcription et d'annotation conceptuelle liées à l'imperfection des systèmes qui les produisent. Les taux d'erreurs observés sur le test sont rappelés dans le tableau 12.1.

Type de données	SLU	ASR + SLU
Taux d'erreurs mots (%)	0,0	27
Taux d'erreurs concepts (%)	10,6	24,3

TABLE 12.1 – Taux d'erreurs en mot et en concept observés sur les données SLU et ASR+SLU de l'ensemble de test MEDIA.

Pour chaque ensemble de données MAN, SLU et ASR+SLU, les fragments sémantiques sont générés par le modèle DBN compact du chapitre 8.

Les différents niveaux d'évaluation sont détaillés ci-dessous :

- **Frames** : les hypothèses de frames sont considérées comme correctes dès lors que les frames correspondantes sont présentes dans la référence.
- **FE** : les hypothèses de FE sont considérées comme correctes dès lors que les FE correspondants sont présents dans la référence ;

- **FE{Frames}** : seules les hypothèses de FE appartenant à des hypothèses de frames correctes sont examinées. L'ensemble de référence est restreint aux FE appartenant aux frames correspondantes dans la référence ;
- **Liens** : les hypothèses de liens sont considérées comme correctes dès lors que les liens correspondants sont présents dans la référence.
- **Liens{Frames}** : seules les hypothèses de liens reliant des hypothèses de frames et FE correctes sont examinées. L'ensemble de référence est restreint aux liens reliant les frames et FE correspondants dans la référence.

Toutes les expérimentations reportées dans ce document ont été réalisées en utilisant la librairie `libSVM` (Chang et Lin, 2001; EL-Manzalawy et Honavar, 2005) pour WEKA (Witten et Frank, 2005; Bouckaert et al., 2008).

12.3 Résultats

Le tableau 12.2 regroupe les résultats issus de l'application des méthodes de connexion forte et SVM sur les fragments sémantiques issus du système basé sur le modèle DBN compact présenté au chapitre 8.

Les résultats sont mesurés en termes de précision, rappel, F-mesure et leurs valeurs moyennes sur un tour de parole pour les trois ensembles de données (MAN, SLU et ASR+SLU)

La méthode de connexion SVM fait usage de 105 classifieurs appris sur le corpus d'entraînement et répartis comme suit :

- 44 classifieurs sont dédiés à l'identification de frames (18) ou de FE (26),
- 61 classifieurs sont dédiés à la liaison de frames et FE.

Les deux algorithmes proposés obtiennent des résultats similaires sur l'ensemble de test MEDIA. Ces résultats confirment l'aptitude de ces algorithmes à composer les fragments sémantiques pour former une représentation complète consistante du message de l'utilisateur.

La structure de la base de connaissance et le contexte relativement fermé des messages de test peuvent expliquer les performances quasi identiques des deux méthodes. En effet, les opérations de regroupement et de liaison des objets sémantiques contenus dans les fragments étant presque toujours justifiées, la méthode de connexion forte commet finalement peu d'erreurs.

Les résultats obtenus par la méthode de connexion SVM permettent d'envisager l'évaluation de cette méthode selon plusieurs axes. L'influence de l'augmentation du dimensionnement de l'espace de travail des classifieurs SVM sur les performances de la méthode pourra être évaluée grâce à la base connaissance LUNA, plus vaste que celle utilisée dans ce travail. Parallèlement, la sélection des tours de parole les plus complexes de l'ensemble de test est en cours. Les deux méthodes pourront prochainement être évaluées sur ce sous-ensemble de tours de parole.

Données	CONNEXION FORTE			CONNEXION SVM		
	MAN					
	$\bar{r} / r$	$\bar{p} / p$	$\overline{F-m} / F-m$	$\bar{r} / r$	$\bar{p} / p$	$\overline{F-m} / F-m$
Frames	0.93/0.88	0.95/0.88	0.93/0.88	0.95/0.92	0.94/0.86	0.93/0.89
FE	0.86/0.73	0.94/0.88	0.88/0.80	0.87/0.77	0.94/0.87	0.88/0.81
FE{Frames}	0.91/0.84	0.99/0.99	0.94/0.91	0.91/0.84	0.99/0.98	0.94/0.90
Liens	0.82/0.64	0.91/0.77	0.82/0.70	0.82/0.66	0.91/0.76	0.82/0.71
Liens{Frames}	0.88/0.73	0.98/0.96	0.91/0.83	0.88/0.75	0.98/0.96	0.91/0.84

Données	SLU					
	$\bar{r} / r$	$\bar{p} / p$	$\overline{F-m} / F-m$	$\bar{r} / r$	$\bar{p} / p$	$\overline{F-m} / F-m$
Frames	0.90/0.84	0.92/0.85	0.89/0.85	0.91/0.87	0.92/0.83	0.89/0.85
FE	0.83/0.70	0.91/0.84	0.84/0.76	0.84/0.73	0.91/0.82	0.84/0.77
FE{Frames}	0.91/0.84	0.98/0.97	0.93/0.90	0.91/0.83	0.98/0.96	0.93/0.89
Liens	0.80/0.61	0.90/0.74	0.79/0.67	0.80/0.63	0.89/0.73	0.79/0.67
Liens{Frames}	0.88/0.74	0.98/0.95	0.90/0.83	0.88/0.75	0.98/0.96	0.91/0.84

Données	ASR+SLU					
	$\bar{r} / r$	$\bar{p} / p$	$\overline{F-m} / F-m$	$\bar{r} / r$	$\bar{p} / p$	$\overline{F-m} / F-m$
Frames	0.82/0.77	0.86/0.78	0.80/0.77	0.83/0.78	0.86/0.76	0.80/0.77
FE	0.79/0.61	0.86/0.75	0.78/0.67	0.80/0.62	0.86/0.73	0.78/0.67
FE{Frames}	0.90/0.78	0.97/0.93	0.92/0.85	0.90/0.78	0.97/0.93	0.92/0.85
Liens	0.77/0.53	0.88/0.68	0.75/0.59	0.77/0.53	0.87/0.66	0.75/0.59
Liens{Frames}	0.87/0.68	0.97/0.94	0.90/0.79	0.87/0.68	0.97/0.94	0.90/0.79

TABLE 12.2 – Précision (p), rappel (r), F -mesure ($\overline{F-m}$), précision moyenne ($\bar{p}$), rappel moyen ($\bar{r}$) et F -mesure moyenne ($\overline{F-m}$) sur l'ensemble de test MEDIA après application des méthodes de connexion forte et SVM aux fragments sémantiques générés par le système basé sur le modèle DBN compact. Trois type de données ont été considérés : MAN, SLU et ASR+SLU.

		CONNEXION FORTE		CONNEXION SVM	
Données		MAN			
Frames	Total REF	Insertions	Suppressions	Insertion	Suppressions
FACILITY	299	1	59	10	5
HOTEL	604	5	25	97	19
LODGING	859	61	286	148	278
LOCATION	574	0	43	7	34
NUMBER	82	3	2	7	1
RESERVE	478	65	4	76	2
ROOM	641	39	35	39	35

		CONNEXION FORTE		CONNEXION SVM	
Données		SLU			
Frames	Total REF	Insertions	Suppressions	Insertion	Suppressions
FACILITY	299	10	78	24	30
HOTEL	604	22	42	95	36
LODGING	859	75	302	184	293
LOCATION	574	21	58	27	53
NUMBER	82	43	20	45	19
RESERVE	478	78	13	86	11
ROOM	641	40	58	40	58

		CONNEXION FORTE		CONNEXION SVM	
Données		ASR+SLU			
Frames	Total REF	Insertions	Suppressions	Insertion	Suppressions
FACILITY	299	25	95	39	57
HOTEL	604	46	88	124	85
LODGING	859	84	324	182	318
LOCATION	574	84	111	87	108
NUMBER	82	126	23	143	22
RESERVE	478	91	39	98	39
ROOM	641	96	70	96	70

TABLE 12.3 – Insertions et suppressions pour 7 types de frames sur l'ensemble de test MEDIA après application des méthodes de connexion forte et par classifieur SVM aux fragments sémantiques générés par le système basé sur le modèle DBN compact. Trois types de données ont été considérés : MAN, SLU et ASR+SLU.

Le tableau 12.3 présente pour les deux méthodes quelques résultats détaillés en termes d'occurrences de frames insérées et supprimées pour sept frames instanciées sur l'ensemble de test MEDIA. Ces frames offrent des situations très différentes. Si au niveau global les systèmes affichent des niveaux d'insertions et d'omissions assez équilibrés (comme le révèlent les valeurs de précision/rappel du tableau 12.2), on con-

state que les situations sont très variables selon les frames. Ainsi `LODGING` est supprimée massivement, `RESERVE` est essentiellement insérée et `ROOM` est autant insérée que supprimée (avec toutefois un taux d'identification très élevé, 90% sur les données manuelles). Une rapide analyse des erreurs nous a aussi permis de constater, sans surprise, que les principales difficultés concernaient les frames les plus "éloignées" des unités de base (i.e. celles qui doivent être inférées et ne peuvent être simplement déduites d'un concept présent dans l'hypothèse). Par exemple, `LODGING` a un taux d'identification de 60% sur les données manuelles tandis que celui de la frame `HOTEL` est de 95% .

Il est intéressant de remarquer que, contrairement aux résultats globaux, les différences de comportement entre les deux méthodes de connexion sont très visibles. Pour un même niveau de performance global, les deux méthodes prennent des décisions très différentes. Ce constat tend à renforcer notre hypothèse que la nature des données de test ne permet pas encore de mettre en avant plus clairement l'avantage de la méthode de connexion par classification SVM sur la méthode de connexion forte.

12.4 Conclusion

Les algorithmes de composition sémantiques proposés dans ce travail ont été évalués sur l'ensemble des 3005 dialogues du test `MEDIA` dans trois conditions de transcription et d'annotation conceptuelle différentes (tâches réalisées manuellement et/ou automatiquement). Les fragments sémantiques associés aux messages de test sont produits par le modèle DBN compact exposé au chapitre 8.

Les résultats obtenus par les deux algorithmes sont similaires. Les deux algorithmes s'avèrent capables de produire des représentations sémantiques complètes des messages utilisateur à partir des fragments sémantiques générés par le décodage séquentiel à base de DBN. Ils confirment la viabilité de l'approche de composition par combinaison d'arbres sémantiques partiels.

La méthode de connexion par classifieur SVM doit permettre une détection plus fine des relations entre les composants sémantiques au sein de la phrase. L'absence d'amélioration de performance dans nos expériences est en grande partie imputable à la nature des données utilisées qui comportent encore trop peu de situations suffisamment complexes en terme de représentation sémantique pour ne plus se satisfaire uniquement des relations déduites de l'ontologie.