



Annotation et recherche contextuelle des documents multimédias socio-personnels

Sonia Lajmi

► To cite this version:

Sonia Lajmi. Annotation et recherche contextuelle des documents multimédias socio-personnels. Autre [cs.OH]. INSA de Lyon, 2011. Français. NNT : 2011ISAL0025 . tel-00668689

HAL Id: tel-00668689

<https://theses.hal.science/tel-00668689>

Submitted on 10 Feb 2012

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

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



N° d'ordre : 2011-ISAL-0025



Année 2011



Thèse

Annotation et recherche contextuelle des documents multimédias socio-personnels

Présentée devant
L'Institut National des Sciences Appliquées de Lyon
(INSA de Lyon)

En cotutelle avec

La Faculté des Sciences Economiques et de Gestion de Sfax

Pour obtenir

Le grade de Docteur

Spécialité : Informatique

Par

Sonia LAJMI

Soutenue le 11 mars 2011

Commission d'Examen

Mr. Abdelmajid BEN HAMADOU	Professeur à l'ISIMSF de Sfax	Co-directeur de thèse
Mme Corine CAUVET	Professeur à l'université Paul Cézanne - Aix-Marseille 3	Examinatrice
Mr. Elöd EGYED-ZSIGMOND	Maître de conférences à l'INSA de Lyon	Co-encadrant
Mr. Mohamed Mohsen GAMMOUDI	Maître de conférences HDR à l'ESSAI, Université de Carthage, Tunisie	Rapporteur
Mme Christine LARGERON épouse LETENO	Professeur à l'université Jean Monnet - Saint-Etienne (UJM)	Examinatrice
Mr. Philippe MULHEM	Chargé de recherche Première Classe au CNRS, LIG, Université de Grenoble	Rapporteur
Mr. Jean Marie PINON	Professeur à INSA de Lyon	Co-directeur de thèse

Dédicaces

*A la mémoire de mon cher défunt père Saïd
qui m'a appris la noblesse de la science
qui m'a guidé dans mes premiers pas dans la vie*

*A la mémoire de mes chères grandes mères Zara et Baya, et
mon cher grand père Ahmed
qui m'ont élevé
Et m'ont tant appris dans ma tendre enfance.*

*A ma chère mère Salwa
le soleil de ma vie..
qui m'a insufflé toute une énergie pour aller jusqu'au bout...*

*A mes frères Mokhless et Nader
A ma belle sœur Ibtissem
qui a toujours été une vraie sœur*

*A ma cousine Ikram
qui a toujours été là pour moi
qui m'a supporté dans des épreuves très difficiles...*

Au reste de ma famille...

Remerciements

Ces travaux de recherche ont été réalisés dans le cadre d'une cotutelle entre le LIRIS (Laboratoire d'InfoRmatique en Images et Systèmes d'information), Université de Lyon et le MIRACL (Multimedia InfoRmation systems and Advanced Computing. Laboratory), Université de Sfax.

Je remercie avec beaucoup de reconnaissance et de considération Pr. Abdelmajid Ben Hamadou et Pr. Jean Marie Pinon qui ont montré beaucoup de patience surtout au démarrage difficile de ma thèse. Je leur remercie de m'avoir donné toute cette confiance et d'avoir cru à mes compétences.

Je voulais témoigner ma reconnaissance à Dr. Elöd Egyed Zsigmond, pour avoir assuré le suivi de mes travaux de recherche. Ses qualités humaines et ses conseils m'ont beaucoup aidé. Je le remercie, aussi, pour la confiance qu'il n'a cessée de me renouveler. Pour la grande liberté qu'elle m'a laissée et pour sa patience et son soutien dans les moments difficiles.

Je remercie très chaleureusement mes rapporteurs Pr. Mohamed Mohsen GAMMOUDI et Pr. Philippe MULHEM qui ont accepté de rapporter mon mémoire ainsi que les honorables membres de jury.

Je remercie chaleureusement Pr. Lionel Brunie et Pr. Sylvie Calabretto de m'avoir accueilli dans leur équipe de recherche DRIM (Distribution et Recherche d'Information Multimédia) au sein du LIRIS. Mille merci pour votre sympathie, d'avoir été des amis avec les thésards plus que des chefs d'équipe. Merci à Bénédicte la femme de Lionel pour sa personnalité très sympathique. Pour avoir été là pour soutenir la plupart des thésards de l'équipe. Merci pour tes délicieux fondant au chocolat.. emm un vrai régal !!

Je n'oublie pas tous les gens avec qui j'ai eu l'occasion de travailler. Je tiens à exprimer ma profonde gratitude à mes collègues Johann Stan et Pierre-Edward Portier ainsi que mes collaborateurs Dr. Hakim Hacid, Pr. Pierre Maret et Pr. Christine Solnon pour leurs esprits de collaboration, pour leurs disponibilité et leurs précieux conseils.

Merci à tous les membres des deux équipes de recherche dans les deux laboratoires MIRACL (université de Sfax) et LIRIS (INSA de Lyon) ainsi qu'aux membres du personnels des l'ISIM de Sfax et de l'INSA de Lyon. Merci pour les discussions scientifiques, pour les moments de détente parfois bien mérités. Je remercie particulièrement Achraf, Ahlem et Mohamed de m'avoir introduit le domaine des ontologies. Je remercie Dr. Nadia Benani et Dr. Sonia Benmokhtar pour leurs aides précieux et pour leurs disponibilités. Je remercie, également, Pierre-Antoine Champin pour tous les discussions scientifiques que nous avons élaborés autours des outils du Web sémantique. Je remercie Venessa pour sa sympathie et son café, Talar pour son amitié, son soutien et surtout ses corrections, Zena pour sa sociabilité et sa sympathie. Sana, Salma, Rami et Farah pour leur solidarité Tunisienne. Merci, aussi, à mes co-thésards; Omar pour ses conseils et Jingoway pour sa générosité et sa sympathie. Merci particulièrement à Jeanine Ressa, Charlotte Noireaux, Sandrine et Raymond Cortes.

Mes vifs remerciements s'adressent à mes amis Addisalem et Yaser qui m'ont accompagné tout le long de ce bout de chemin avec des conseils d'aînés, pour tous les moments inoubliables que nous avons passé ensemble. Ce travail n'aurait pas été possible sans votre soutien et votre aide morale et scientifique. Je n'oublierais jamais votre amitié. Même si le destin a fait que nos chemins se séparent, tout ce que vous avez fait pour moi que ce soit dans les bons ou difficiles moments restera gravé dans ma mémoire pour toujours.

Les conférences, journées de recherche, écoles d'automne sont l'occasion de faire des rencontres intéressantes. C'était le cas pour moi, merci à ceux qui ont rendu festive l'ambiance studieuse ! La liste est longue un merci particulier à Sawsan, Mouna et Eric.

Ma famille a été pour moi le principal moteur, chacun à sa manière de près ou de loin a su m'insuffler l'énergie nécessaire.

Je ne sais comment remercier ma très chère maman, le soleil de ma vie. Ce travail n'était pas possible sans la motivation que tu n'as cessé de me la renouvelé. Tu es ma source d'inspiration et d'énergie. Je t'aime plus que tous au monde. Merci pour tes prières, merci de te soucier toujours pour moi. Merci pour إنشاء الله يا بنيتي إنشاء الله , de

m'avoir toujours rappelé de la miséricorde de Dieu. J'espère que tu es fière de moi comme tu l'as toujours dit..ce travail était, essentiellement, pour toi.

Merci à mes amis en France, particulièrement à Mabrouka qui était plus qu'un membre du personnel du LIRIS. Elle était une vrai Amie. Merci pour ta serviabilité, ta générosité, d'être là quant j'avais besoin de toi. Merci à Sleh pour ses conseils d'aîné, pour son soutien et son aide, pour sa générosité et surtout de m'avoir considéré comme sa propre fille. J'espère vous avoir gardé comme amis pour toujours.

Merci à tous mes amis, particulièrement à Sameh, Manel, Mariam, Ikram Moalla, ma chère cousine Ikram, Khouloud Ghorbel, Emna, Mohamed Elleuch, Ibtissem, Habiba, ... Merci de votre soutien de ne pas m'avoir oublié pendant ma longue période d'hibernation, d'avoir toléré l'absence de nouvelles, les faux plans, les sauts d'humeur, les moments de doutes ... Les moments que j'ai passés avec vous furent les meilleurs, j'espère vous garder pour toujours.

Je voulais remercier Annie de m'avoir écouté et de faire tous pour m'aider quant j'avais besoin et ma chère Geneviève qui m'a ouvert grand les portes de sa maison et m'a accueillie. J'espère que vous serez toutes les deux dans le paradis de dieu pour votre générosité et bon cœur. Je ne vous oublierai jamais. Je me souviendrai toujours de Geneviève ; une maman qui a veillé sur moi en France. Je vous aime de tout mon cœur.

La thèse a été égaillée par les enseignements que j'ai pu donner. Je remercie vivement mes collègues enseignants à l'Université de Sfax (ISIM de Sfax), à l'université de Lyon (INSA de Lyon et Lyon 1 IUT B) et à l'université Jean Monnet (Faculté des sciences de saint Etienne).

Je remercie surtout tous mes étudiants, vous avez été mon rayon de soleil... Vous êtes ma véritable source de jouvence, merci pour toute cette énergie. J'ai hâte d'entendre encore "MADAME ça ne marche pas !" ..

La rumeur court qu'il y a une autre vie après la thèse .. Encore une hypothèse à formaliser, tester, expérimenter et prouver... tout un programme à implémenter ☺.

Sonia LAJMI, 11 mars 2011, Lyon, France

A handwritten signature in blue ink, appearing to read 'Sonia LAJMI', with a long horizontal stroke extending to the right.

Résumé

L'objectif de cette thèse est d'instrumentaliser des moyens, centrés utilisateur, de représentation, d'acquisition, d'enrichissement et d'exploitation des métadonnées décrivant des documents multimédias socio-personnels.

Afin d'atteindre cet objectif, nous avons proposé un modèle d'annotation, appelé *SeMAT* avec une nouvelle vision du contexte de prise de vue. Nous avons proposé d'utiliser des ressources sémantiques externes telles que GeoNames¹ et Wikipédia² pour enrichir automatiquement les annotations partant des éléments de contexte capturés. Afin d'accentuer l'aspect sémantique des annotations, nous avons modélisé la notion de profil social avec des outils du web sémantique en focalisant plus particulièrement sur la notion de liens sociaux et un mécanisme de raisonnement permettant d'inférer de nouveaux liens sociaux non explicités. Le modèle proposé, appelé *SocialSphere*, construit un moyen de personnalisation des annotations suivant la personne qui consulte les documents (le consultateur). Des exemples d'annotations personnalisées peuvent être des objets utilisateurs (e.g. maison, travail) ou des dimensions sociales (e.g. ma mère, le cousin de mon mari). Dans ce cadre, nous avons proposé un algorithme, appelé *SQO*, permettant de suggérer au consultateur des dimensions sociales selon son profil pour décrire les acteurs d'un document multimédia. Dans la perspective de suggérer à l'utilisateur des événements décrivant les documents multimédias, nous avons réutilisé son expérience et l'expérience de son réseau de connaissances en produisant des règles d'association. Dans une dernière partie, nous avons abordé le problème de correspondance (ou appariement) entre requête et graphe social. Nous avons proposé de ramener le problème de recherche de correspondance à un problème d'isomorphisme de sous-graphe partiel. Nous avons proposé un algorithme, appelé *h-Pruning*, permettant de faire une correspondance rapprochée entre les nœuds des deux graphes : motif (représentant la requête) et social.

Pour la mise en œuvre, nous avons réalisé un prototype à deux composantes : web et mobile. La composante mobile a pour objectif de capturer les éléments de contexte lors de la création des documents multimédias socio-personnels. Quant à la composante web, elle est dédiée à l'assistance de l'utilisateur lors de son annotation ou consultation des documents multimédias socio-personnels. L'évaluation a prouvé : (i) l'efficacité de notre approche de recherche dans le graphe social en termes de temps d'exécution ; (ii) l'efficacité de notre approche de suggestion des événements (en effet, nous avons prouvé notre hypothèse en démontrant l'existence d'une cooccurrence entre le contexte spatio-temporel et les événements) ; (iii) l'efficacité de notre approche de suggestion des dimensions sociales en termes de temps d'exécution.

Mots-clés: Annotation et Recherche Sémantique, Web 2.0, Web Sémantique, Ontologies, Réseaux sociaux, Raisonnement Sémantique, Sensibilité au contexte social, Profil social, Graphe, Appariement de graphe.

¹ www.geonames.org/

² fr.wikipedia.org/

Abstract

The overall objective of this thesis is to exploit a user centric means of representation, acquisition, enrichment and exploitation of multimedia document metadata.

To achieve this goal, we proposed an annotation model, called *SeMAT* with a new vision of the snapshot context. We proposed the usage of external semantic resources (e.g. GeoNames³, Wikipedia⁴, etc.) to enrich the annotations automatically from the snapshot contextual elements. To accentuate the annotations semantic aspect, we modeled the concept of ‘social profile’ with Semantic web tools by focusing, in particular, on social relationships and a reasoning mechanism to infer a non-explicit social relationship. The proposed model, called *SocialSphere* is aimed to exploit a way to personalize the annotations to the viewer. Examples can be user’s objects (e.g. home, work) or user’s social dimensions (e.g. my mother, my husband’s cousin). In this context, we proposed an algorithm, called *SQO* to suggest social dimensions describing actors in multimedia documents according to the viewer’s social profile. For suggesting event annotations, we have reused user experience and the experience of the users in his social network by producing association rules. In the last part, we addressed the problem of pattern matching between query and social graph. We proposed to steer the problem of pattern matching to a sub-graph isomorphism problem. We proposed an algorithm, called *h-Pruning*, for partial sub-graph isomorphism to ensure a close matching between nodes of the two graphs: motive (representing the request) and the social one.

For implementation, we realized a prototype having two components: mobile and web. The mobile component aims to capture the snapshot contextual elements. As for the web component, it is dedicated to the assistance of the user during his socio-personnel multimedia document annotation or socio-personnel multimedia document consultation. The evaluation have proven: (i) the effectiveness of our exploitation of social graph approach in terms of execution time, (ii) the effectiveness of our event suggestion approach (we proved our hypothesis by demonstrating the existence of co-occurrence between the spatio-temporal context and events), (iii) the effectiveness of our social dimension suggestion approach in terms of execution time.

Keywords: Semantic annotation and retrieval, Web 2.0, Semantic Web, Ontology, Social Network, Semantic reasoning, Social-aware, Social profile, Graph, Graph matching.

³ www.geonames.org/

⁴ fr.wikipedia.org/

Table des Matières

DEDICACES	III
REMERCIEMENTS	IV
RESUME	IX
ABSTRACT.....	XI
TABLE DES MATIERES	XIII
LISTE DES FIGURES	XIX
LISTE DES TABLEAUX.....	XXI
LISTE DES ALGORITHMES	XXIII
LISTE DES DEFINITIONS.....	XXV
CHAPITRE 1. INTRODUCTION	1
1.1. CONTEXTE D’ETUDE	3
1.2. ENJEUX	4
1.3. OBJECTIFS	6
1.4. CONTRIBUTIONS	7
1.5. ORGANISATION DE LA THESE.....	8
CHAPITRE 2. TECHNOLOGIES ET FORMALISMES DE REPRESENTATION DE DOCUMENTS ENRICHIS	11
2.1. INTRODUCTION.....	13
2.2. DOCUMENT MULTIMEDIA SOCIO-PERSONNEL	13
2.3. TAGGING, METADONNEES ET ANNOTATIONS	14
2.4. TECHNOLOGIES DU WEB SEMANTIQUE.....	15
2.4.1. <i>Les formats RDF et RDFS</i>	16
2.4.2. <i>Ontologies et le format OWL</i>	18
2.4.3. <i>Framework Jena</i>	18
2.4.4. <i>Langage de requête SPARQL</i>	19
2.5. FORMALISMES DE REPRESENTATION D’ANNOTATIONS	20
2.5.1. <i>Dublin Core</i>	20
2.5.2. <i>Information Interchange Model (IIM)</i>	21
2.5.3. <i>EXIF- EXchangeable Image File</i>	21
2.5.4. <i>MPEG-7</i>	22

2.5.5.	<i>L'ontologie DIG35</i>	22
2.5.6.	<i>PhotoRDF</i>	23
2.5.7.	<i>VRA Core</i>	24
2.5.8.	<i>XMP</i>	24
2.5.9.	<i>OntoAlbum</i>	24
2.5.10.	<i>Discussion</i>	26
2.6.	FORMALISMES DE REPRESENTATION DES LIENS SOCIAUX	26
2.6.1.	<i>CSV et format .dot</i>	27
2.6.2.	<i>DyNetML</i>	27
2.6.3.	<i>Social Content graph</i>	28
2.6.4.	<i>XFN</i>	28
2.6.5.	<i>FOAF et ses extensions</i>	29
2.6.6.	<i>Discussion</i>	29
2.7.	RESUME	30
CHAPITRE 3.	MESURES DE SIMILARITE SEMANTIQUE	31
3.1.	INTRODUCTION	33
3.2.	TERMINOLOGIE ET NOTATIONS	33
3.3.	CALCUL DE SIMILARITE PAR LE NOMBRE D'ARCS	35
3.4.	CALCUL DE SIMILARITE PAR LE CONTENU INFORMATIF	37
3.5.	CALCUL DE SIMILARITE HYBRIDE	39
3.6.	RESUME ET DISCUSSION	39
CHAPITRE 4.	ENRICHISSEMENT DES ANNOTATIONS	41
4.1.	INTRODUCTION	43
4.2.	ANNOTATIONS BASEES SUR LES RESSOURCES SEMANTIQUES	44
4.2.1.	<i>RisuPicWeb</i>	44
4.2.2.	<i>Lim et al. 2003</i>	45
4.2.3.	<i>Hollink et al. 2003</i>	46
4.2.4.	<i>Shevade et al. 2005</i>	46
4.2.5.	<i>Discussion</i>	47
4.3.	ANNOTATIONS BASEES SUR UNE COLLABORATION A L'ANNOTATION MANUELLE	47
4.4.	ANNOTATIONS BASEES SUR DES DONNEES PROVENANT D'UN RESEAU SOCIAL	47
4.4.1.	<i>Shevade et al. 2007</i>	48
4.4.2.	<i>Zunjarward et al. 2007</i>	48
4.4.3.	<i>Stone et al 2008</i>	48
4.4.4.	<i>Elliot et al. 2009</i>	48

4.4.5.	<i>Discussion</i>	49
4.5.	SUGGESTION DE TAGS LIBREMENT SAISIS	49
4.5.1.	<i>MiAlbum</i>	49
4.5.2.	<i>Sigurbjörnsson et Zwol 2008</i>	50
4.5.3.	<i>Kucuktunc et al. 2008</i>	50
4.5.4.	<i>Lvanov et al. 2010</i>	50
4.5.5.	<i>Discussion</i>	51
4.6.	ANNOTATIONS BASEES SUR LA CAPTURE DU CONTEXTE PHYSIQUE	52
4.6.1.	<i>Enrichissement de contexte physique par des ressources impersonnelles</i>	52
4.6.2.	<i>Enrichissement du contexte physique par des ressources spécifiques à l'utilisateur</i>	56
4.6.3.	<i>Discussion</i>	57
4.7.	RESUME.....	58
CHAPITRE 5.	MODELISATION ET ENRICHISSEMENTDES METADONNEES	61
5.1.	INTRODUCTION.....	62
5.2.	DEFINITIONS.....	62
5.3.	MODELISATION DES METADONNEES EN CONSIDERANT LE CONTEXTE	63
5.3.1.	<i>Vue globale des métadonnées et de leur relation avec le contexte de prise de vue</i>	63
5.3.2.	<i>Modèle centré utilisateur de représentation des métadonnées</i>	66
5.3.3.	<i>Modélisation du profil social pour enrichir les métadonnées</i>	68
5.3.3.1.	Description des liens sociaux et des catégories sociales	70
5.3.3.2.	Règles d'inférence de sens commun	72
5.3.3.3.	Contraintes d'intégrités.....	78
5.3.4.	<i>Discussion</i>	78
5.4.	PLATEFORME D'ENRICHISSEMENT DES METADONNEES.....	81
5.4.1.	<i>Architecture de la plateforme d'enrichissement</i>	81
5.4.2.	<i>Génération des éléments de la couche physique du contexte</i>	83
5.4.2.1.	Exemple de capture du contexte physique lors de la création	84
5.4.3.	<i>Contexte de création enrichi</i>	85
5.4.3.1.	Enrichissement du contexte temporel	86
5.4.3.2.	Enrichissement du contexte spatial	86
5.4.3.3.	Création du contexte spatio-temporel	89
5.4.3.4.	Exemple d'enrichissement automatique du contexte de prise de vue.....	90
5.4.4.	<i>Contexte de création personnalisé</i>	91
5.4.4.1.	Enrichissement par des objets utilisateurs.....	92
5.4.4.2.	Enrichissement par événement.....	93
5.4.4.3.	Enrichir par l'identité des personnes proches.....	94
5.4.4.4.	Exemple d'enrichissement personnalisé pour Bob du contexte physique	94

5.4.5.	<i>Enrichissement manuel des métadonnées du contenu</i>	95
5.5.	RESUME ET DISCUSSION	96
CHAPITRE 6. SENSIBILITE SOCIALE DES RECOMMANDATIONS D'ANNOTATIONS 99		
6.1.	INTRODUCTION	100
6.2.	MOTIVATIONS.....	101
6.3.	RECOMMANDATION DES EVENEMENTS POUR ANNOTER DES PHOTOS	102
6.3.1.	<i>Nature des données considérées</i>	104
6.3.2.	<i>Découverte des motifs fréquents</i>	106
6.3.3.	<i>Découverte des règles d'association</i>	107
6.3.4.	<i>Suggestion de tag événement basée sur la proximité sociale</i>	109
6.3.5.	<i>Discussion</i>	113
6.4.	ENRICHISSEMENT DES TAGS IDENTITE LORS DE LA CONSULTATION DES PHOTOS	114
6.4.1.	<i>Heuristiques considérées</i>	116
6.4.2.	<i>Algorithme SQO</i>	118
6.4.3.	<i>Analyse de la complexité de l'algorithme SQO</i>	121
6.4.4.	<i>Discussion</i>	122
6.5.	RESUME	122
CHAPITRE 7. EXPLOITATION DU GRAPHE SOCIAL.....123		
7.1.	SCENARIO INTRODUCTIF	125
7.2.	UNE CONCEPTUALISATION ABSTRAITE DU GRAPHE SOCIAL.....	126
7.3.	GRAPHE SOCIAL	128
7.3.1.	<i>Arcs</i>	129
7.3.2.	<i>Nœuds</i>	129
7.3.3.	<i>Attributs d'un nœud</i>	129
7.3.4.	<i>La liste des attributs pour chaque nœud du graphe social</i>	131
7.4.	GRAPHE MOTIF POUR L'EXPRESSION DES REQUETES.....	132
7.4.1.	<i>Nœuds du graphe motif</i>	133
7.4.2.	<i>Attributs des nœuds du graphe motif</i>	133
7.4.3.	<i>Fonctions de comparaison du graphe motif</i>	133
7.4.4.	<i>Les arcs du graphe motif</i>	134
7.5.	INSTANCIER LES GRAPHE MOTIFS	135
7.5.1.	<i>Isomorphisme de sous-graphe partiel</i>	136
7.5.2.	<i>Notations et définitions</i>	138
7.5.3.	<i>Algorithme d'isomorphisme de sous-graphe partiel</i>	140
7.5.3.1.	<i>Construction du domaine D_x initial</i>	142
7.5.3.2.	<i>Propagation des contraintes de différence FC_{Diff}</i>	146

7.5.3.3.	Propagation des contraintes d'arrêt FC_{Edges}	147
7.5.3.4.	Propagation des contraintes d'arrêt AC_{Edges}	148
7.5.3.5.	Considération de la similarité entre nœuds pour filtrer les domaines de valeurs.....	150
7.5.3.6.	Analyse de la complexité de notre algorithme h-Pruning	150
7.5.4.	<i>Exemple complet</i>	151
7.6.	RESUME ET DISCUSSION	157
CHAPITRE 8.	IMPLEMENTATION ET EVALUATION	159
8.1.	INTRODUCTION.....	160
8.2.	PRESENTATION DU PROTOTYPE	161
8.2.1.	<i>Application mobile : PASMi</i>	163
8.2.2.	<i>Composante web pour l'organisation et la publication de photos : Phoox</i>	166
8.2.2.1.	Fonctionnalités.....	166
8.2.2.2.	Les bases de données	172
8.2.2.3.	Représentation et accès au profil social	175
8.3.	EXPERIMENTATION ET EVALUATION DES PERFORMANCES.....	176
8.3.1.	<i>Construction de collection de test</i>	177
8.3.2.	<i>Évaluation de l'approche de l'exploitation du graphe social</i>	178
8.3.3.	<i>Evaluation de l'approche de suggestion des dimensions sociales</i>	181
8.3.4.	<i>Evaluation de l'approche de suggestion d'événements</i>	183
8.4.	DISCUSSION.....	188
8.5.	RESUME.....	189
CHAPITRE 9.	CONCLUSION ET PERSPECTIVES	191
9.1.	CONCLUSION	193
9.2.	BILAN DES CONTRIBUTIONS.....	194
9.3.	PERSPECTIVES	197
GLOSSAIRE DES ACRONYMES	199	
BIBLIOGRAPHIE	205	
ANNEXES	219	
ANNEXE I.	L'ONTOLOGIE SOCIALSPHERE	220
ANNEXE II.	L'ONTOLOGIE SEMAT	227
ANNEXE III.	REGLES D'INFERENCES EXPLICITES DE SENS COMMUN	231
ANNEXE IV.	EXEMPLES D'UTILISATION DU FRAMWORK JENA	237

Liste des Figures

Figure 1. Les différentes couches du web sémantique.....	16
Figure 2. Exemple de graphe RDF	17
Figure 3. Ontologie d'annotation de photos proposée dans [3]	25
Figure 4. Taxonomie des personnes proposée dans [3]	25
Figure 5. Structure du schéma XML du DyNetML présentée dans [46].....	28
Figure 6. Partie de la structure du thesaurus WordNet représentant le tag ‘party’ dans deux sens.....	34
Figure 7. Les relations conceptuelles [52].....	37
Figure 8. Extrait de WordNet présenté dans [53] avec les probabilités correspondantes.....	38
Figure 9. Framework général de suggestion d’annotation de photos proposé par [60]	44
Figure 10. Processus d'annotation et de recherche de photos basé sur l'apprentissage [62]	45
Figure 11. Processus d'annotation et de recherche des photos de l’art proposé dans [42].....	46
Figure 12. Framework d'annotation décrit dans [78]	51
Figure 13. Processus d'annotation de photocopain décrit dans [86]	54
Figure 14. Interface d'assistance à l'annotation de MediaAssist décrit dans [90], [91].....	55
Figure 15. Architecture d’ACRONYM proposée dans [92].....	56
Figure 16. Vue générale du système d'annotation et de recherche proposée par [98].....	57
Figure 17. Les éléments de contexte et ses différentes couches.....	66
Figure 18. Squelette des concepts de l’ontologie <i>SeMAT</i> et l’ontologie <i>SocialSphere</i>	67
Figure 19. Réseau social enrichie avec l'ontologie <i>SocialSphere</i>	72
Figure 20: Comparaison entre les deux représentations des métadonnées sur les relations avec UML et avec RDF.....	80
Figure 21. Architecture de la plateforme d'enrichissement semi-automatique des annotations.....	82
Figure 22. Exemple d'annotation automatique du contexte de prise de vue de la photo	84
Figure 23. Exemple de création automatique de contexte de prise de vue	85
Figure 24. Entrées et sorties du module CCE	86
Figure 25. Enrichissement de contexte physique capturé	90
Figure 26. Composant du module CCP	92
Figure 27. Exemple d'enrichissement personnalisé se servant du profil social: l’objet d'utilisateur de May, une participante à la prise de vue est rattachée à la photo	93
Figure 28. Insertion de métadonnées sur le contenu par l'utilisateur :actor et socialInteraction	96

Figure 29. Étapes d'extraction de règles à partir de données	104
Figure 30. Enrichissement des annotations avec dimension sémantique pour deux utilisateurs	115
Figure 31. Exemple de réseau social parcouru depuis le profil social de <i>Bob</i>	116
Figure 32: Modèle du graphe social à deux niveaux	128
Figure 33. Définition de nœuds avec un ensemble d'attributs et d'instances.....	131
Figure 34. Graphe motif ramenant un réseau social.....	134
Figure 35. Exemple de graphe motif représentant la requête de David	135
Figure 36. Isomorphisme de sous graphe partiel entre GM et GS.....	137
Figure 37. Etapes d'exécution de l'algorithme <i>h-Pruning</i>	142
Figure 38. Architecture du prototype à deux composants: (a) <i>PASMi</i> (b) <i>Phoox</i>	163
Figure 39. Exemple de fichier de représentation de métadonnées	164
Figure 40. Quelques captures d'écran de l'application <i>PASMi</i>	165
Figure 41. Interface d'upload de l'application <i>Phoox</i>	167
Figure 42. Interface d'annotation automatique de contexte temporel de prise de vue	167
Figure 43. Interface d'annotation de contexte environnemental de prise de vue	168
Figure 44. Interface d'annotation et suggestion de contexte spatial de prise de vue	169
Figure 45. Interface d'annotation de contexte social de la prise de vue	170
Figure 46. Interface de suggestion des événements	171
Figure 47. Interface de proposition des dimensions sociales	171
Figure 48. Diagramme de classe de l'API graphe.....	173
Figure 49. Diagramme de classes de l'application <i>Phoox</i>	174
Figure 50. Interface de présentation et édition de profil sociale	175
Figure 51. Représentation du profil social sous format RDF/XML.....	176
Figure 52. Temps d'exécution des deux algorithmes en fonction du nombre de nœuds dans G_S	180
Figure 53. Temps d'exécution des deux algorithmes en fonction du nombre de nœuds dans G_M	181
Figure 54. Temps d'exécution des deux algorithmes BFS et SQO en fonction de la taille du graphe ..	183
Figure 55. Nombres d'utilisateurs ayant x règles produites	186
Figure 56. Nombre d'utilisateurs ayant une proximité sociale donnée avec d'autres utilisateurs	188

Liste des Tableaux

Tableau 1. Exemple de résultat d'une requête SPARQL	20
Tableau 2. Les techniques d'annotations.....	43
Tableau 3: les différentes relations sociales ainsi que leurs différentes catégories	71
Tableau 4. Exemple de règles et contraintes en OWL.....	77
Tableau 5. Exemples de lignes des données considérées pour générer les règles d'associations	105
Tableau 6. Proximité sociale entre Alice, Carol, Bob et May	111
Tableau 7. Types Att_n (et leur multiplicité $\ast=0$ ou n) pour chaque spécification de nœuds.....	132
Tableau 8. . Statistiques sur la collection de données construite à partir de Flickr.....	178
Tableau 9. Temps d'exécution des deux algorithmes en fonction du nombre de nœuds dans G_S	179
Tableau 10. Temps d'exécution des algorithmes en variant le nombre de nœuds dans G_M	180

Liste des Algorithmes

Algorithme 1. Algorithme <i>RSP</i> de recommandation de tag événement via la proximité sociale	112
Algorithme 2. Algorithme <i>SQO</i> de recherche en largeur avec heuristique d'optimisation.....	120
Algorithme 3. Algorithme <i>MarquerNoeuds</i> de sélection des les nœuds potentiellement intéressants .	144
Algorithme 4. Algorithme <i>RSS</i> de réduction de l'espace de recherche.....	144
Algorithme 5. Algorithme <i>Construction D_x</i> de construction des domaines de valeurs initiaux	146
Algorithme 6. Algorithme Filtrage FC_{Diff}	147
Algorithme 7. Algorithme Filtrage FC_{Edges}	148
Algorithme 8. Algorithme Filtrage AC_{Edges}	149

Liste des Définitions

Définition 1. Document multimédia socio-personnel.....	14
Définition 2. Activité documentaire	62
Définition 3. Interaction sociale.....	62
Définition 4. Évènement.....	62
Définition 5. Proximité sociale.....	63
Définition 6. Dimension sociale.....	63
Définition 7. Tag	63
Définition 8. Contexte de prise de vue	64
Définition 9. Profil social.....	69
Définition 10. Objet d'utilisateur :	69
Définition 11. Règles d'association spécifiques à un utilisateur.....	111
Définition 12. Graphe Social.....	128
Définition 13. Graphe Motif	132
Définition 14. Réseau Social.....	134
Définition 15. Correspondance.....	138
Définition 16. Domaine de Valeurs D_x	139
Définition 17. Appariement	139
Définition 18. Adjacent $\text{adj}(x)$	139
Définition 19. Distance entre deux nœuds $d(x, x')$	139
Définition 20. Excentricité de x	140
Définition 21. Degré de x $\text{deg}(x)$	140
Définition 22. Restriction sur N_S	140
Définition 23. Projection d'un nœud selon un attribut.....	140

Chapitre 1. Introduction

1.1.	CONTEXTE D'ETUDE	3
1.2.	ENJEUX	4
1.3.	OBJECTIFS	6
1.4.	CONTRIBUTIONS	7
1.5.	ORGANISATION DE LA THESE	8

1.1. Contexte d'étude

De nos jours, les ordinateurs jouent un rôle important dans la gestion des documents multimédias. Ces documents multimédias, qui peuvent être des vidéos, des enregistrements sonores, des photos, etc., sont stockés, partagés, recherchés et visualisés dans un format numérique.

La gestion des grandes collections de documents multimédias devient de plus en plus une tâche difficile. Comme le coût de capture et de stockage des documents multimédias devient de moins en moins cher, nous allons à grands pas vers un stockage de toute une vie de documents [1]. En même temps, la question de l'utilisabilité de ces documents se pose vue que les méthodes d'accès et de la recherche restent encore limitées. Avec les documents multimédias numériques, la possibilité de trouver ce qu'on souhaite est séduisante, mais encore loin d'être satisfaisante.

En même temps, l'ouverture sur le web a fournit la possibilité de partager, plus aisément, les documents multimédias. Actuellement, Facebook est élu premier pour le partage des photos avec plus de trois milliards de photos au total et plus de quatre cents millions d'utilisateurs actifs dont 50% sont des utilisateurs mobiles [2]. Les utilisateurs sont confrontés à une gestion d'une masse gigantesque de documents multimédias et exigent non seulement de gérer leurs propres documents multimédias (dont ils sont les créateurs) mais aussi ceux de leurs réseaux sociaux ce qui amplifie encore le problème de gestion des documents multimédias. Ainsi, les systèmes de gestion des documents multimédias offerts doivent prendre en considération l'aspect collaboratif aussi bien au niveau de l'annotation que celui de la recherche et la visualisation.

Dès lors, des systèmes tels que Flickr, Facebook et d'autres comme Getty Images et Photoree se sont intéressés à proposer des moyens d'organisation de corpus de photos en mettant à la disposition de l'utilisateur des fonctionnalités classiques comme la protection de leurs publications, la classification en galeries, etc. Néanmoins, la recherche des photos reste la fonctionnalité la moins satisfaisante dans les systèmes commerciaux de gestion de photos [3]. La plupart des systèmes existants se basent sur les annotations (tags) des utilisateurs. Les photos annotées dans les sites les plus populaires restent relativement peu nombreuses en comparaison avec le nombre très élevé de photos non annotées. Rares sont les utilisateurs qui remplacent par un nom significatif celui attribué par défaut, par leurs

appareils-photo numériques. Les utilisateurs ont tendance à passer très peu de temps à annoter leurs photos [4]. Pourtant, l'annotation de photos reste un enjeu majeur pour assurer la pertinence de la recherche de ces photos.

1.2. Enjeux

Les enjeux majeurs des applications de gestion des documents multimédias peuvent être résumés par les points suivants :

- **Consistance versus subjectivité** : l'annotation est subjective et il faut qu'elle le reste. Tant que la personne qui annote est toujours l'unique utilisateur de ses documents multimédias, une annotation libre (à la manière de Flickr par exemple) pourrait être satisfaisante dans certaines conditions (i.e. corpus de photos limité et récent). Cependant, ce type d'annotation devient inefficace en passant à l'échelle et l'exploitation collective. En effet, dans le cadre de leur utilisation actuelle, les documents multimédias sont, le plus souvent, partagés. Le contexte de partage et d'utilisation à plusieurs impose une certaine cohérence au niveau de l'annotation. Comment assurer une annotation cohérente des documents multimédias tout en sachant qu'elle restera toujours subjective ? En d'autres termes il s'agit de trouver un compromis entre la cohérence des annotations et leurs caractères subjectifs.
- **Exploitation des documents multimédias par contexte** : des études empiriques d'utilisabilité de l'accès par contexte ont été fournies dans [5] et [6] et [1]. Ces études constatent qu'une partie importante des informations qui aident les personnes à se souvenir d'une photo sociale font référence au contexte. Les éléments de ce contexte sont le moment de la prise de vue (quand), la localisation (où) et les personnes y figurant (qui). De plus, vu le nombre important d'utilisateurs mobiles actifs sur les systèmes de réseaux sociaux, l'hypothèse de capture du contexte via l'environnement ambiant devient de plus en plus réaliste [2]. Comment pourrait-on, alors, se servir de capteurs physiques pour inférer des informations contextuelles utiles et exploitables lors de la recherche des photos ?
- **Exploitation des documents multimédias par interactions sociales**: les interactions humaines sont présentes partout: dans les photos sociales, dans les films, dans les émissions de télévision, dans les programmes radio, dans des enregistrements de réunions, etc. [7]. En outre, des études psychologiques ont

montré que les interactions sociales sont l'un des principaux canaux par lesquels nous comprenons la réalité [8]. Par exemple, les téléspectateurs sont le plus souvent intéressés par les interactions qui relient les acteurs (celui-ci aime celle là, celui-ci a tué celui-là, etc.) plus que par ce qu'ils disent. Dans le cas des documents multimédias sociaux, ce qui intéresse en premier les utilisateurs est de comprendre les interactions sociales entre les acteurs. Comment pourrait-on alors rendre les documents multimédias exploitables via les interactions entre acteurs ?

- **Aspect personnel de la description du contenu et du contexte** : des études empiriques [9] ont prouvé que l'utilisateur préfère organiser, annoter et exploiter (rechercher) les photos sociales en utilisant des axes sémantiques personnalisés comme l'événement (par exemple, 'mon anniversaire', 'excursion en famille au parc', etc.), les acteurs (par exemple, 'moi', 'mon fils', 'ma mère', etc.), les indicateurs temporels (par exemple, 'mois passé', 'cette année', etc.) et les indicateurs spatiaux (par exemple, 'ma maison', 'lieu de travail de mon frère'). Ces mots-clés, qui représentent le contexte et le contenu des photos décrit d'une manière personnalisée, ne font sens que pour l'annotateur ou le consultateur mais demeurent insignifiants pour les autres utilisateurs. Comment pourrait-on rendre les documents sociaux exploitables via des indicateurs personnalisés qui représentent le contenu et/ou le contexte tout en gardant la consistance des annotations ?
- **Caractère dynamique des annotations** : les annotations doivent non seulement être personnelles mais aussi dynamiques. En effet, si nous considérons les indicateurs contextuels personnels, ces derniers peuvent changer d'essence dans le temps, par exemple, changement de lieu de travail, déménagement, etc. Dans ce cas, il faut investiguer des moyens et des outils permettant de prendre en considération ce changement.
- **Résolution des requêtes portés sur les documents multimédias** : la théorie des graphes est un moyen puissant qui a servi, depuis longtemps, à apporter des solutions efficaces pour de nombreux problèmes. Le corpus de documents multimédias peut être modélisé en graphe. L'interrogation de ce dernier peut être réalisée en investiguant des algorithmes de la théorie des graphes.

1.3. Objectifs

Compte tenu des différents enjeux soulevés, un système de gestion de documents multimédias doit être conçu dans la finalité d'aider l'utilisateur à retrouver plus facilement ses documents multimédias ou les documents multimédias de son réseau de connaissance.

L'objectif global de cette thèse est de proposer des moyens de représentation, d'acquisition, d'enrichissement et d'exploitation des métadonnées qui décrivent les documents multimédias tout en veillant à prendre en considération les nombreux enjeux soulevés. Cet objectif peut être réalisé à travers :

- la proposition d'un modèle de représentation des métadonnées permettant à l'utilisateur d'accéder aux documents multimédias par contexte, par contenu et par interaction sociale. Ce modèle doit être interopérable avec les standards existants.
- la proposition des mécanismes d'acquisition automatique et d'interprétation du contexte afin d'enrichir les métadonnées. Nous devons considérer tous les éléments du contexte qui peuvent être capturés via l'environnement ambiant de l'utilisateur lors de la prise de vue et qui peuvent être enrichis afin de faciliter la tâche de gestion des documents multimédia.
- la construction des moyens de visualisation des documents multimédias et de leurs métadonnées appropriées pour chaque utilisateur.
- l'édification des moyens d'assistance à l'annotation et à la recherche en se servant de l'intelligence collective.
- l'amélioration des moyens d'assistance à l'annotation et à la recherche en se servant des ressources sémantiques.
- l'édification d'algorithmes de recommandations via la réutilisation de l'expérience de l'utilisateur et de son réseau de connaissance.
- la résolution des requêtes portées sur le graphe social formé par l'ensemble de documents multimédia, d'annotations, d'utilisateurs et de liens sociaux. Les moyens investigués doivent profiter du domaine de la théorie des graphes.

1.4. Contributions

Nous avons proposé un modèle d’annotation avec une nouvelle vision du contexte de prise de vue (section 5.3 Chapitre 5). Nous avons, également, modélisé la notion de profil social (section 5.3.3 Chapitre 5) avec des outils du web sémantique en focalisant, plus particulièrement, la notion de liens sociaux et un mécanisme de raisonnement permettant d’inférer de nouveaux liens sociaux non explicites.

La modélisation du profil social vise à instrumentaliser un moyen de personnalisation des annotations suivant le consultateur. Par exemple, si nous considérons l’identité des acteurs figurant dans un document multimédia, il serait intéressant de présenter au consultateur des dimensions sociales expliquant la nature du lien social qui le relie à cet acteur là.

Nous avons également, proposé d’utiliser des ressources externes (comme GeoNames⁵, Wikipédia⁶, etc.) pour enrichir les annotations partant des éléments de contexte capturés lors de la prise de vue (section 5.4.3 Chapitre 5).

Pour accentuer l’aspect sémantique des annotations, nous avons proposé d’utiliser le profil social dans le but de suggérer à l’utilisateur une liste d’annotations sémantiques (section 5.4.4 Chapitre 5). Des exemples peuvent être des objets utilisateurs (maison, travail, etc.) ou des dimensions sociales (ma femme, le cousin de mon mari, etc.).

Dans la section 6.3 du Chapitre 6, nous avons, aussi, réutilisé l’expérience de l’utilisateur et des membres de son réseau de connaissance en appliquant un algorithme de fouille de données sur l’ensemble de sa base de documents multimédias annotés. La réutilisation de l’expérience de l’utilisateur est effectuée dans l’objectif de lui proposer de nouvelles annotations. Plus particulièrement, nous nous sommes intéressés par la suggestion des événements qui décrivent les photos et les collections.

Dans une dernière partie de notre travail de thèse présenté dans le Chapitre 7, nous avons abordé le problème de correspondance (ou appariement) entre requête et graphe social formé par l’ensemble de photos, d’annotations, d’utilisateurs et de liens sociaux. Nous avons proposé de ramener le problème de recherche de correspondance à un problème d’isomorphisme de sous-graphe partiel. Nous avons, alors, proposé un

⁵ www.geonames.org/

⁶ fr.wikipedia.org/

algorithme, appelé *h-Pruning*, d'isomorphisme de sous-graphe partiel permettant de faire une correspondance rapprochée entre les nœuds des deux graphes : motif (représentant la requête) et social. La comparaison entre les nœuds des deux graphes prend en considération la similarité de leurs attributs. Cette similarité peut être calculée en spécifiant un ou plusieurs mesures de similarité sémantique.

Les approches proposées ont été mises en œuvre dans un prototype à deux composants, web et mobile. La composante Mobile, baptisée PASMi, présentée dans la section 8.2.1 Chapitre 8, a pour objectif de capturer les éléments de contexte lors de la création des photos et de capturer des traces de partage et d'interaction entre les utilisateurs mobiles. Quant au composant web, appelée Phoox, présenté dans la section 8.2.2 Chapitre 8, il est dédié à l'assistance de l'utilisateur lors de son annotation, ou consultation des photos. Ce composant permet, aussi, l'exploitation du graphe social en permettant à l'utilisateur d'exprimer sa requête sous forme de graphe motif et de faire la correspondance.

Nous avons, aussi, évalué les approches proposées en utilisant une collection de tests construite à partir du célèbre service de médias sociaux Flickr⁷. Le protocole de test est présenté dans la section 8.3.1 du Chapitre 8. Les tests ont prouvés :

- l'efficacité de notre approche de recherche dans le graphe social en termes de temps d'exécution (i.e. résultats présentés dans la section 8.3.2 Chapitre 8).
- l'efficacité de notre approche de suggestion de dimensions sociales en termes de temps d'exécution (i.e. résultats présentés dans la section 8.3.3 Chapitre 8).
- L'efficacité de notre approche de suggestion des événements. En effet, nous avons validé notre hypothèse de départ en démontrant (i) l'existence d'une cooccurrence entre le contexte spatio-temporel et les événements et (ii) l'importance de la notion de proximité sociale dans les sites de médias sociaux (i.e. résultats présentés dans la section 8.3.4 Chapitre 8).

1.5. Organisation de la thèse

Le manuscrit de la thèse est organisé comme suit:

- Dans le Chapitre 2, nous présentons un état de l'art concernant les technologies et formalismes de représentation des documents multimédias et de liens sociaux. Une attention particulière a été attribuée aux technologies du web sémantique.

⁷ www.flickr.com/

- Dans le Chapitre 3, nous présentons les mesures de similarité sémantique les plus connues.
- Dans le Chapitre 4, nous présentons un état de l’art sur les moyens d’enrichissement des annotations automatiques ou semi-automatiques des documents multimédia.
- Le Chapitre 5 propose un modèle d’annotation ainsi qu’un modèle de profil social. Le chapitre présente aussi notre plateforme d’enrichissement automatique et semi-automatique des annotations.
- Le Chapitre 6 détaille les algorithmes de recommandation sensible au contexte social.
- Le Chapitre 7 présente notre contribution sur l’appariement entre requête et graphe social.
- L’implémentation et l’évaluation des performances du système proposé sont discutées dans le Chapitre 8.
- Le bilan des contributions, la conclusion et les perspectives sont présentés dans le Chapitre 9.

La thèse comporte à sa fin une bibliographie suivie de quatre annexes.

Chapitre 2. Technologies et Formalismes de représentation de documents enrichis

2.1.	INTRODUCTION.....	13
2.2.	DOCUMENT MULTIMEDIA SOCIO-PERSONNEL	13
2.3.	TAGGING, METADONNEES ET ANNOTATIONS	14
2.4.	TECHNOLOGIES DU WEB SEMANTIQUE.....	15
2.4.1.	Les formats RDF et RDFS	16
2.4.2.	Ontologies et le format OWL	18
2.4.3.	Framework Jena	18
2.4.4.	Langage de requête SPARQL	19
2.5.	FORMALISMES DE REPRESENTATION D'ANNOTATIONS	20
2.5.1.	Dublin Core	20
2.5.2.	Information Interchange Model (IIM)	21
2.5.3.	EXIF- EXchangeable Image File.....	21
2.5.4.	MPEG-7.....	22
2.5.5.	L'ontologie DIG35.....	22
2.5.6.	PhotoRDF	23
2.5.7.	VRA Core	24
2.5.8.	XMP	24
2.5.9.	OntoAlbum.....	24
2.5.10.	Discussion	26
2.6.	FORMALISMES DE REPRESENTATION DES LIENS SOCIAUX	26
2.6.1.	CSV et format .dot.....	27
2.6.2.	DyNetML.....	27
2.6.3.	Social Content graph	28
2.6.4.	XFN.....	28
2.6.5.	FOAF et ses extensions.....	29
2.6.6.	Discussion.....	29
2.7.	RESUME.....	30

2.1. Introduction

La représentation de métadonnées qui décrivent un document multimédia est fondamentale surtout quand il s'agit de document non textuel tel que les photos. En effet, les photos ne portent aucune sémantique en elles-mêmes contrairement au texte. Sans métadonnées, les documents multimédias restent inexploitable.

Dans ce chapitre, nous dressons l'état de l'art des documents multimédias, des technologies et des formalismes de représentation. Tout d'abord, nous introduisons dans la section 2.2 un type spécifique de document multimédia : le document multimédia socio-personnel. Ensuite, nous présentons dans la section 2.3 le tagging, les métadonnées et les annotations. La section 2.4, est consacrée à la présentation des technologies du web sémantique comme moyen efficace de représentation des annotations et/ou métadonnées. Nous enchainons par un exposé des formalismes de représentation des annotations dans la section 2.5. Un intérêt spécifique est accordé aux formalismes de représentation des images ou des photos puisqu'ils constituent le cadre d'application de notre travail. Puis, nous détaillons dans la section 2.6 les formalismes existants dans la littérature visant à décrire les liens sociaux entre utilisateurs. Enfin, nous résumons ce chapitre dans la section 2.7.

2.2. Document Multimédia socio-personnel

Le terme document multimédia a été initialement adopté lorsqu'un document numérique regroupait divers contenus multimédias diversifiés (i.e. images, sons, textes, vidéos, animations, flux de texte). Des exemples de ces documents selon cette définition peuvent être les documents *SMIL* (Synchronized Multimedia Integration Language) [10], les documents Macromedia Flash FLA [11] ou les documents HTML5 [12].

Depuis, l'essence de la notion de document multimédia a évolué. Actuellement, elle est adoptée pour décrire tout document qui représente au moins un type de média.

Les utilisateurs créent un document multimédia pour plusieurs finalités. En effet, un document multimédia peut être créé, par exemple, pour mémoriser des événements familiaux ou de soirées entre amis, ou encore il peut être crée en vu de produire un film, ou pour produire un cours interactif ou un documentaire sur un thème de la médecine, etc.

La notion de document multimédia personnel est liée aux documents multimédias produits par les utilisateurs eux-mêmes, sans un objectif, a priori, professionnel [13], [14], [15]. Quant à la notion de document multimédia social, plus particulièrement, photo

sociale, elle a été introduite par Crampes et al. 2009 [16] pour désigner les photos qui sont prises lors d'événements familiaux ou de soirées entre amis et dans lesquelles figurent des individus ou des groupes d'individus.

S'inspirant de ces définitions nous définissons, dans ce qui suit, la notion de document multimédia socio-personnel.

Définition 1. Document multimédia socio-personnel

Le document multimédia socio-personnel (en anglais, *Socio-Personal Multimedia Document* abrégé *SPMD*) est un document multimédia produit pour un usage non professionnel pour mémoriser un événement social (e.g. des soirées entre amis, des excursions en famille) dont les éléments les plus importants sont les individus et qui sera destiné à être partagé par un réseau de connaissance.

Notons que, pour faciliter la gestion de ces types de documents (Annotation, recherche visualisation, etc.) des outils spécifiques ont été développés. La finalité de ces outils dépend étroitement du type du document multimédia mais surtout de l'usage [11].

2.3. Tagging, métadonnées et annotations

Le tagging est la pratique qui consiste à associer des mots-clés (appelé encore tags) librement choisis à un document multimédia. Le tagging constitue donc une forme de classification appelée folksonomie [17]. Une folksonomie est un système de classification collaborative, décentralisée et spontanée, basé sur une indexation effectuée par des non-spécialistes [18].

Une métadonnée est littéralement une donnée sur une donnée. Plus précisément, c'est un ensemble structuré (ou pas) de données décrivant une ressource quelconque [19]. Les méta-données sont particulièrement importantes pour les documents multimédias comme les images et les vidéos qui, sans elles, peuvent demeurer pratiquement inexploitable et difficiles à retrouver.

Une méta-donnée peut être utilisée à des fins diverses (e.g. la description et la recherche des ressources, la gestion de collections de ressources, la préservation des ressources). Ces informations peuvent être évidentes (telles que l'auteur, la date de publication, l'éditeur d'un livre, etc). Elles peuvent aussi être difficiles à spécifier et à gérer tels les avis sur un document ou une portion d'un document.

Plusieurs standards qui concernant les méta-données existent. La section 2.5 a pour objectif de présenter ces standards avec un intérêt particulier accordé à celles qui sont destinées à la description des photos.

Une annotation est une information graphique ou textuelle attachée à un document [20]. C'est une note critique ou explicative accompagnant un document. Desmontils et Jacquin présentent les différentes dimensions d'une annotation suivant son utilisation, le formalisme dans lequel elle est décrite et son rôle.

Dans le cadre du web sémantique, on parle d'annotations sémantiques (dans le sens formel). Dans ce cadre on essaie de représenter le contenu du document par une description formelle. Annoter des documents c'est les décrire par des métadonnées. Les termes annotation et métadonnée sont souvent utilisés indifféremment.

2.4. Technologies du web sémantique

Selon la vision de Berners-Lee, le web sémantique est le web du futur dans lequel les informations possèdent un sens explicite facilitant ainsi aux machines leur traitement et leur intégration [21] et [22]. Le web sémantique offre une infrastructure qui permet une utilisation de connaissances formalisées en plus du contenu informel du web actuel. Selon le W3C, le web sémantique a une architecture en couches (voir Figure 1), qui s'appuie sur une pyramide de langages, dont les couches basses sont les plus stabilisées. Cette architecture repose sur la notion d'URI (Uniform Resource Identifier) qui permet de localiser une ressource sur le web en lui attribuant un identifiant unique. XML [23] et XML-Schema [24] représentent la couche syntaxique à laquelle vient s'ajouter les langages qui permettent d'annoter les ressources (RDF) et de définir le vocabulaire des annotations (RDF-S).

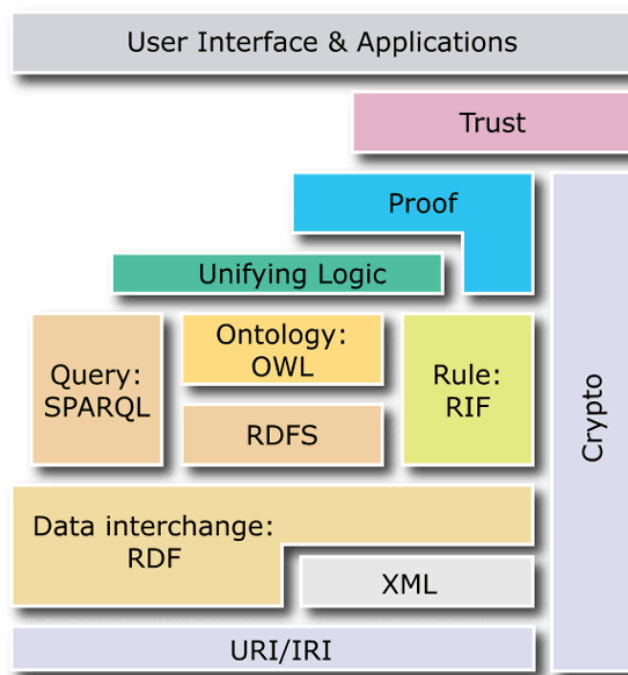


Figure 1. Les différentes couches du web sémantique

La pyramide du W3C inclut aussi la possibilité d'effectuer des déductions en se basant sur un langage standardisé de règles, comme RuleML⁸ (Rule Markup Language) [25], SWRL [26] ou encore Jena Rule Language [27]. A plus long terme la capacité de produire des preuves des déductions faites pourra augmenter le niveau de confiance des utilisateurs dans ces déductions.

2.4.1. Les formats RDF et RDFS

Resource Description Framework (RDF) [28] est une recommandation W3C pour décrire les ressources du WWW. Ces dernières sont représentées sous forme de graphes orientés et étiquetés. La principale entité en RDF est appelée ressource. RDF utilise le principe de l'URI⁹ pour identifier d'une manière unique les ressources. Le langage RDF consiste en des déclarations (*statements*) faites au sujet des ressources. Ces déclarations sont représentées sous la forme de triplets composés d'un sujet, d'un prédicat et d'un objet. Le sujet est la ressource dont la déclaration est faite. Le prédicat est un attribut ou une caractéristique de l'objet. L'objet peut être soit une autre ressource ou une valeur littérale, comme une chaîne ou un entier.

⁸ <http://ruleml.org/>

⁹ Uniform Resource Identifier

Comme exemple de déclaration considérons : "il y a une personne identifiée par <http://www.w3.org/People/EM/contact#me>, dont le nom est Eric Miller, et l'adresse e-mail est em@w3.org, et le titre est le Dr . "

La représentation RDF correspondante à cette déclaration est donnée dans la Figure 2.

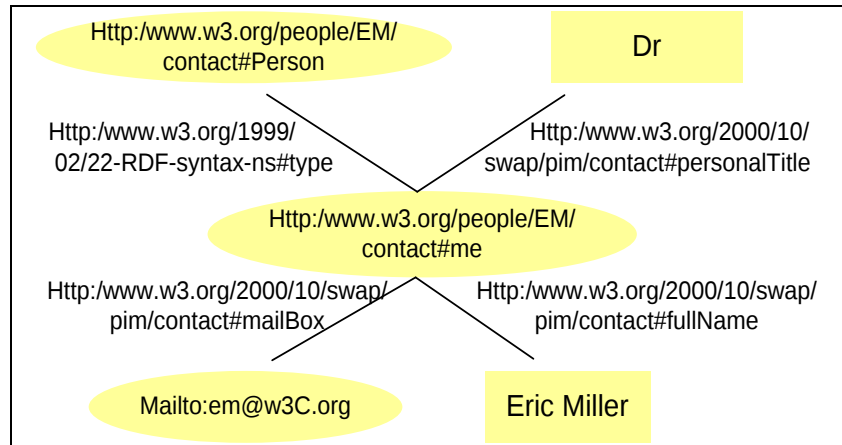


Figure 2. Exemple de graphe RDF

- Individus (en anglais, *individuals*), e.g. Eric Miller, identifiés par <http://www.w3.org/People/EM/contact#me>
- Des spécifications de ressources, e.g. Person, identifiés par <http://www.w3.org/2000/10/swap/pim/contact#Person>
- Propriétés (en anglais, *property*), e.g. mailbox, identifiées par <http://www.w3.org/2000/10/swap/pim/contact#mailbox>
- Valeurs des propriétés (i.e. valeurs littérales), e.g. <mailto:em@w3.org> comme la valeur de la propriété mailbox (RDF utilise aussi des chaînes de caractères telles que "Eric Miller", et des valeurs ayant d'autres types de données comme des entiers et des dates).

RDF utilise des types de données XML-Schema pour indiquer le type de ces valeurs littérales.

De plus, afin d'apporter plus de sémantique à RDF, RDFS (RDF-Schema) a été développé permettant de créer des relations sémantiques entre classes de ressources. Ainsi, RDFS permet aux auteurs de créer des hiérarchies simples en utilisant les ressources *rdfs:Class* et les propriétés *rdfs:subClassOf*. Les ressources peuvent être déclarées comme des instances de *rdfs:Class* via la propriété *rdfs:type*. RDFS permet, également, aux concepteurs de spécifier le domaine de définition (en anglais, *domaine*) et le domaine de

valeur (en anglais, *range*) des propriétés. Enfin, les propriétés en RDFS peuvent être déclarées en tant que sous-propriétés d'autres propriétés via *rdfs:subPropertyOf*.

2.4.2. Ontologies et le format OWL

L'ontologie a été définie formellement, pour la première fois, par Gruber [29] comme étant "*une spécification explicite d'une conceptualisation*". Le terme Ontologie est emprunté de la philosophie, où elle est une description systématique de l'Existence. Dans le domaine de l'intelligence artificielle, ce qui "existe" désigne ce qui peut être représenté. Lorsque les connaissances d'un domaine sont représentées dans un formalisme déclaratif, l'ensemble des objets qui peuvent être représentés est appelé l'univers du discours. Cet ensemble d'objets et les relations qui leurs relient, se représentent par un vocabulaire. Des noms spécifiques tels que des classes, des relations, des fonctions sont associés aux éléments d'un vocabulaire. Des axiomes formels, qui font partie du vocabulaire, permettent l'interprétation et l'utilisation bien formée de ces éléments.

En raison de la puissance expressive limitée de RDFS, les chercheurs dans le domaine du web sémantique ont travaillé sur l'extension de RDFS pour représenter l'ontologie. Les fruits de recherches ont donné naissance au langage d'ontologies du web OWL, qui est devenu un standard de W3C. Il est conçu pour être utilisé par les applications qui doivent traiter le contenu de l'information au lieu de simplement présenter des informations à l'utilisateur. OWL permet une plus grande interopérabilité entre les applications du Web comparée à elle offerte par XML, RDF et RDF Schema (RDF-S) [30] en fournissant un vocabulaire supplémentaire avec une sémantique formelle. OWL est construit au-dessus de RDF / RDFS dans les couches du web sémantiques (voir Figure 1), ce qui signifie qu'une ontologie OWL peut comporter des déclarations faites sur les ressources.

Notons que, parmi les apports du langage OWL: la possibilité d'importation de plusieurs ontologies au sein d'une même, la définition de plusieurs types de propriétés, la définition de cardinalités, la définition des caractéristiques de propriétés, la définition des restrictions, la définition de l'équivalence entre classes et individus.

2.4.3. Framework Jena

Jena est un framework qui permet de lire et écrire des fichiers dans l'un des formats standard de stockage RDF. Il est considéré comme un des plus complet Framework de manipulation de données RDF. En outre, Jena peut stocker et lire des données RDF dans une base de données relationnelles (MySQL, PostgreSQL et Oracle sont supportés). Il supporte la manipulation de données RDF et OWL (y compris les littéraux typés, RDFS et

OWL spécifiques, et vocabulaires normalisés), et intègre un langage de requête RDF : SPARQL. Jena fournit un environnement de programmation pour RDF, RDFS, OWL et SPARQL et comprend un moteur d'inférence à base de règles. Il supporte le raisonnement utilisant des règles implicites ou explicites. Les règles explicites peuvent être exprimées sous plusieurs langages. Jena supporte son propre langage Jena Rule Language mais aussi d'autres comme SWRL et RuleML. Il supporte aussi des moteurs d'inférences non propres à lui comme Pellet. Jena possède des classes d'objets pour représenter des graphes, des ressources, des propriétés, et des littéraux. Les interfaces représentant les ressources, les propriétés et les littéraux sont appelés respectivement *Ressource*, *Property* et *Literal*. Dans Jena, un graphe est appelé *model*. Plus de détails sur l'utilisation du Framework Jena figurent dans l'Annexe IV.

2.4.4. Langage de requête SPARQL

SPARQL est le langage de requête pour les graphes RDF [31]. Il est basé sur l'adéquation de motifs (en anglais, *patterns*) de graphes qui peuvent contenir des motifs triplets, des conjonctions, des disjonctions, et des motifs optionnels. Les motifs triplets sont comme des triplets RDF, mais avec la possibilité de remplacer les sujets, les prédicats ou les objets par des variables.

Par exemple, la requête SPARQL décrite par le Code 1 retourne tous les appareils mobiles ayant un écran VGA et une vitesse de processeur plus de 400Mhz. Il affiche également le nom du propriétaire si cette information est disponible. Le résultat de la requête est donné dans le Tableau 1

```
//SPARQL query formation
PREFIX ns:
<http://www.myexample.com/SocialSphere.owl#>
SELECT ?mdv ?ps ?person
WHERE {
    ?mdevice rdfs:subClassOf ns:Device.
    ?mdevice ns:deviceType ns:Mobile.
    ?mdv rdf:type ?mdevice.
    ?mdv ns:displayType ns:VGA.
    ?mdv ns:processorSpeed ?ps.
    OPTIONAL{?mdv ns:ownedBy ?person.}
    FILTER(?ps>400)
```

Code 1. Exemple de requête en SPARQL

Tableau 1. Exemple de résultat d'une requête SPARQL

mdv	ps	person
PDA001	480	Bob
TabletPC009	700	
SmartPhone023	420	

La combinaison des motifs triplets donnent un graphe motif élémentaire (en anglais, *basic pattern graph*), où une correspondance exacte pour un graphe est nécessaire. La manière de résolution de la correspondance entre un graphe motif élémentaire et un graphe RDF est appelée problème de homomorphisme de sous-graphe. Ce problème est différent de celui de l'isomorphisme de sous-graphe que nous avons abordé dans le Chapitre 7. Les algorithmes qui permettent de résoudre ce problème ne font pas partie de la spécification du standard SPARQL [32].

2.5. Formalismes de représentation d'annotations

Dans cette section nous nous intéressons aux formats et standards existants dans la littérature qui permettent de représenter des métadonnées sur les photos.

2.5.1. Dublin Core

Dublin Core (Dublin Core Metadata Initiative) [33] est une organisation dédiée à la création et la diffusion de vocabulaires destinés à la description des métadonnées sur des ressources numériques y compris les photos. L'objectif des vocabulaires proposés est de faciliter l'échange de documents entre des individus et des organisations et d'améliorer le fonctionnement des méthodes de gestion et de recherche de ces documents.

Parmi les vocabulaires proposés, la norme Dublin Core Metadata Element Set (DCMES) est la plus adoptée. Elle comprend 15 termes dont la sémantique a été établie par un consensus international de professionnels provenant de diverses disciplines telles que la bibliothéconomie et l'informatique. Ces termes décrivent : i) le contenu du document ("description", "format", "relation", "source", "subject", "title"), (ii) la propriété intellectuelle ("contributor", "creator", "publisher", "rights") et (iii) le document numérique lui-même ("date", "identifier", "language", "type"). Plusieurs formats peuvent être utilisés pour la représentation du vocabulaire DCMES, tels que XML, ou même une simple liste d'attributs-valeurs à l'intérieur du document. Grâce à sa simplicité et à sa généralité, le DCMES est fréquemment incorporé à d'autres schémas d'annotation pour la

description de vidéos, textes, images et fichiers audio. Ce vocabulaire possède, par exemple, une large adoption au sein des standards du W3C.

2.5.2. Information Interchange Model (IIM)

IIM (Information Interchange Model) [34] est un vocabulaire de métadonnées proposé par l'IPTC (International Press Telecommunications Council) au début des années 1990 dans l'objectif d'améliorer les échanges d'informations internationaux et d'annoter dans plusieurs types de support multimédias. Néanmoins, ces métadonnées ont principalement été utilisées dans le monde par les journaux, les agences de presse et les photographes. Les métadonnées d'IIM les plus utilisées sont : le nom de l'auteur ou du photographe, des informations sur le copyright et des descriptions textuelles. Des mots-clés et les descriptions textuelles peuvent être ajoutés manuellement au fichier de l'image. Adobe Photoshop et d'autres logiciels offrent des interfaces graphiques de visualisation et de saisie de métadonnées IIM.

2.5.3. EXIF- EXchangeable Image File

Le vocabulaire EXchangeable Image File -EXIF [35] constitue désormais le vocabulaire le plus utilisé par les fabricants d'appareils photos numériques pour la description de métadonnées relatives aux images prises par leurs équipements. Ce vocabulaire est une partie de la recommandation EXIF dont sa dernière spécification est EXIF 2.2. EXIF a été transcrit, récemment, en format RDF/XML. Le vocabulaire EXIF a été proposé par le Japan Electronics and Information Technology Industries Association (JEITA) afin d'encourager l'interopérabilité entre dispositifs de capture et systèmes de gestion de photos. Lors de la prise d'une vue, l'appareil photo numérique ajoute automatiquement à l'en-tête de l'image une description au format EXIF. La description de l'image en utilisant ce vocabulaire concerne (i) des informations basique sur l'image (e.g. le nom, la taille, la date de création) ; (ii) des informations sur l'appareil photo (e.g. modèle, fabricant, résolution maximale) ; et (iii) les réglages utilisés pour prendre la photo (e.g. l'intensité du flash, l'ouverture du diaphragme).

Il est à noter aussi que, EXIF permet d'ajouter, aussi, les coordonnées GPS. Par ailleurs, EXIF est largement exploité par divers outils de gestion multimédias.

2.5.4. MPEG-7

Les vocabulaires EXIF, IIM et Dublin Core s'intéressent très peu à la description du contenu d'un document multimédia. En effet, leur description de contenu est limitée à l'attachement d'un texte ou de quelques mots-clés. Pour palier ces limites, le standard MPEG-7 a été proposé et standardisé en ISO/IEC depuis 2001 dans l'objectif de fournir une description sémantique des ressources multimédias, spécialement les fichiers audio et vidéo [36]. MPEG-7 est le résultat des efforts de standardisation du groupe de travail Moving Pictures Expert Group (MPEG).

À la différence des autres standards MPEG, MPEG-7 est plutôt destiné à l'encodage du contenu des documents audiovisuels, MPEG-7 en vue de faciliter la recherche, la navigation, et le filtrage.

MPEG-7 fournit un vocabulaire de description de contenu qui comprend : des descripteurs (D), des schémas de description (DS) et un langage de définition des descriptions (DDL). Un descripteur représente une caractéristique visuelle de bas niveau (i.e. couler, texture) ou auditive (i.e. timbre, mélodie) d'un document multimédia. Les schémas de description (DS) offrent une représentation d'informations, assez complexe, de plus haut niveau de contenu (objets, événements, interactions entre les objets, etc.). Le vocabulaire MPEG-7 peut être étendu en utilisant le langage DDL qui définit un schéma pour la création de nouveaux descripteurs et de schémas de description.

MPEG-7 permet, également, de segmenter un document en parties qui correspondent à des segments spatiaux, temporels ou spatio-temporels. À chaque segment créé, un ensemble d'annotations peuvent lui être attaché.

Plusieurs outils utilisent MPEG-7 pour l'annotation des vidéos et des images. Citons à titre d'exemple, l'application de bureau Caliph [37] et [38] permettant aux utilisateurs d'annoter des photos en utilisant une large partie du schéma description sémantique de MPEG-7.

2.5.5. L'ontologie DIG35

L'ontologie DIG35 [39], développé par le laboratoire IBBT Multimedia Lab (University of Ghent) dans le cadre du W3C Multimedia Semantics Incubator Group, fournit un schéma OWL (en OWL Full) couvrant la spécification DIG35. Pour la représentation formelle de DIG35, pas d'autres ontologies ont été utilisées. Toutefois, les relations avec

d'autres ontologies tels que Exif, FOAF, etc. est prévue dans des versions futures afin de donner à l'ontologie DIG35 un plus large éventail sémantique.

La spécification DIG35 comprend un ensemble "*standard set of metadata for digital images*" une norme de métadonnées pour les images numériques qui favorise l'interopérabilité et l'extensibilité, ainsi que d'un "*uniform underlying construct to support interoperability of metadata between various digital imaging devices*" une uniforme sous-jacente de construire pour soutenir l'interopérabilité des métadonnées entre les divers appareils d'imagerie numérique. Les métadonnées DIG35 couvrent une description de caractéristiques basiques de l'image, une description de métadonnées relatives à (i) la création de l'image (e.g. la caméra) ; (ii) une description générale du contenu de l'image, en particulier, le qui (who), quoi (what), le quand (when) et le où (where); (iii) historiques de modification de l'image ; (vi) les droits d'auteur; et (v) les types de métadonnées fondamentaux. Il est à noter que la spécification 1.1 du DIG35 est payante.

2.5.6. PhotoRDF

PhotoRDF [40] est une tentative de normaliser un ensemble de catégories et des étiquettes spécifiques aux collections de photos. La norme a été proposée au début de 2002, mais elle n'a pas évolué depuis. La norme fonctionne déjà comme une brique de base pour différentes autres normes qui devrait permettre de décrire les photos en RDF.

Les métadonnées dans PhotoRDF sont divisées en trois schémas différents à savoir : (i) une partie du schéma Dublin Core ; (ii) un schéma technique ; et (iii) un schéma de contenu. Le schéma Dublin Core est adopté pour décrire le créateur, éditeur, titre, date de publication, etc. En ce qui concerne les aspects techniques de la photo la norme PhotoRDF redéfinit une petite partie de métadonnées EXIF. En ce qui concerne la description du contenu, le schéma de contenu définit un ensemble très restreint de mots-clés qui seront utilisés dans le champ "*subject*" du schéma Dublin Core.

PhotoRDF couvre les différents aspects d'une photo qui vont du réglage de l'appareil photo aux objets représentés sur la photo. La norme échoue, cependant, à couvrir la description des aspects fondamentaux de photos. Par exemple, le lieu de prise de vue de la photo n'est pas abordé. Par ailleurs, la description du contenu est limitée à un petit nombre de mots-clés et ne supporte pas le principe du tagging car il n'a pas été prévu au moment de l'élaboration de la norme.

2.5.7. VRA Core

VRA Core (Visual Resources Association) actuellement dans sa version 4.0 [41] est un standard de données pour la communauté du patrimoine culturel qui a été élaboré par le Comité des ressources visuelles de l'Association de normalisation des données.

Le vocabulaire VRA Core se compose d'un ensemble d'éléments de métadonnées (unités d'information comme le titre, lieu, date, etc), ainsi que des instructions permettant d'organiser les documents sous une structure hiérarchique. Les ensembles d'éléments, ainsi produits, permet un accès par catégories aux images et aux documents de l'œuvres de l'art.

La plupart des travaux qui traitement les images des œuvres d'art utilisent ce standard pour la description des métadonnées. Parmi ces travaux nous citons [42] proposant un outil d'annotation sémantique et de recherche dans des collections de l'art se basant sur des métadonnées VRA.

2.5.8. XMP

Adobe System a lancé, vers la fin de l'année 2001, un format baptisé XMP "*Extensible Metadata Platform*" [43]. Ce format a la particularité de : (i) enregistrer les métadonnées décrivant l'image au sein du même fichier sous format JPEG, JPEG 2000, GIF, PNG, HTML, TIFF, Adobe Illustrator, PSD, PostScript, etc. (ii) utilise XML pour représenter les métadonnées. Il est donc aisément extensible sans crée de problèmes de compatibilité. (iii) l'usage de l'unicode qui permet l'intégration de texte en n'importe quelle langue contrairement au format IPTC ou EXIF.

Bien qu'ouvert à tout type de données pouvant intégrer un document XML, XMP prédéfinit la façon de stocker un certain nombre d'informations les plus courantes (e.g. le titre du fichier, l'auteur et l'historique des modifications) en reprenant, en particulier, des éléments de Dublin Core et d'EXIF.

2.5.9. OntoAlbum

OntoAlbum [3] est un système de gestion de photos basé sur deux ontologies : une pour la description d'une photo et une taxonomie des membres d'une famille (père, mère, fils, etc.). La première ontologie est illustrée sur la Figure 3.

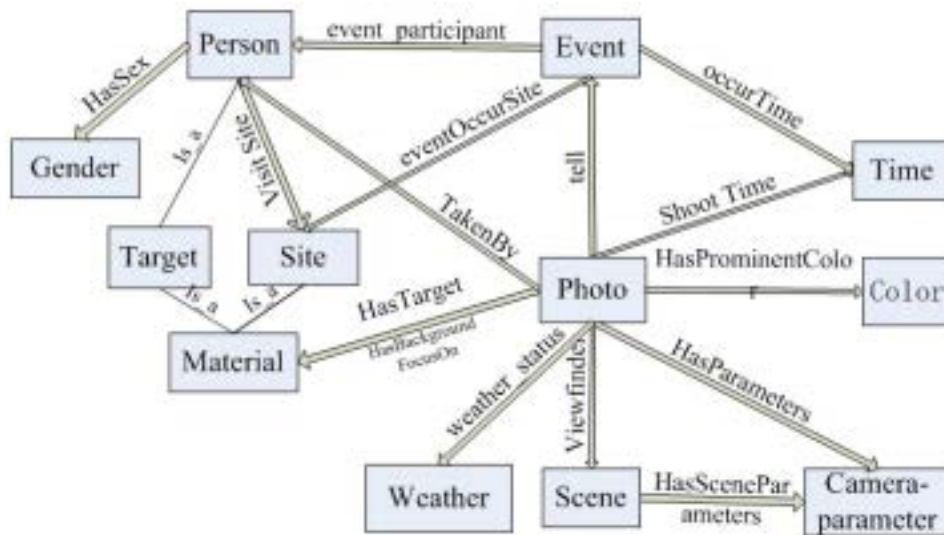


Figure 3. Ontologie d'annotation de photos proposée dans [3]

Les auteurs modélisent une photo comme étant un événement qui peut avoir plusieurs participants. La classe "Person" peut être relié à la taxonomie exprimée dans la Figure 4.

L'outil *OntoAlbum* présente une interface graphique permettant aux utilisateurs de se servir des concepts de deux ontologies dans la mise en œuvre de l'annotation. Par exemple, on peut exprimer manuellement que ce mot-clé "Alice" décrit le nom de mère du propriétaire de la photo et aussi que cette personne est un des participants de la scène capturée. L'arbre hiérarchique des concepts des ontologies est utilisé pour la navigation de la collection.

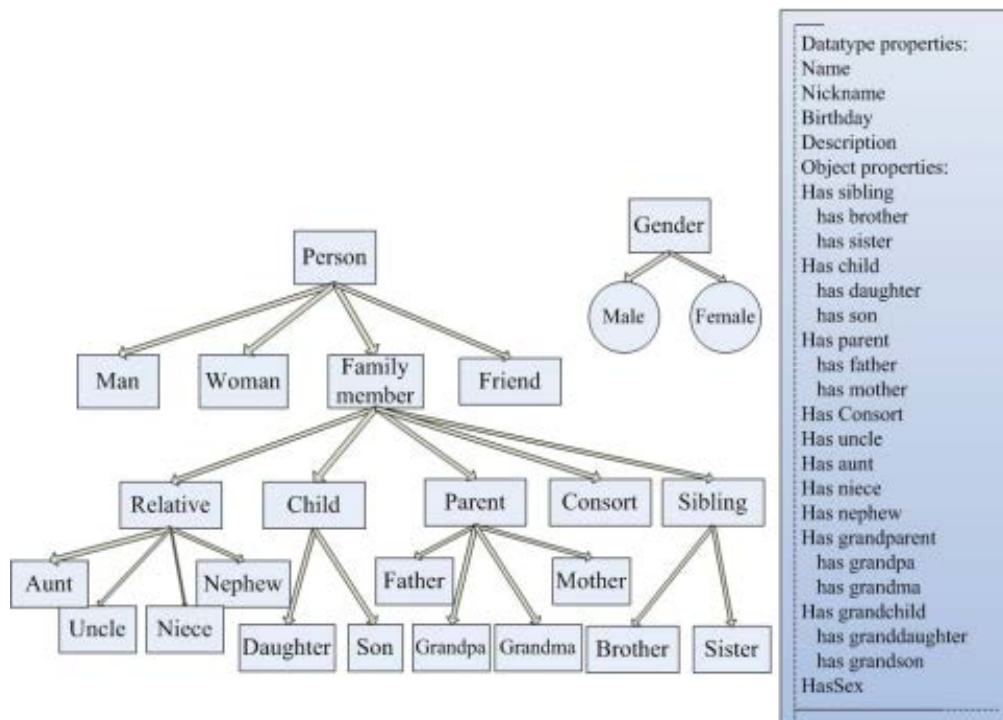


Figure 4. Taxonomie des personnes proposée dans [3]

2.5.10. Discussion

Divers modèles de représentation des métadonnées sont proposés. Les plus simples sont basés sur les mots-clés (tags librement saisis), mais ne sont pas suffisamment expressifs pour garantir de bons résultats de recherche. Des modèles plus structurés et plus expressifs sont proposés ces dernières années basés sur les ontologies dans ses formes variées (allant d'une structure hiérarchique simple, jusqu'aux réseaux sémantiques). L'objectif étant d'introduire plus de sémantique dans la description des ressources du web : documents multimédias (i.e. page web, images, musique, etc.), services (applications, services web, etc.) et utilisateurs (profils utilisateurs, réseau social, etc.) [22]. Avec ces nouveaux modèles, les moteurs de recherches peuvent fournir à l'utilisateur une information plus pertinente et il devient possible d'activer des mécanismes d'inférence afin de dériver de nouvelles informations non initialement explicitées sur les ressources décrites.

Par ailleurs, en vue de garantir l'interopérabilité des systèmes de gestion des documents multimédias, le W3C Media Annotations Working Group a développé des standards de description et de représentation de ressources du web tels que RDF (Ressource Description Framework) et OWL (Ontology Web Langage). Ce sont ces deux standards que nous avons utilisé pour la représentation formelle du modèle d'annotation que nous avons proposé.

Concernant les modèles d'annotation, nous avons constaté que ceux qui ont été proposés dans la littérature, n'offrent quasiment pas de description dédiée au contexte de création des documents multimédias socio-personnels. Le seul travail que nous avons repéré est celui de Viana et al. [44] qui proposent une ontologie baptisée *ContextMultimedia* pour décrire le contexte de prise de vue. La vision du contexte proposée par ces auteurs reste incomplète : elle ne donnent pas aux annotations un aspect personnalisé à l'utilisateur (voir détail dans la section 5.3 du Chapitre 5). De plus, leur modélisation n'est pas centrée sur l'utilisateur mais plutôt sur le document multimédia. Le modèle que nous proposons essaie de pallier ces insuffisances au travers une vision en couche qui favorise l'enrichissement progressif des annotations.

2.6. Formalismes de représentation des liens sociaux

Peu sont les travaux qui ont proposé des formalismes pour représenter les liens sociaux entre utilisateurs d'un réseau social.

Les travaux proposés peuvent être classifiés selon le type de modèle de données utilisé. On distingue particulièrement ceux basés sur les graphes et ceux basés sur le modèle sémantique. Dans la littérature, le modèle des réseaux sociaux le plus répandu est celui basé de graphes où les nœuds représentent les individus (les personnes) et les arêtes (i.e. relations binaires) représentant des relations sociales. Ces relations sont soit déclarées explicitement par des utilisateurs soit déduites implicitement par le système en analysant, par exemple, les degrés de similarité entre utilisateurs. Dans ce qui suit, nous allons présenter les plus importants les formalismes et les standards utilisés pour représentation les liens sociaux.

2.6.1. CSV et format .dot

Le format CSV¹⁰ est l'abrégié de *Comma Separated Values* en français valeurs séparées par des virgules. CSV est un format de fichiers ouvert qui permet de stocker les données d'un tableau. Chaque ligne du fichier correspond à une ligne du tableau ; les colonnes sont en général séparées par des virgules. On peut très bien remplacer les virgules par des tabulations ou tout autre caractère. La plupart des chercheurs en sciences sociales représentent souvent leurs données en utilisant les feuilles de calcul Microsoft Excel, qui peuvent être exportées dans le format CSV (*Comma Separated Values*). Comme ce format n'a pas été proposé spécifiquement pour les structures de graphe, il est simplement un moyen d'exporter un tableau de données. La visualisation graphique des réseaux sociaux est par la suite assurée dans la plupart du temps en moyen de format .dot utilisé par l'Open Source paquet *graphviz* développé chez *AT&T Labs Research*¹¹.

2.6.2. DyNetML

Le DyNetML (Rich Social Network Data Interchange Language) [46] est un formalisme basé sur XML pour stocker un réseau social et représenter les liens sociaux entre les utilisateurs. Il est proposé par Tsvetovat et al. en 2004. La Figure 5 présente les différentes entités de DyNetML. En utilisant ce formalisme, le concepteur d'un réseau social peut spécifier (i) des nœuds de différents types (i.e. "personne", "ressource", "organisation", etc) ; (ii) des groupes de nœuds de même type ; (iii) plusieurs types d'attributs par nœud ;

¹⁰ <http://www.eila.univ-paris-diderot.fr/sysadmin/gestion-docs/formats-ouverts/csv>

¹¹ <http://www.research.att.com/>

(iv) des relations (ou liens) typés ; (v) des attributs de relation ; (vi) plusieurs graphes sur plusieurs fichiers ; et (vi) une entité réseau dynamique (DynamicNetwork sur la Figure 5). Le formalisme proposé est caractérisé par sa flexibilité et sa simplicité.

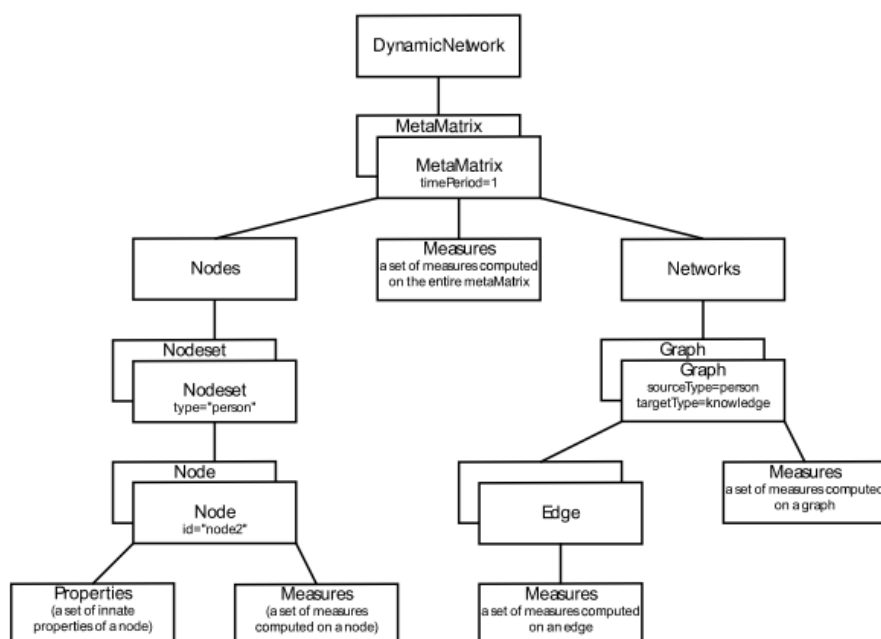


Figure 5. Structure du schéma XML du DyNetML présentée dans [46]

2.6.3. Social Content graph

Social Content graph est un formalisme intégré au sein de l'architecture SocialScope proposée par Amer-Yahia et al. [47]. Ce formalisme permet de représenter des entités des réseaux sociaux (les individus et les groupes) et d'autres informations comme les tags et les documents manipulés par les individus. Le formalisme proposé est complété par une algèbre basée sur l'algèbre relationnelle permettant l'exploitation, la transformation et l'interrogation du réseau représenté par le "Social Content graph".. Malgré l'effort considérable fourni par les auteurs pour formaliser des moyens d'exploitation du réseau social, une instrumentation de moyens intelligents de déduction implicite de liens sociaux est manquante. En effet, l'utilisation d'outils du web sémantique et de logique de description pouvait être envisagée comme solution pour réduire l'effort de l'utilisateur.

2.6.4. XFN

XHTML Friends Network abrégé XFN est un moyen simple de représentation les relations humaines en utilisant les hyperliens. XFN est un micro-format destiné à indiquer les relations existant entre les individus via les liens établis entre leurs sites ou blogs.

Par exemple, pour indiquer dans ma page personnelle que je suis une amie de Alice "*friend*", et je l'ai déjà rencontrée (*meet*), je vais ajouter dans le code suivant du lien les deux attributs :

```
<a href="Alice.example.com" rel="friend met">le site de Alice</a>
```

Malgré que plusieurs services commerciaux comme Amazon s'intéressent aujourd'hui à XFN pour décrire la liste de connaissances de leurs clients, XFN reste un moyen très simple et peu efficace de description d'un réseau social du fait qu'il n'utilise pas un outil d'inférence (tel que les outils du web sémantique) de relation sociale qui réduit l'effort manuel de l'utilisateur à décrire manuellement les relations sociales.

2.6.5. FOAF et ses extensions

Le standard FOAF (Friend Of Friend) [48] est un formalisme basé sur OWL/RDF pour la représentation du réseau social. Bien que beaucoup d'applications du web 2.0 le supportent, il s'avère très rudimentaire au niveau de la représentation de relations sociales. En effet, seule la propriété "*knows*" permet d'exprimer le lien entre les individus. Dans la perspective d'enrichir ces relations, Mika [49] discute plusieurs possibilités de représentation des réseaux sociaux. La solution proposée est basée sur des ontologies et de raisonnements automatisés grâce à la prédéfinition des règles.

L'extension proposée par Mika [49] baptisée FOAF-X repose sur la représentation de l'intensité des liens entre les individus. L'auteur montre la difficulté et la complexité de la représentation de cette notion. Du fait qu'elle : (i) nécessite l'utilisation de la notion de réification dans RDF (ii) une construction complexe qui peut affecter beaucoup d'aspects comme la passion, la fréquence de contact, la mutualité, la confiance, la complémentarité, l'adaptation, l'endettement, la collaboration, la transaction d'investissement, l'attente/l'espérance, le capital social, le voisinage.

2.6.6. Discussion

Plusieurs formats de description d'un réseau social existent. Certains sont génériques et n'ont pas été proposés pour décrire le réseau social spécifiquement mais pour la description des graphes en général. D'autres sont assez pauvres et ne couvrent pas tous les aspects de la notion de relations entre individus.

Nous envisageons utiliser cette représentation pour faciliter l'annotation automatique, la suggestion de tags et donner un plus haut niveau sémantique aux métadonnées.

Nous avons besoin de réduire l'effort fourni par l'utilisateur pour décrire manuellement son réseau social. De ce fait, *foaf* est le mieux placé pour satisfaire cette exigence du fait qu'il se base sur OWL/RDF et peut être adapté pour décrire d'une manière plus précise et intelligente les relations sociales. De plus, afin de satisfaire la même exigence il est nécessaire que le formalisme choisi soit répandu en terme de support dans les outils de butinage (en anglais, *crowling*). Dans cette perspective, XFN et FOAF s'avèrent les mieux placés.

2.7. Résumé

Dans ce chapitre, nous avons présenté des standards et formalismes permettant de décrire les documents multimédias. Nous avons présenté aussi les technologies du web sémantique. Ces technologies proposent une vision prometteuse pour structurer les métadonnées sur les documents multimédias et permettent ainsi de les gérer plus efficacement. Les utilisateurs et leur réseau social faisant partie intégrante de l'enrichissement des documents. Nous avons aussi mis l'accent sur d'autres types de données : les liens sociaux entre utilisateurs.

Nous avons présenté les modèles et les standards de représentation des métadonnées des photos existants comme MPEG-7, Dublin Core, PhotoRDF, VRA Core, EXIF, etc. Ces standards ont pour vocation la description du contenu des images ou la description des caractéristiques de caméras. Aucun de ces modèles ne décrit explicitement la notion du contexte de prise de vue (ou contexte de création). Il nous semble, aussi, que les modèles fondés sur des techniques du web sémantique sont les plus prometteurs du fait qu'elles offrent plus d'expressivité et d'interopérabilité avec d'autres systèmes. Le travail de Viana et al. [44] est le seul travail proposé une modélisation se focalisant sur la caractérisation du contexte de création et agrégant plusieurs ontologies pour le décrire. Néanmoins, la vision du contexte proposée par ces auteurs reste incomplète : elle ne permet pas d'enrichir les annotations d'une manière personnalisée.

Le chapitre suivant présente des mesures de similarités sémantiques entre mots-clés. Ces mesures ont été nécessaires pour notre approche d'appariement entre requête et graphe social décrite dans le Chapitre 7.

Chapitre 3. Mesures de similarité sémantique

3.1.	INTRODUCTION.....	33
3.2.	TERMINOLOGIE ET NOTATIONS.....	33
3.3.	CALCUL DE SIMILARITE PAR LE NOMBRE D'ARCS	35
3.4.	CALCUL DE SIMILARITE PAR LE CONTENU INFORMATIF	37
3.5.	CALCUL DE SIMILARITE HYBRIDE	39
3.6.	RESUME ET DISCUSSION.....	39

3.1. Introduction

Les utilisateurs de systèmes de gestion de documents multimédias socio-personnels emploient une grande variété de tags pour exprimer le même concept. De plus, ils peuvent utiliser les mêmes tags pour exprimer des concepts différents. Le but de manier des thesaurus est d'apporter une solution à ces problèmes et de permettre de faire une correspondance approchée fondée sur la sémantique des concepts à la place des tags.

Dans ce cadre, WordNet peut donc être vu comme un thesaurus, composé de synsets et des relations. La composante atomique sur laquelle repose le système entier est le synset (*synonym set*) : un groupe de mots interchangeables, dénotant un sens ou un usage particulier. Notons S l'ensemble des synsets, et R l'ensemble de relations telles que hyperonymie, hyponymie, antonymie, etc. Chaque synset dénote une acception différente d'un mot, décrite par une courte définition. Une occurrence particulière de ce mot dénotant par exemple le premier sens (le plus courant), dans le contexte d'une phrase ou d'un énoncé, serait ainsi caractérisée par le fait qu'on pourrait remplacer le mot polysémique par l'un ou l'autre des mots du synset sans altérer la signification de l'ensemble.

Plusieurs mesures de similarité sémantique existent dans la littérature, comme nous allons montrer dans ce chapitre, la plupart d'entre elles reposent sur le lien taxonomique "est un" (en anglais, *is-a*) qui peut désigner l'hyperonymie, l'hyponymie ou la spécialisation pour le calcul de similarité. Elles ne considèrent pas, alors, les autres types de relations dans WordNet.

3.2. Terminologie et notations

Les approches présentées ci-dessous manipulent des concepts et des tags. Un concept, noté C , fait référence à un sens particulier d'un tag t donné. Par contre un tag t peut avoir plusieurs sens donc faire référence à plusieurs concepts. Par exemple, si nous avons le tag $t=party$; ce terme fait référence à plusieurs concepts. Selon le concept C_1 (premier sens dans WordNet) *party* fait référence à *political party* qui signifie une organisation voulant atteindre une puissance politique (en anglais, *an organization to gain political power*), selon le concept C_2 (deuxième sens dans WordNet) *party* fait référence à une occasion lors de laquelle les individus se rassemblent pour des interactions sociales et du divertissement (en anglais, *an occasion on which people can assemble for social interaction and*

entertainment). Quant on dit que deux tags sont similaires, c'est dans le sens que les concepts qui leur sont reliés sont similaires.

La Figure 6 illustre une partie du thesaurus WordNet représentant des tags reliés avec des liens "*is-a kind of*" depuis la racine "*entity*" vers les tags "*party*" dans les deux sens : sens 1 (une partie politique) et sens 2 (une occasion lors de laquelle les individus se rassemblent pour des interactions sociales et du divertissement).

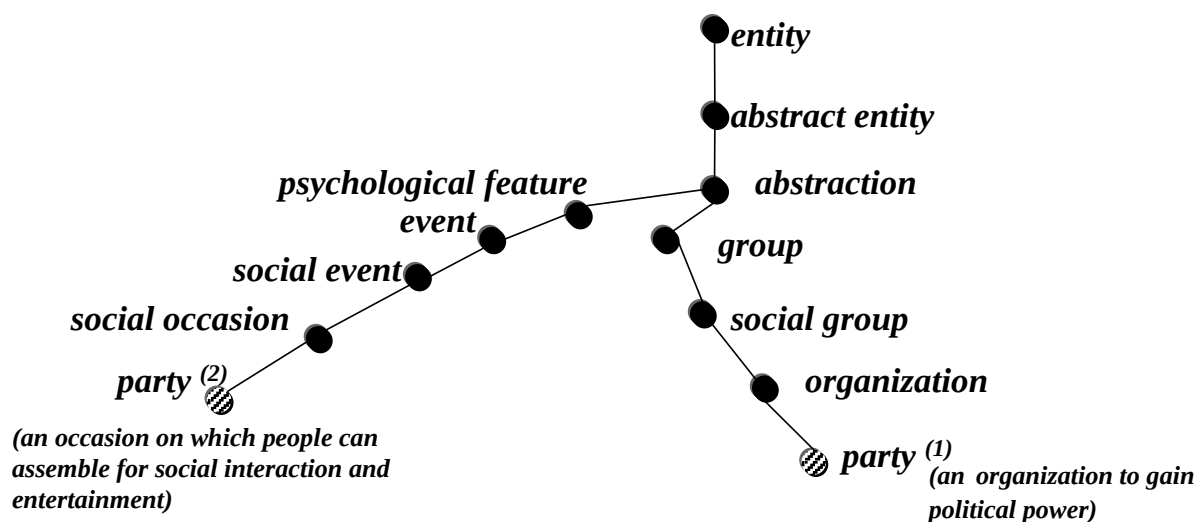


Figure 6. Partie de la structure du thesaurus WordNet représentant le tag 'party' dans deux sens

Il est généralement admis qu'une mesure de similarité est une fonction $sim : S^2 \rightarrow [0,1]$ avec S l'ensemble de concepts. Une mesure de similarité doit être réflexive et symétrique:

- $sim(x,x)=1$: réflexivité
- $sim(x,y)=sim(y,x)$: symétrie

Nous utiliserons les notations suivantes :

–*chemin*(C_1, C_2) est longueur du chemin le plus court en nombre d'arcs qui mène d'un concept C_1 à un concept C_2 .

–*profondeur*(C) est la longueur du plus court chemin entre le concept C et la racine.

–*ppac*(C_1, C_2) est le plus petit ancêtre commun des concepts C_1 et C_2 , c'est le concept le plus spécifique qui les subsume.

–*dppac*(C_1/C_2) est la distance en nombre d'arcs qui sépare C_1 au *ppac*(C_1, C_2).

Dans les formules que nous présentons, nous utilisons la similarité sémantique entre deux tags, t_1 et t_2 , calculée en utilisant la similarité maximale entre les concepts qu'elles présentent:

$$Sim(t_1, t_2) = \max_{C_1 \in S(t_1), C_2 \in S(t_2)} sim(C_1, C_2)$$

où $S(t_i)$ est l'ensemble des concepts qui sont des sens de t_i . La similarité entre deux mots-clés est, donc, égale à la similarité maximale entre les différents sens des mots-clés.

Les mesures de similarité sémantique utilisant des thésaurus peuvent être classées en trois différentes catégories : (i) calcul de similarité par le nombre d'arcs ; (ii) calcul de similarité par le contenu informatif ; (iii) calcul de similarité hybride (contenu informatif et nombre d'arcs).

3.3. Calcul de similarité par le nombre d'arcs

L'idée sous-jacente du calcul de similarité par le nombre d'arcs est de compter le nombre de liens taxonomiques "is-a" entre deux tags. Un des moyens les plus évidents pour évaluer la similarité sémantique dans une taxonomie est de calculer la distance entre les concepts par le chemin le plus court [50]. Dans [50] les auteurs soulignent que cette proposition est valable pour tous les liens de type hiérarchique (est-un "is-a", sorte-de "kind-of", partie-de "part-of") mais doit être adaptée pour d'autres types de liens tels que la cause. Le succès d'une telle approche est dû probablement à la spécificité du domaine comme la médecine qui assure une homogénéisation de la hiérarchie.

Hirst et St-Onge [51] calculent la proximité sémantique qui est une notion plus large que la similarité sémantique. En effet, dans cette mesure, toutes les relations dans WordNet sont prises en compte. Les liens sont classés comme haut (hyperonymie et meronymie), bas (hyponymie et homonymie) et horizontal.

En plus du classement des types de relation, les auteurs définissent un degré de relation :

–Très-fort : (en anglais, *extra-strong*) fondé sur la forme de surface des mots.

–Fort : (en anglais, *strong*) selon trois cas : (i) quand les mots font partie du même synset, (ii) quand les mots ont une relation horizontale entre leurs synsets, ou (iii) quand un mot fait partie d'un mot-composé.

–Moyen-fort : (en anglais, *medium-strong*) quand il y a un chemin acceptable entre deux mots.

–Faible : il s'agit de tous les autres couples de mots.

Les mots qui ont une relation très-forte ou forte ont une similarité constante de $3 * T$ et $2 * T$ respectivement (T étant une constante). Les mots qui ont une relation faible ont une

similarité nulle. Pour les relations moyennes, la proximité est calculée entre les mots par le poids du chemin le plus court qui mène du synset du mot à un autre synset. Elle est calculée en fonction des classification qui indiquent les changements de direction.

$$Prox_{Hirst\ et\ St-Onge}(C_1, C_2) = T - che\ min(C_1, C_2) - k * d$$

T et k sont des constantes. Dans la pratique, les auteurs fixent T à 8 et k à 1 et d le nombre de changements de direction. L'idée est que deux mots sont proches sémantiquement si leurs synsets sont connectés par un chemin qui n'est pas très long et qui ne change pas souvent de direction. S'il n'y a pas de chemin, la proximité est nulle. Cette mesure s'éloigne de la similarité proprement dite car elle traite tout type de relations, comparée aux autres mesures de similarité. Elle ne donne pas, de ce fait, de bons résultats et est très lente en exécution.

Une des limites d'utiliser uniquement le chemin le plus court entre deux concepts lors du calcul de similarité entre eux est que l'on ne prend pas en considération la position des concepts dans l'arborescence du thesaurus. Intuitivement, deux concepts classés en bas du thesaurus sont très spécifiques et sont donc à un degré de granularité plus fin que deux concepts classés en haut du thesaurus.

La mesure de Wu et Palmer [52] apporte une réponse à ce problème en tenant compte de la position des concepts par rapport à la racine du thesaurus.

Wu et Palmer [52] ont défini une mesure de similarité entre concepts pour la traduction automatique entre l'anglais et le mandarin chinois. Pour éviter les problèmes d'ambiguïté, leur mesure s'applique à un domaine conceptuel qui correspond à un point de vue donné pour lequel un mot a un seul sens et correspond donc à un seul concept. La similarité est définie à partir de la distance qui sépare deux concepts par rapport au concept le plus spécifique qui subsume les deux concepts dans le thesaurus, ainsi que la racine de la hiérarchie. La similarité entre C_1 et C_2 est (voir Figure 7):

$$sim_{Wu\ et\ Palmer}(C_1, C_2) = \frac{2 * N3}{N1 + N2 + (2 * N3)}$$

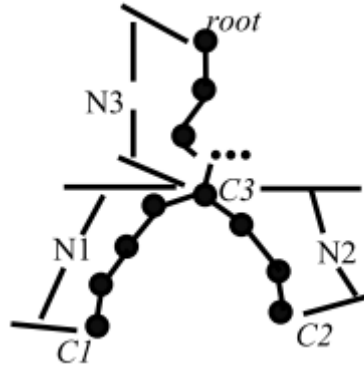


Figure 7. Les relations conceptuelles [52]

Plus formellement cette mesure revient à:

$$sim_{\text{Wu et Palmer}}(C_1 + C_2) = \frac{2 * profondeur(ppac(C_1, C_2))}{profondeur_C(C_1) + profondeur_C(C_2)}$$

Où $profondeur(ppac(C_1, C_2))$ est le nombre minimal d'arcs qui séparent $ppac(C_1, C_2)$ de la racine (root dans la Figure 7) et $profondeur_C(C_1)$ (ou $profondeur_C(C_2)$) est le nombre minimal d'arcs qui séparent C_i de la racine en passant par $ppac(C_1, C_2)$. Deux concepts identiques ont une similarité maximale de 1. Plus les concepts sont éloignés, plus la mesure décroît ; elle atteint 0 pour deux concepts qui ont la racine comme $ppac$.

Cette mesure a l'avantage d'être rapide à calculer, en restant aussi expressive que les autres méthodes (elle a d'ailleurs de bonnes performances que les autres mesures de similarité) [53]. Nous avons vérifié expérimentalement ce propos (voir section 8.3.4 Chapitre 8).

3.4. Calcul de similarité par le contenu informatif

La notion de contenu informatif, dénoté CI , a été introduite pour la première fois par Resnik[54]. Le contenu informatif d'un concept traduit la pertinence d'un concept dans le corpus en tenant compte de la fréquence de l'apparition des mots auxquels il se réfère ainsi que de la fréquence d'apparition des concepts qu'il généralise. Plus précisément le contenu informatif de Resnik se calcule par la formule suivante :

$$CI(C) = -\log(prob(C))$$

Où $prob(C)$ est la probabilité de retrouver qu'un mot du corpus soit une instance du concept C (un des mots référés par le concept C ou par un de ses descendants). Elle est monotone quant on remonte dans la hiérarchie ($prob(A) \leq prob(B)$) si A est au dessus de B dans le thesaurus (e.g. dans la Figure 8), $prob(hill) < prob(geological-formation)$. Dans les expérimentations de Resnik [54], ces probabilités sont calculées par :

$$prob(C) = \frac{frequency(C)}{N},$$

où N est le nombre total de concepts. Plus un concept est général, plus son contenu informatif est faible, ainsi, $CI(prob(racine))=0$ parce qu'il a un concept informatif nul. À l'inverse, plus le concept est spécifique plus son contenu informatif est important (e.g. le contenu informatif de *hill* est plus important que celui de *geological-formation*). L'intuition de la notion de contenu informatif est que la similarité entre deux concepts est la portion d'information qu'ils ont en commun qui, dans le cadre d'un thesaurus, peut être déterminée par le concept le plus spécifique qui le subsume (*ppac*). Cette intuition est indirectement appliquée par les mesures présentées dans la section précédente qui calculent la similarité par la distance en nombre d'arcs qui séparent les deux concepts.

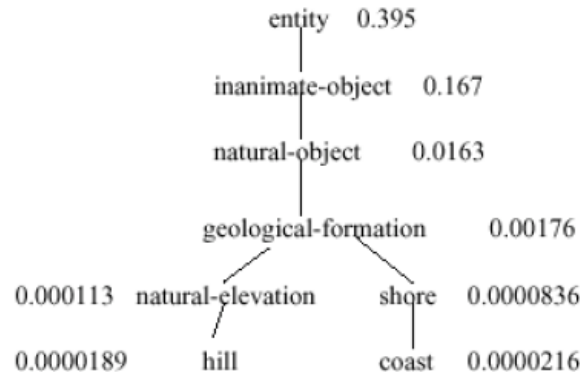


Figure 8. Extrait de WordNet présenté dans [53] avec les probabilités correspondantes

Resnik définit la similarité sémantique entre deux concepts par la quantité d'information qu'ils partagent. Cette information partagée est évaluée numériquement par le contenu informatif du plus petit ancêtre commun (*ppac*).

$$Sim_{Resnik}(C_1, C_2) = CI(ppac(C_1, C_2))$$

Ainsi, si deux termes sont très éloignés et ont comme *ppac* la racine, leur similarité est égale à 0. L'approche de Resnik essaie d'éviter le problème de granularité, sus-indiqué, en diminuant le rôle des arcs dans le calcul de similarité. En effet, les arcs ne sont utilisés que pour retrouver le *ppac*, elle est, de ce fait, un peu sommaire car nous pouvons avoir $ppac(hill, shore) = ppac(shore, natural - elevation)$ même si *shore* et *natural-elevation* sont plus proches de leur *ppac* (*geological-formation*) que *shore* et *hill*.

3.5. Calcul de similarité hybride

Nous entendons par calcul de similarité hybride un calcul combinant le contenu informatif et le nombre d'arcs. Dans ce contexte, nous pouvons citer la mesure de Jiang et Conrath [55] qui pallie les limites de la mesure de Resnik en combinant le contenu informatif du *ppac* à ceux des concepts. Elle prend en considération, aussi, le nombre d'arcs en calculant le contenu informatif de chaque concept.

La mesure de similarité revient, donc, à calculer :

$$Sim_{jiang\ et\ Conrath}(C_1, C_2) = \frac{1}{CI(C_1) + CI(C_2) - (2 \cdot CI(ppac(C_1, C_2)))}$$

Lin [53] propose une mesure de similarité qui calcule la proportion d'information commune entre deux concepts par rapport à leur description.

$$Sim_{Lim}(C_1, C_2) = \frac{2 * CI(ppac(C_1, C_2))}{CI(C_1) + CI(C_2)}$$

Notons que si la probabilité $prob(C)$ ¹² est la même pour toutes les paires de concepts, cette mesure coïncide avec celle de Wu et Palmer.

3.6. Résumé et Discussion

Le calcul de la similarité par nombre d'arcs avec une restriction sur le lien "is-a" pose deux problèmes :

- La granularité : les concepts placés en bas du thesaurus, donc plus spécifiques, sont plus similaires que les concepts en haut du thesaurus.
- La densité : pour un concept qui a plusieurs fils, il paraît logique de supposer que sa similarité avec ses fils est plus grande qu'avec les autres concepts. Ainsi, dans les parties plus denses du thesaurus, la similarité doit augmenter.

Ces problèmes peuvent être résolus en associant des poids aux nœuds. L'affectation de ces poids peut être basée sur : les types de liens présents qui est le même si nous ne considérons que le lien "is-a", la profondeur du nœud dans la taxonomie et la densité du concept par ses voisins immédiats [56].

¹² Rappelons que $CI(C) = -\log(prob(C))$

La performance des mesures hybrides ainsi que celles fondées sur le contenu informatif dépend de la qualité du corpus sur lequel sont calculés les contenus informatifs. Il faut disposer d'un corpus qui contient des occurrences des mots qui vont servir pour le calcul de similarité avec une répartition qui donnerait une idée de la relation entre les concepts. C'est la raison pour laquelle nous avons préféré l'utilisation d'une mesure qui prend en compte le nombre d'arcs.

Chapitre 4. Enrichissement des annotations

4.1.	INTRODUCTION.....	43
4.2.	ANNOTATIONS BASEES SUR LES RESSOURCES SEMANTIQUES.....	44
4.2.1.	<i>RisuPicWeb</i>	44
4.2.2.	<i>Lim et al. 2003</i>	45
4.2.3.	<i>Hollink et al. 2003</i>	46
4.2.4.	<i>Shevade et al. 2005</i>	46
4.2.5.	<i>Discussion</i>	47
4.3.	ANNOTATIONS BASEES SUR UNE COLLABORATION A L'ANNOTATION MANUELLE.....	47
4.4.	ANNOTATIONS BASEES SUR DES DONNEES PROVENANT D'UN RESEAU SOCIAL	47
4.4.1.	<i>Shevade et al. 2007</i>	48
4.4.2.	<i>Zunjarward et al. 2007</i>	48
4.4.3.	<i>Stone et al 2008</i>	48
4.4.4.	<i>Elliot et al. 2009</i>	48
4.4.5.	<i>Discussion</i>	49
4.5.	SUGGESTION DE TAGS LIBREMENT SAISIS	49
4.5.1.	<i>MiAlbum</i>	49
4.5.2.	<i>Sigurbjörnsson et Zwol 2008</i>	50
4.5.3.	<i>Kucuktunc et al. 2008</i>	50
4.5.4.	<i>Lvanov et al. 2010</i>	50
4.5.5.	<i>Discussion</i>	51
4.6.	ANNOTATIONS BASEES SUR LA CAPTURE DU CONTEXTE PHYSIQUE	52
4.6.1.	<i>Enrichissement de contexte physique par des ressources génériques</i>	52
4.6.2.	<i>Enrichissement du contexte physique par des ressources spécifiques à l'utilisateur</i>	56
4.6.3.	<i>Discussion</i>	57
4.7.	RESUME.....	58

4.1. Introduction

Plusieurs travaux de recherche ont été menés afin d'améliorer le processus d'annotation des photos personnelles. D'un côté, des techniques visent à extraire les métadonnées les plus pertinentes depuis le contenu de l'image en utilisant une extraction des couleurs/textures, identification des objets, reconnaissance de visage, catégorisation basée sur le contenu, etc. En 2000, Smeulders et al. [57] publient un état de l'art assez riche décrivant ces techniques. Notons que le propos de cette thèse exclut les techniques d'annotation par contenu.

En général, les techniques d'annotation de photos peuvent être classifiées en manuelle, semi-automatique et automatique. Le Tableau 2 donne une comparaison de différentes catégories d'annotation en termes d'effort humain fourni et de l'assistance de la machine.

Tableau 2. Les techniques d'annotations

	Effort humain	Assistance fournie par le système
Manuelle	Ajouter des annotations structurées avec une information sémantiquement assez pertinente pour l'utiliser lors de la recherche.	Stocker les annotations dans une base de données.
Semi-automatique	Ajouter des annotations structurées ou sous formes de texte libre ou mots-clés. Interagir avec le système dans un processus itératif.	Analyser les annotations et en extraire des informations sémantiques.
Automatique	Vérifier la consistance des annotations et y faire des corrections si nécessaire.	Ajouter des annotations automatiquement en utilisant des capteurs de contexte physiques ou des techniques de reconnaissance de formes.

Dans le présent chapitre, nous décrivons des techniques d'annotations automatique et semi-automatique en mettant l'accent sur l'assistance fournie par le système et sur les ressources d'annotation. La classification des travaux portant sur l'annotation de documents multimédias s'annonce difficile car les travaux proposés combinent plusieurs méthodes, techniques et manières pour enrichir les annotations. En général, nous pouvons distinguer différentes manières pour enrichir les annotations. En effet, certains travaux enrichissent les annotations en se basant sur les ressources sémantiques. D'autres les

enrichissent en se basant sur une collaboration lors d’une annotation manuelle ou se basant sur une similarité par le contenu. Certaines impliquent les données provenant du réseau social pour les enrichir. Des travaux se basent également sur la capture du contexte de prise de vue pour enrichir les annotations.

Ce chapitre est organisé comme suit : la section 4.2 présente les travaux d’enrichissement d’annotations basés sur les ressources sémantiques. La section 4.3 expose les travaux basés sur la collaboration à l’annotation manuelle. Dans la section 4.4, nous exposons les travaux d’enrichissement d’annotation basés sur des données provenant du réseau social. Les travaux de suggestion d’annotations dans le cadre de tags librement saisis sont, quant à eux, détaillés dans la section 4.5. Dans la section 4.6, nous présentons les travaux qui partent de la capture du contexte physique pour enrichir les annotations. Enfin, un résumé du chapitre est présenté dans la section 4.7.

4.2. Annotations basées sur les ressources sémantiques

Ce type d’approche utilise des ressources sémantiques externes qui peuvent être génériques comme WordNet [58] ou spécifiques à un domaine donné comme le domaine de l’art AAT (Art and Architecture Thesaurus) [59]. Le recours à l’utilisation de ces ressources externes est fait dans l’objectif de faciliter la saisie manuelle des annotations ou d’enrichir automatiquement les annotations. Dans ce qui suit, nous allons présenter les travaux les plus connus basés sur les ressources sémantiques.

4.2.1. RisuPicWeb

Elliot et al. [60] proposent l’utilisation des fonctions de comparaison multicritères et de thesaurus dans le calcul de la mesure de similarité entre concepts.

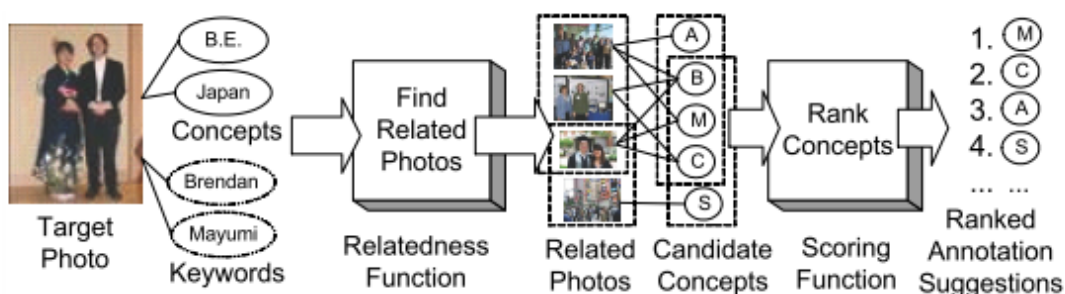


Figure 9. Framework général de suggestion d’annotation de photos proposé par [60]

La Figure 9 illustre le processus de suggestion d’annotation de photos proposé par Elliot [60]. Ce processus part de la construction d’une matrice qui représente l’existence ou

l'inexistence de concepts dans chaque description de photos. La considération des concepts est basée sur le thésaurus WordNet. Aucun détail n'a été fourni à propos du processus d'extraction de concepts à partir de la description textuelle des photos. De plus, la matrice est d'une largeur indéterminée. Pour minimiser le nombre de concepts dans la matrice, les auteurs utilisent une technique appelée Latent Semantic Indexing [61].

La fonction multicritère prend en considération toutes les photos dans la base y compris celles qui appartiennent au même album que la photo en cours d'annotation (i.e. Target Photo sur la Figure 9) et celles qui partagent des concepts communs avec celle en cours d'annotation.

Pour trier les résultats, les auteurs utilisent le cosinus entre chaque concept de la photo en cours d'annotation et chaque photo appartenant aux critères en question. Pour dégager la similarité entre deux concepts, ils calculent la somme des similarités des photos qui contiennent le deuxième concept.

4.2.2. Lim et al. 2003

Lim et al. [62] présentent un processus d'annotation et de recherche de photos illustré par la Figure 10. L'approche repose sur une taxonomie prédéfinie d'événements permettant de grouper les images déjà annotées. Le regroupement est basé sur des descripteurs appelés mots-clés visuels déterminés à partir de l'application d'une technique d'apprentissage. Le *mot-clé visuel* est une technique qui permet de décrire un ensemble de mots-clés à partir d'un ensemble d'images exemples. Lorsqu'une photo non annotée se présente un calcul de similarité par contenu, entre elles et les photos de chaque classe, est effectué. À l'issue de ce calcul sont déterminés les mots-clés les plus proches.

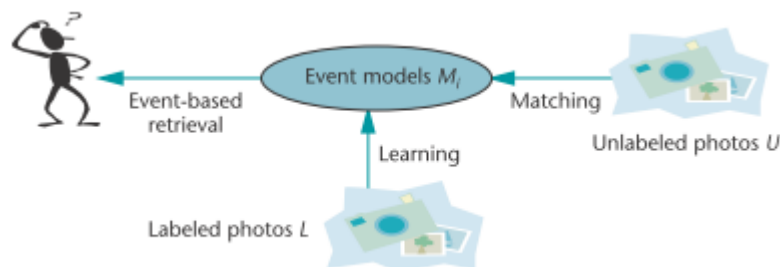


Figure 10. Processus d'annotation et de recherche de photos basé sur l'apprentissage [62]

4.2.3. Hollink et al. 2003

Hollink et al. [42] présentent un outil d'annotation et de recherche dans les collections d'images d'art. L'outil proposé permet à l'utilisateur de sélectionner des termes appartenant à divers thesaurus existants comme AAT (*Art and Architecture Thesaurus*) [59], WordNet, ULAN (*Union List of Artist Names*) [63] et Icon-class (*An iconographic classification system*) [64].

Comme illustré dans la Figure 11, l'outil se base sur une transcription des métadonnées VRA en RDF schéma. L'utilisateur pourra renseigner les métadonnées des images de l'art en se servant de divers thesaurus. Des liens entre les différents thesaurus ont été créés comme les liens d'équivalence, de généralisation/spécialisation grâce à l'utilisation d'OWL [65]. Ces liens permettent l'interopérabilité entre les différentes ontologies utilisées.

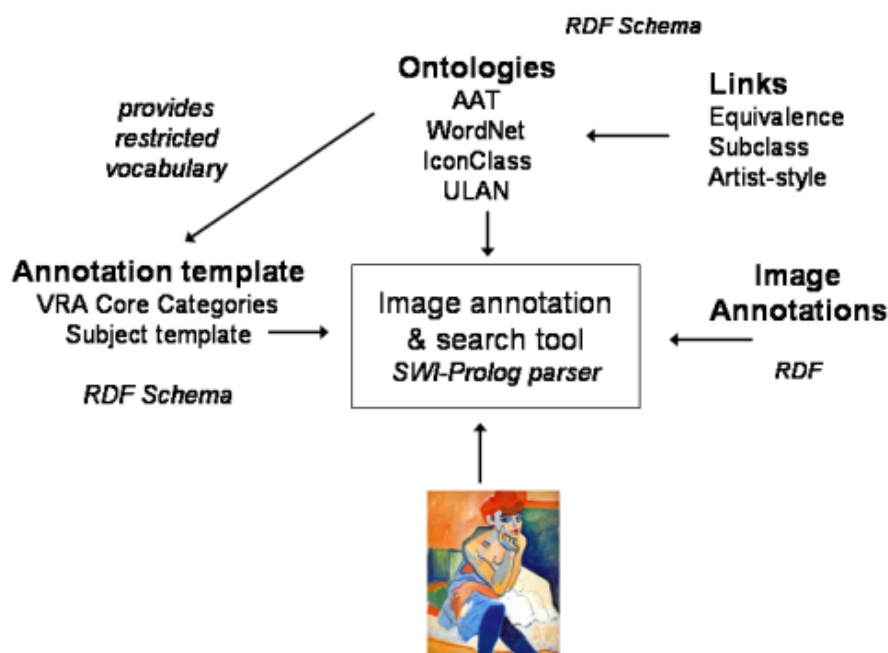


Figure 11. Processus d'annotation et de recherche des photos de l'art proposé dans [42]

4.2.4. Shevade et al. 2005

Le travail de Shevade et al. [66] combine des techniques de comparaison d'images par contenu avec des similarités de tags en utilisant des ressources sémantiques comme le thesaurus WordNet [58] et ConceptNet [67] pour proposer des suggestions. Pour effectuer la comparaison par contenu, les images sont regroupées en classes (en anglais, clusters) basé sur l'annotation partagée.

4.2.5. Discussion

L'implication d'une assistance à base de ressources sémantiques qui n'a pas recours à une annotation automatique basique nécessite une participation massive de la part des utilisateurs et consomme énormément leurs temps et leurs énergies. De plus, ce genre d'approches ne peut pas s'appliquer seul dans le cadre d'annotation de documents multimédias socio-personnels. Les documents multimédias socio-personnels nécessitent des informations sur la vie privée de la personne concernée. L'utilisation de ressources sémantiques de nature spécifique telles que AAT ne permet pas à lui tout seul, de fournir le taux d'assistance suffisante et nécessaire à l'annotation des documents multimédias socio-personnels.

4.3. Annotations basées sur une collaboration à l'annotation manuelle

Ces annotations sont apparues à travers des travaux qui partent de l'idée de déterminer une annotation 'idéale' pour une photo. L'idée est de fournir des recommandations d'annotations/tags basées sur des collaborations entre utilisateurs [68], [69] et [70]. Dans le projet ESP [68], les auteurs mettent en place un jeu en ligne inventif dans lequel les utilisateurs jouent les uns contre les autres. Le jeu consiste à libeller des images. À la fin d'une partie, le système stocke seulement les tags communs de part et d'autres des utilisateurs participants au jeu. Ces tags stockés pour une image serviront ultérieurement à l'identifier lors d'une recherche par mots-clés. Dans la même perspective, Google a mis en place Google Image Labeler [69] un jeu qui consiste à déterminer le sens commun des images pour deux joueurs distants. Ce type de travaux néglige que le sens commun entre les utilisateurs est, à notre avis, fortement lié à un aspect social et une connaissance préalable.

4.4. Annotations basées sur des données provenant d'un réseau social

Un nombre limité de travaux utilise des données provenant d'un réseau social pour améliorer la suggestion des annotations. Pour ce type d'approche, nous pouvons citer les travaux de Shevade et al. 2007 [71], de Zunjarward et al. 2007 [72], de Stone et al. 2008 [73] et enfin de Elliot et al. 2009 [60] que nous détaillons dans la suite.

4.4.1. Shevade et al. 2007

Shevade et al. 2007 [71] combinent des mesures de similarité entre utilisateurs. Ces mesures comprennent des proximités entre utilisateurs dans le réseau social, des similarités sémantiques entre concepts utilisant ConceptNet, et des similarités entre événements. Les annotations sont générées en utilisant : (i) les annotations de l'utilisateur le plus similaire dans le réseau social ; (ii) la similarité des images basées sur leur contenu ; et (iii) l'application de la propagation d'activation (en anglais, "*activation spreading*") [74] sur le graphe les concepts les plus similaires. La propagation d'activation est un processus de recherche initiée par l'étiquetage d'un ensemble de nœuds source avec des poids. La propagation aux autres nœuds est itérative, elle prend en considération la relation entre les nœuds.

4.4.2. Zunjarward et al. 2007

Zunjarward et al. 2007 [72] ont repris le travail de Shevade et al. 2007 [71]. L'apport de Zunjarward et al. se résume en ce qu'ils appellent réseau social de confiance. Cette approche se base sur deux mesures. La première est fondée sur une valeur binaire attribuée manuellement par l'utilisateur à chaque personne membre de son réseau social. Cette valeur binaire correspond à une valeur de confiance. La deuxième mesure se base sur la cooccurrence de l'utilisateur avec d'autres dans le même événement. Les événements sont des tags attribués par les utilisateurs. Une valeur de confiance est accordée à un utilisateur, s'il a annoté des photos avec un événement qui correspond à l'évènement en cours d'annotation.

4.4.3. Stone et al 2008

Stone et al. 2008 [73] identifient des personnes sur Facebook en utilisant la reconnaissance et la similarité de visages par contenu. Pour suggérer les noms d'utilisateurs figurants, ils estiment statistiquement, l'intensité de la relation inter-utilisateurs. Pour ce faire, cette méthode prend en considération deux métriques qui sont : le nombre de photos d'un utilisateur u sur lesquelles une personne est identifiée ou dans les photos de ses amis et celle de le nombre de photos où l'utilisateur u est présent ensemble avec cette personne.

4.4.4. Elliot et al. 2009

Dans une seconde approche, implémentée dans le système RisuPicWeb, présenté dans la section 4.2.1, Elliot et al. 2009 [60] prennent en considération une proximité sociale. La

mesure est simple et se base sur le nombre de liens qui relient les propriétaires des photos. La fonction multicritère de l'approche présentée dans la section 4.2.1 est alors, redéfinie pour prendre en considération la proximité sociale.

4.4.5. Discussion

Certes, l'implication des informations provenant du réseau social dans le processus de suggestion d'annotations des documents multimédias socio-personnels est une direction très prometteuse.

Malgré l'originalité des approches proposées par Shevade et al. [71] et Zunjarward et al. [72], les expérimentations présentées ne semblent pas indiquer clairement l'efficacité de leurs méthodes. De plus, aucune structure de réseau social n'a été proposée ou utilisée. Par ailleurs, le calcul de la similarité entre utilisateurs est très basique et ne dépend pas du contexte.

Dans Stone et al. [73], les métriques proposées, pour estimer l'identité de la personne dans la photo, semblent raisonnables. Cependant, elles peuvent être améliorées en considérant l'événement ou bien le contexte de cooccurrence des personnes dans une même photo. L'approche se base sur des techniques de reconnaissance de visages qui nécessitent d'excellentes conditions de lumière et d'exposition, ce qui n'est pas toujours le cas des documents multimédias socio-personnels.

4.5. Suggestion de tags librement saisis

Ce genre d'approche ne considère pas la sémantique des relations entre tags. Les tags qui annotent les photos sont une liste de mots-clés librement saisis. Nous présentons différentes techniques de recommandation de tags proposées dans la littérature.

4.5.1. MiAlbum

Le système MiAlbum [75] implémente une approche de réinjection de pertinence (en anglais, "*relevance Feedback*") pour améliorer la recherche des images. Grâce à un avis de pertinence donné par l'utilisateur lors de sa recherche, le système procède à une augmentation d'indice de pertinence des mots-clés communs entre les annotations de l'image, jugée pertinente, et la requête. Si les mots-clés de la requête ne se trouvent pas initialement dans les annotations, le système les ajoute dans l'annotation avec un indice de pertinence fort. Dans le cas où l'utilisateur juge une image résultat non pertinente, le

système procède à une discrimination de l'indice de pertinence des mots-clés communs entre la requête et les mots-clés qui décrivent l'image en question.

4.5.2. Sigurbjörnsson et Zwol 2008

Sigurbjörnsson et Zwol 2008 [76] présentent une approche de recommandation de tags pour assister la tâche de tagging dans le service de médias sociaux Flickr. Cette méthode capture l'intelligence collective en utilisant la cooccurrence du tag saisi par l'utilisateur et ceux qui apparaissent dans d'autres photos. La liste des candidats tags obtenue est triée suivant un facteur qui prend en compte la fréquence des tags dans la base. La méthode de recommandation ignore la similarité des photos lors du choix de cooccurrence ainsi que la proximité sociale entre utilisateurs.

4.5.3. Kucuktunc et al. 2008

Kucuktunc et al. 2008 [77] présentent un système de suggestion de tags baptisé Tag Suggestr. Ce système a été expérimenté sur une collection de test Flickr. Le processus de recommandation est semblable à celui de Sigurbjörnsson et Zwol avec une différence au niveau de la comparaison des images par contenu lors du filtrage des tags. En effet, un indice de pertinence fort est attribué aux tags qui décrivent les images les plus similaires.

4.5.4. Lvanov et al. 2010

Lvanov et al. 2010 [78] détaillent une approche de suggestion de tags pour assister l'utilisateur à annoter les photos. Le framework d'annotations est illustré par la Figure 12. Dans cette approche, une étape de prétraitement est nécessaire afin d'extraire les caractéristiques bas-niveaux et de classifier les photos annotées issues de Flickr. Le processus (en ligne) part de la sélection d'une région d'intérêt. Une comparaison au niveau du contenu entre des images, déjà annotées et classifiées, et la région d'intérêt sélectionnée par l'utilisateur est effectuée afin de proposer les images les plus similaires avec leurs annotations. Après validation par l'utilisateur, un processus dit de propagation de tags est effectué afin d'annoter par les tags choisis par l'utilisateur d'une photo non annotée, similaire à la région d'intérêt.

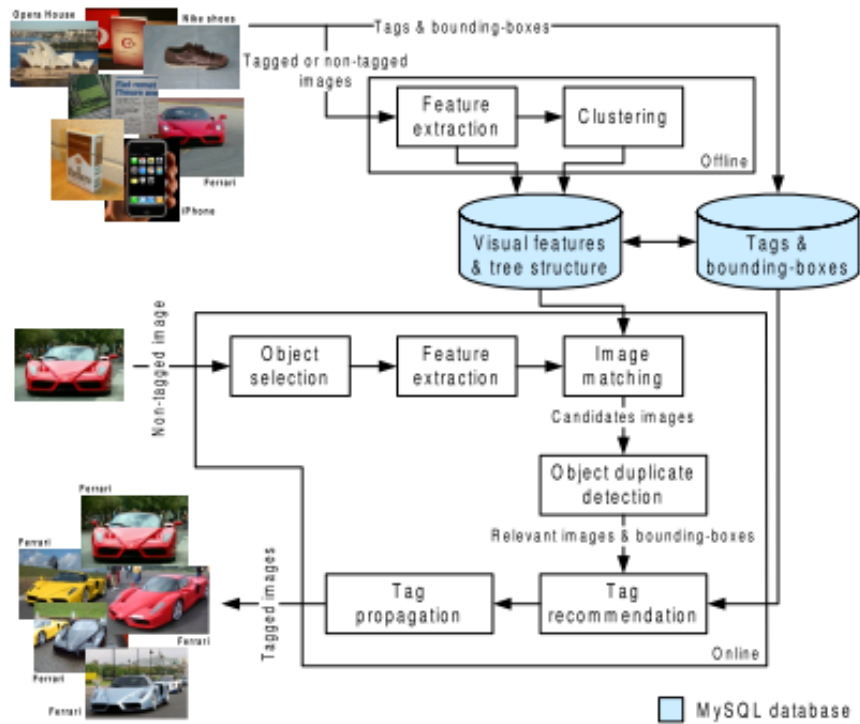


Figure 12. Framework d'annotation décrit dans [78]

4.5.5. Discussion

L'inconvénient majeur de ces approches exposées est qu'elles ne considèrent pas la sémantique des tags et les relations entre tags. La sémantique des tags est très importante dans les documents multimédias socio-personnels. Par exemple, on pourrait décrire "May" comme une personne ou comme un indicateur temporel ou encore en étant associé à un autre tag comme "home" pour parler de la maison de "May". Dans ce cas là, on parlerait d'un indicateur spatial. De plus, dans l'approche d'Ivanov et al. [78] et Kucuktunc et al. [77], la liste des tags proposée est, parfois, hétérogène. Dans l'exemple proposé dans [78], le système propose "Lausanne" et "Paris" comme suggestions possibles. Une approche qui considère une similarité d'un élément du contexte (localisation) pourrait résoudre ce conflit en gardant seulement les tags qui sont similaires par contexte à la photo en cours d'annotation. De plus, l'utilisation exclusive de la similarité par contenu pour propager des annotations s'avère un processus très confus. Par exemple, une brique peut ressembler par contenu à une texture de tapis tissée en paille. Cependant, sémantiquement elle en est très loin. Par ailleurs, à cause des conditions de prise de vue (i.e. zoom, condition de lumière), on peut prendre deux photos qui désignent la même chose, les approches par contenu ne permettront pas dans ce cas de détecter cette similarité.

4.6. Annotations basées sur la capture du contexte physique

Nombreux sont les travaux qui s'appuient sur la notion de contexte dans le domaine de la gestion des documents multimédias socio-personnels. Cependant, ce qu'ils appellent contexte n'est pas formellement défini. En le considérant, on y fait référence à plusieurs aspects hétérogènes. Nous pourrions classer les approches de ce genre par type de ressources utilisés pour enrichir le contexte physique capturé : ces ressources peuvent être génériques comme des services web dédiés à l'enrichissement des coordonnées géographiques ou spécifiques à l'annotateur comme son agenda personnel pour en déduire des annotations qui coïncident avec le temps de la prise de vue.

4.6.1. Enrichissement de contexte physique par des ressources impersonnelles

Plusieurs efforts de recherche ont été menés à fournir une approche automatique ou semi-automatique d'annotation de documents multimédias socio-personnels partant de la capture du contexte physique depuis l'environnement ambiant de l'utilisateur et l'enrichir via des services web impersonnelles. En effet, actuellement les utilisateurs s'intéressent à prendre des photos à l'aide de leurs dispositifs mobiles. Ces dispositifs sont équipés de senseurs comme Bluetooth [79], GPS [80], Cell-ID [81], etc. Ces données dites contextuelles, capturées sont interprétées pour déduire des informations comme les dispositifs autour de la prise de vue, la localisation ou le temps de la prise de vue nécessaires pour annoter automatiquement les photos.

Parmi ces travaux, nous pouvons citer CONFOTO [82] un système doté de services d'annotation et de navigation sémantique pour les photos de conférences.

PhotoCompas [83] utilise l'instant de prise de vue et la localisation via GPS pour interpréter des métadonnées contextuelles de la meilleure façon qu'il soit. En reliant ces photos à des photos précédemment capturées et annotées, le système suggère l'identité des personnes qui y figurent. Les auteurs utilisent des services web externes afin d'enrichir des informations basiques (latitude, longitude et altitude et "*timestamp*"(instant de prise de vue)) avec des données plus riches et significatives pour l'utilisateur. Ces données comprennent les conditions météorologiques, l'état de la lumière (e.g. nuit, jour), la saison (i.e. automne, hiver, printemps et été). Dans ce travail, les auteurs de PhotoCompas

argumentent leur approche par une étude sur un échantillon représentatif des utilisateurs des photos sociales. L'expérimentation a prouvé l'utilité des informations telles que la saison, les conditions de lumière, le moment de la journée comme moyens pour exploiter ultérieurement ces photos. L'expérience a prouvé que les métadonnées les plus efficaces pour se souvenir des photos sociales sont : celle qui indique si la photo est prise à l'intérieur ou à l'extérieur, sa localisation, l'évènement qu'elle représente, le nombre de personnes présentes.

Basée sur un principe similaire, Zonetag [84], une application mobile de Yahoo!, permet aux utilisateurs de charger (en anglais, upload), "*flux montant*" les photos prises avec les smart-phones de marque Motorola ou Nokia vers leurs comptes Flickr. Zonetag considère des éléments de contexte tels que la localisation et le temps pour suggérer d'autres tags Flickr basés sur des tags précédemment saisis par l'utilisateur ou par les membres de son réseau social dans un contexte spatio-temporel similaire.

Le système MMM Image Gallery [85], quant à lui, capture l'ID de la cellule GSM, l'identité de l'utilisateur, l'instant de la prise de vue et la localisation via un module GPS. A l'aide d'un serveur web, ces informations sont analysées pour en déduire si la photo a été prise à l'intérieur ou à l'extérieur.

Dans ce même cadre, nous pouvons citer Photocopain [86], [87] un système d'annotation semi-automatique qui intègre plusieurs outils/services pour assister la tâche de l'annotation. Le prototype s'intègre dans le cadre du projet Advanced Knowledge Technologies Project¹³ (AKT). Le système Photocopain, tel qu'il est décrit dans la Figure 13, repose sur une architecture orientée service qui part de l'extraction des données EXIF des photos et des coordonnées géographiques capturés via GPS pour les enrichir avec des services web dédiés tels qu'un serveur de calendrier qui associe le temps de prise de vue à des événements décrits dans le calendrier de l'utilisateur, et des services spatio-temporels comme Gazetteer¹⁴ pour produire des annotations riches. De plus, il assiste le processus d'annotation manuelle par combinaison de tags et d'images extraites depuis des sites communautaires comme Flickr, avec des techniques d'analyse d'images comme celle de la classification afin de trouver un appariement entre des photos déjà annotées et d'autres en cours d'annotations. L'annotation produite par PhotoCopain est représentée dans le format IIM qui permet l'usage de l'outil AktiveMedia [88], [89].

¹³ <http://www.aktors.org/akt/>

¹⁴ http://gisdata.usgs.net/XMLWebServices/TNM_Gazetteer_Service.php

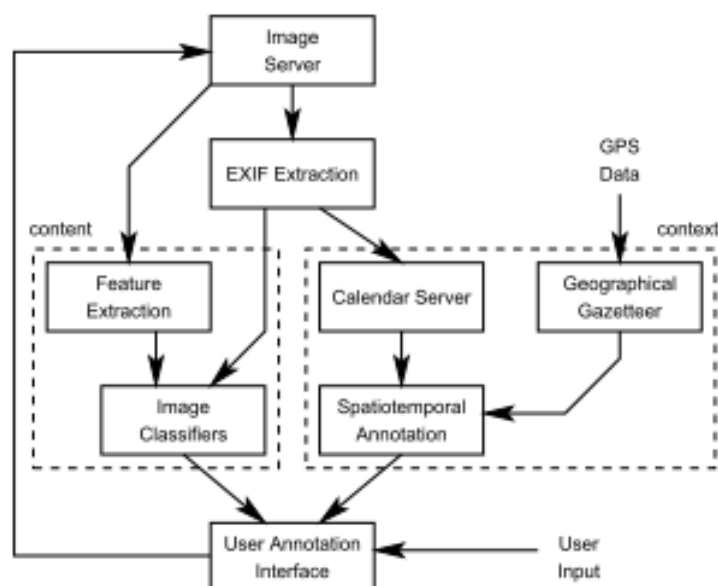


Figure 13. Processus d'annotation de photocopain décrit dans [86]

Le système MediaAssist [90], [91], proposé par le *Center Digital Video Processing*¹⁵ (CDVP) de l'université de Dublin, Ireland, est un système d'assistance à l'annotation et à la recherche des photos sociales qui se base sur le contenu et sur les données contextuelles. Il exploite la date et les coordonnées géographiques disponibles sur les métadonnées EXIF. Ces métadonnées sont exploitées pour en dériver des informations sur le contexte de création, d'une manière identique à celle proposée par PhotoCompas (l'adresse, la saison, le moment de la journée, etc.). De plus, MediaAssist utilise des méthodes de classification d'images reposant sur les informations bas-niveau pour détecter des bâtiments, des visages de personnes et les habits des personnes. Toutes ces informations sont suggérées comme des annotations de photos et peuvent être modifiées sur l'interface de l'application (Figure 14).

¹⁵ <http://www.cdvp.dcu.ie/>

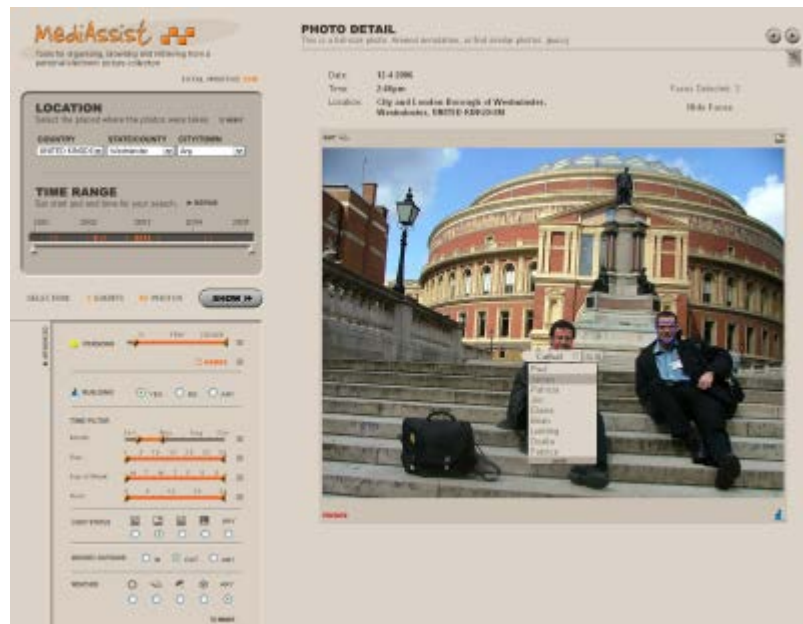


Figure 14. Interface d'assistance à l'annotation de MediaAssist décrit dans [90], [91]

Le recours aux dispositifs mobiles, dans l'annotation semi-automatique par des noms des personnes figurant a été étudié par Monaghan et O'Sullivan [92], [93].

Monaghan et O'Sullivan proposent un système appelé ACRONYM et un algorithme de suggestion des noms des personnes figurant dans les photos à l'aide de deux méthodes statistiques: (i) la première est basée sur la corrélation de Pearson [94] (ii) la deuxième est basée sur le Factor Analysis [95]. L'architecture d'ACRONYM (Figure 15) est une opération qui part de la capture des adresses Bluetooth des dispositifs environnants chaque fois qu'une photo est prise. Grâce à un registre des utilisateurs d'un site de média social, le système ACRONYM serait capable de récupérer l'identité des utilisateurs à partir de ces adresses. En cas d'échec, une méthode d'analyse des photos annotées est mise en place. L'objectif est de sélectionner des photos dont l'adresse Bluetooth a été également retrouvée ailleurs et de proposer l'annotation ajoutée préalablement à ces photos. Monaghan et O'Sullivan suggèrent, également, l'utilisation de méthodes de détection de visage telles que celles disponibles sur la bibliothèque OpenCV [96] afin de mettre en évidence sur l'image des régions contenant un visage. Une interface permettrait à l'utilisateur d'associer à chaque région les noms des personnes détectées par le Bluetooth. Le prototype ACRONYM est disponible en ligne¹⁶ et a obtenu un brevet d'intension.

¹⁶ <http://acronym.deri.org/>.

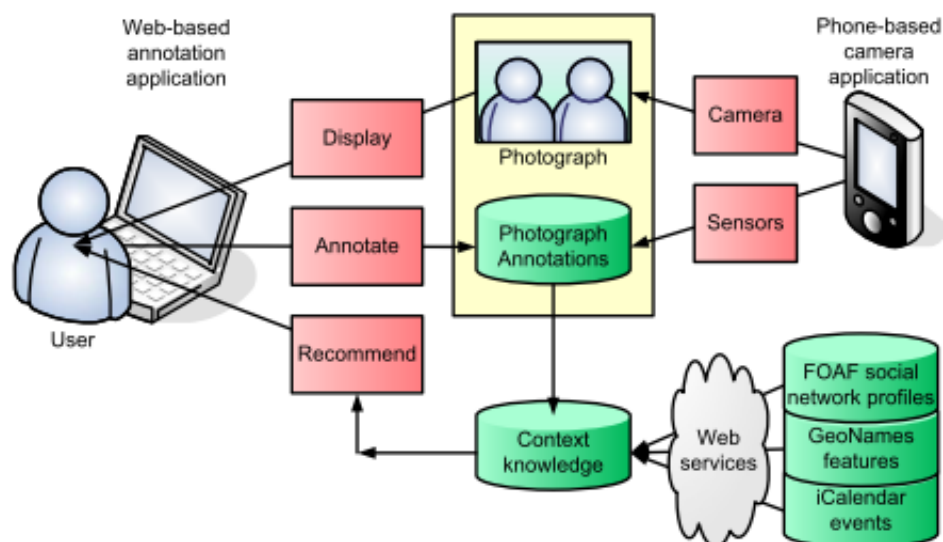


Figure 15. Architecture d'ACRONYM proposée dans [92]

Le système PhotoMap proposé par Viana et al. 2008 [97] utilise une connexion à un module GPS afin d'enrichir les métadonnées EXIF¹⁷ par des métadonnées géographiques et les identités des propriétaires des dispositifs mobiles proches afin d'attribuer des annotations indiquant ce qu'ils appellent le contexte social (e.g. qui est près de moi dans la photo ?).

4.6.2. Enrichissement du contexte physique par des ressources spécifiques à l'utilisateur

Sarin et al. 2008 [98] adoptent une approche basée sur l'utilisation des ressources spécifiques à l'annotateur pour enrichir et proposer des annotations à un utilisateur dans une application de bureau (en anglais, *desktop application*). L'approche proposée utilise Google desktop¹⁸ pour rechercher des fichiers stockés dans la machine de l'utilisateur, de même que la date et la localisation que la photo en cours d'annotation. Dans l'architecture illustrée par la Figure 16, un module d'extraction d'entités se charge d'extraire les index des fichiers les plus pertinents en suivant les catégories, date, localisation, personnes et organisations. Les fichiers les plus pertinents sont envoyés vers un module qui classifie les mots-clés en fonction de leurs fréquences d'apparition dans les sources de données. Les mots-clés ainsi ordonnés sont présentés à l'utilisateur sous la forme d'une liste de

¹⁷ Exchangeable image file format

¹⁸ <http://desktop.google.com/fr/>

recommandations. Après validation par l'utilisateur, les métadonnées, ainsi obtenues, sont représentées dans le format MPEG-7 pour décrire le contenu de la photo.

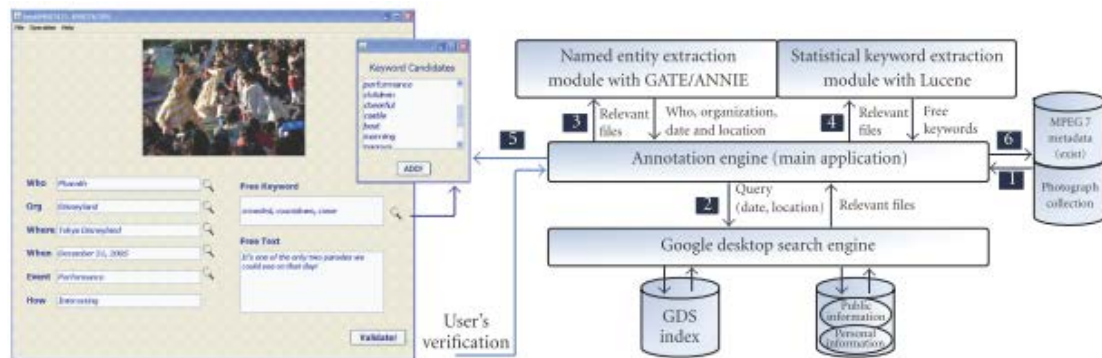


Figure 16. Vue générale du système d'annotation et de recherche proposée par [98]

Par ailleurs, les travaux de Naaman et al. 2005 [99] peuvent être considérés dans cette catégorie d'approches du fait qu'ils exploitent des annotations antérieures du même utilisateur pour en extraire des motifs de re-occurrence et cooccurrence des personnes dans différentes localisations et événements. Ces motifs sont utilisés pour fournir des suggestions d'identités de personnes, des localisations et des événements représentés dans des photos sociales.

4.6.3. Discussion

En ce qui concerne les approches qui prennent en considération la capture automatique des Bluetooth autour de l'appareil-photo comme MMM Image Gallery [85], ZoneTag [84] et PhotoMap [97] souffrent d'un problème de passage à l'échelle si le nombre d'éléments de contexte augmente (e.g. les Bluetooths présent autour de l'utilisateur). Quant à l'approche de Naaman et al. 2005 [99], elle souffre du problème de démarrage à froid si on se confronte à un nouvel utilisateur ou encore si l'utilisateur n'a jamais annoté manuellement ses photos sociales.

Malgré les techniques évoluées proposées par Sarin et al. [98], la source de données utilisée (i.e. les données existant dans la machine de l'annotateur) est insatisfaisante. En effet, les auteurs négligent la nature sociale des documents multimédias socio-personnels du fait que les documents de chaque utilisateur correspondent à des événements qui ocurrent dans son monde réel avec des personnes appartenant à son réseau social. Le même événement vécu dans un document multimédia socio-personnel peut être partagé par plusieurs personnes rassemblées dans un même réseau social. En effet, une personne est

présente dans ses documents multimédias socio-personnels avec des membres de son réseau social dans des événements qui, parfois se répètent plusieurs fois et avec les mêmes personnes.

L'ouverture sur les données du web de plusieurs utilisateurs permet d'enrichir les annotations et d'élever la probabilité de produire des recommandations plus pertinentes.

Afin de garantir une meilleure annotation automatique des documents multimédias socio-personnels, il nous semble que les propositions reposant sur la caractérisation du contexte de prise de vue sont prometteuses du fait que la capture d'éléments de contexte à partir de dispositifs mobiles devient réaliste vue leur intégration dans la plupart des dispositifs mobiles.

Par ailleurs, les sources d'enrichissement de contexte ne doivent pas se restreindre à des services web de type générique pour enrichir par des informations comme le climat, le jour du mois. Ces sources doivent être spécifiques à l'annotateur et/ou à son réseau social. En effet, l'utilisation d'agenda ou de profil d'utilisateur sur le web importé à partir de systèmes de réseaux sociaux ou de services de médias sociaux s'avère très fertile.

4.7. Résumé

Dans ce chapitre, nous avons présenté les travaux les plus importants en rapport avec l'enrichissement des annotations. Nous avons présenté les approches d'annotation automatique et semi-automatique. D'un côté l'annotation automatique permet un gain énorme au niveau du temps et l'énergie fournis par l'utilisateur. D'un autre côté, nous ne pouvons pas l'envisager, seule, pour des annotations ayant un haut niveau sémantique comme les événements.

Afin de garantir une annotation automatique consistante, il nous semble que les approches reposant sur la caractérisation du contexte de prise de vue sont prometteuses du fait que la capture d'éléments de contexte à partir de dispositifs mobiles devient réaliste vue leur intégration dans la plupart des dispositifs mobiles.

En ce qui concerne les approches semi-automatiques d'annotation, peu sont les approches qui exploitent le fait que ces documents sont produits et consommés par un réseau social d'amis, de membres de famille, de collègues, etc. Des annotations peuvent faire sens pour un groupe social comme elles peuvent avoir un autre sens pour un autre groupe. Les approches de suggestion d'annotations doivent, alors, être sensibles au

contexte social. Cette caractéristique de sensibilité sociale est souvent omise dans les systèmes de gestion de documents multimédias socio-personnels existants.

Afin de trouver un compromis entre la cohérence des annotations et leur caractère subjectif, il faut garantir d'une part leur consistance. Ceci peut être assuré en adoptant une approche qui repose sur la caractérisation du contexte physique. D'autre part, il faut enrichir d'une manière personnelle les annotations. Ceci peut être effectué en considérant la nature sociale des documents multimédias socio-personnels. En effet, les documents de chaque utilisateur correspondent à des événements qui ocurrent dans son monde réel avec des personnes appartenant à son réseau social. Le même événement vécu dans un document multimédia socio-personnel peut être partagé par plusieurs personnes rassemblées dans un même réseau social. Les annotations attribuées par un membre 'proche' appartenant au réseau social de l'utilisateur peuvent faire sens pour lui et peuvent être en rapport avec un nouveau document sujet d'annotation.

Chapitre 5. Modélisation et enrichissement des métadonnées

5.1.	INTRODUCTION.....	62
5.2.	DEFINITIONS.....	62
5.3.	MODELISATION DES METADONNEES EN CONSIDERANT LE CONTEXTE	63
5.3.1.	<i>Vue globale des métadonnées et de leur relation avec le contexte de prise de vue</i>	63
5.3.2.	<i>Modèle centré utilisateur de représentation des métadonnées</i>	66
5.3.3.	<i>Modélisation du profil social pour enrichir les métadonnées.....</i>	68
5.3.3.1.	Description des liens sociaux et des catégories sociales	70
5.3.3.2.	Règles d'inférence de sens commun	72
	• Règles implicites à l'ontologie exprimé en OWL	73
	• Règles explicites.....	74
	• Règles d'enrichissement.....	77
5.3.3.3.	Contraintes d'intégrités.....	78
5.3.4.	<i>Discussion.....</i>	78
5.4.	PLATEFORME D'ENRICHISSEMENT DES METADONNEES.....	81
5.4.1.	<i>Architecture de la plateforme d'enrichissement</i>	81
5.4.2.	<i>Génération des éléments de la couche physique du contexte.....</i>	83
5.4.2.1.	Exemple de capture du contexte physique lors de la création	84
5.4.3.	<i>Contexte de création enrichi.....</i>	85
5.4.3.1.	Enrichissement du contexte temporel	86
5.4.3.2.	Enrichissement du contexte spatial	86
5.4.3.3.	Création du contexte spatio-temporel	89
5.4.3.4.	Exemple d'enrichissement automatique du contexte de prise de vue.....	90
5.4.4.	<i>Contexte de création personnalisé.....</i>	91
5.4.4.1.	Enrichissement par des objets utilisateurs.....	92
5.4.4.2.	Enrichissement par événement.....	93
5.4.4.3.	Enrichir par l'identité des personnes proches.....	94
5.4.4.4.	Exemple d'enrichissement personnalisé pour Bob du contexte physique	94
5.4.5.	<i>Enrichissement manuel des métadonnées du contenu.....</i>	95
5.5.	RESUME ET DISCUSSION	96

5.1. Introduction

Dans ce chapitre, nous proposons un modèle conceptuel permettant de structurer les métadonnées des documents multimédias socio-personnels et faciliter, ainsi, leur gestion. Ce modèle est peuplé, partiellement, d'une manière automatique pour garantir la cohérence des annotations. Il est, également, peuplé d'une manière semi-automatique. Le processus d'enrichissement des métadonnées s'accomplit au moyen de deux phases. La première étant la capture d'éléments de l'environnement ambiant de l'utilisateur. La notion est appelée dans la littérature contexte de prise de vue. La deuxième est l'enrichissement par des ressources externes telles que les ressources spatiales, et la personnalisation via les profils sociaux.

5.2. Définitions

Nous allons, dans cette section, introduire quelques définitions utiles pour la présentation ultérieure de notre modèle d'annotation.

Définition 2. Activité documentaire

Une activité documentaire désigne une interaction entre l'utilisateur et le document. Dans notre cadre d'étude, nous nous intéressons aux documents multimédias socio-personnels, aux photos plus précisément. Cependant, le modèle mis en place peut convenir à tout document multimédia socio-personnel c'est-à-dire aussi bien aux vidéos qu'aux enregistrements sonores. L'activité documentaire peut être une création "*creation*", une annotation "*annotation*", une recherche "*search*" ou une consultation "*view*", etc.

Définition 3. Interaction sociale

Les *interactions sociales* relient les acteurs. Elles sont exprimées sous forme de verbes tels qu'*aimer*, *embrasser* ou des paraphrases telles qu'*offrir un cadeau*. Le vocabulaire décrivant les interactions sociales peut être emprunté à WorldNet.

Définition 4. Évènement

Selon Deleuze [100] *un évènement est un croisement entre une occurrence spatio-temporelle et un observateur qui lui prête une signification*. Un évènement peut, donc, signifier une interprétation d'un contexte environnemental par un consultateur.

Définition 5. Proximité sociale

Une *proximité sociale* fait référence à l'intensité de lien social reliant deux utilisateurs u_1 et u_2 . La proximité sociale est représentée par une valeur numérique comprise entre 0 et 1.

Définition 6. Dimension sociale

Une *dimension sociale* correspond à un chemin composé d'un nombre de relations sociales et identités de personnes. Ce chemin permet de relier le consultateur d'un document multimédia socio-personnel et les acteurs de ce document là.

Définition 7. Tag

Un tag correspond à un mot-clé associé à un objet (e.g. document multimédia socio-personnel, utilisateur, objet géo-localisé). Désormais, il est utilisé comme synonyme d'annotation.

5.3. Modélisation des métadonnées en considérant le contexte

5.3.1. Vue globale des métadonnées et de leur relation avec le contexte de prise de vue

La notion de contexte a été invoquée dans plusieurs domaines de recherche, essentiellement, dans le domaine des systèmes d'information sensibles au contexte (en anglais, *contexte-aware systems*). La notion de contexte fait référence, principalement, à la caractérisation de la situation de l'utilisateur lors de l'accès à un système. D'après Dey [101], *le contexte est construit à partir de tous les éléments d'information qui peuvent être utilisés pour caractériser la situation d'une entité. Une entité correspond ici à toute personne, tout endroit, ou tout objet (en incluant les utilisateurs et les applications eux-mêmes) considéré comme pertinent pour l'interaction entre l'utilisateur et l'application*¹⁹. Le qualificatif 'sensible au contexte' peut, donc, être associé aux systèmes qui guident leur comportement selon leur contexte d'utilisation.

Dans le domaine de la gestion des documents multimédias socio-personnels, le contexte de création n'a pas été précisément défini. Néanmoins, nous considérons les travaux de Viana et al. [102] comme les plus riches sur la modélisation du contexte. Les auteurs présentent le contexte comme un ensemble d'éléments qui sont

¹⁹ Traduction de l'auteur

- *Contexte social* : il représente les personnes présentes au moment de la prise de vue.
- *Contexte temporel* : il représente le jour, le mois, l'année, le jour de la semaine et la période du jour.
- *Contexte spatial* : il représente la latitude, la longitude, l'altitude, les objets proches et les relations spatiales.
- *Contexte spatio-temporel*: représente la saison, les conditions météorologiques, l'état de la lumière et les objets mobiles proches.
- *Contexte comportemental* : il représente les propriétés de la caméra, les Bluetooth proches.

En nous inspirant de cette définition, pour caractériser dans ce qui suit, le contexte de création des documents multimédias socio-personnels ou de la prise de vue comme :

Définition 8. Contexte de prise de vue

Le contexte de prise de vue désigne l'ensemble des éléments qui peuvent être capturés via l'environnement ambiant de l'utilisateur lors de la prise de vue et qui peuvent être enrichis afin de faciliter la tâche de la gestion des documents multimédias socio-personnels.

Nous stratifions, le contexte de création en trois couches (voir Figure 17) qui sont :

- **Couche physique** : dans la couche physique, le contexte est représenté par des éléments capturés via l'environnement ambiant. Nous distinguons, trois éléments principaux qui sont : l'élément spatial, l'élément temporel, et l'élément social. Le premier est capturé via des senseurs physiques comme GPS, GSM Cell-ID ou des senseurs logiciels comme le service Skyhook²⁰. Il est représenté par les coordonnées géographiques : latitude, longitude et altitude. Quant au second, il représente l'instant de prise de vue. Le dernier élément, est capturé via des senseurs Bluetooth. Nous supposons que les utilisateurs présents lors de la prise de vue sont identifiés à partir de leurs adresses Bluetooth ou leurs noms de dispositifs mobiles (en anglais, *friendly name*).
- **Couche enrichie** : dans cette couche, le contexte physique sera enrichi par d'autres informations. Ces informations proviennent des ressources sémantiques. Grâce à l'utilisation des ressources sémantiques, nous créons, à partir des éléments de contexte de couche physique, un contexte environnemental. Le contexte

²⁰ <http://www.skyhookwireless.com/>

environnemental est défini par la saison, les conditions météorologiques et l'état de la lumière. Pour Viana et al. [102], le contexte environnemental correspond à l'élément spatio-temporel. Pour le contexte spatial, nous l'enrichissons par les noms des objets proches tels que les monuments célèbres ou les musées (e.g. musée Dar Jellouli à Sfax) ou bien par les noms de l'arrondissement, de la ville et du pays. Enfin, en ce qui concerne le contexte temporel, il peut être enrichi via le jour, le mois, l'année, le jour de la semaine et la période du jour.

- **Couche personnalisée** : dans cette couche le contexte prend une autre forme. Il est personnalisé selon le profil social de l'annotateur. En effet, à travers un processus d'exploitation du profil social de l'utilisateur, nous explicitons des événements. Les profils sociaux d'utilisateurs exploités peuvent coïncider avec ceux des participants à la prise de vue ou à celui d'un consultateur *a posteriori*. Nous citons comme exemple d'événements liés aux utilisateurs l'anniversaire de ma cousine, le mariage de ma meilleure amie. Se basant sur les mêmes sources de données, le contexte spatial peut être interprété en des noms d'objets d'utilisateur géo-localisés. Nous citons comme exemple de ces objets d'utilisateur sa maison, le lieu de travail de mon frère. Enfin, le contexte social est personnalisé suivant les liens sociaux avec l'utilisateur. Par exemple, mon frère, le mari de ma meilleure amie. Ces liens sociaux, caractérisant l'identité des personnes figurantes ou celles des témoins à la prise de vue, sont qualifiés dans la suite de dimensions sociales.

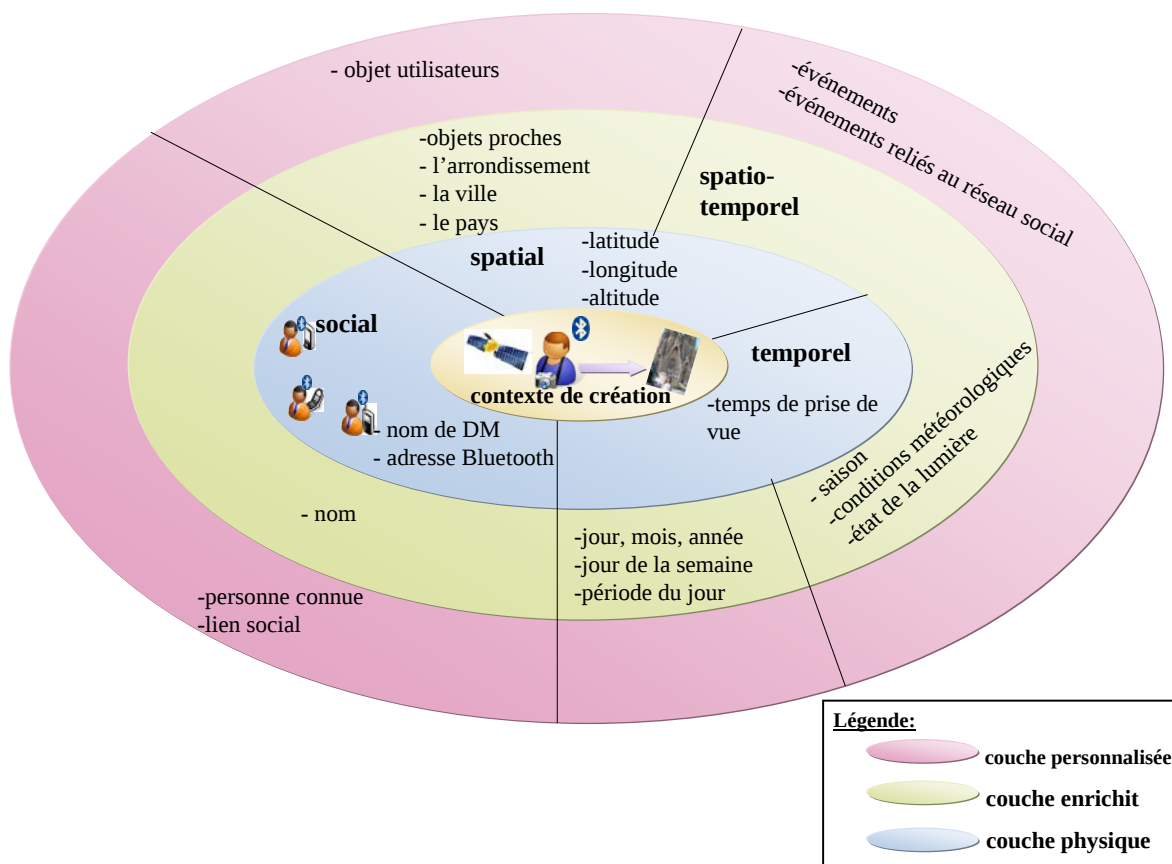


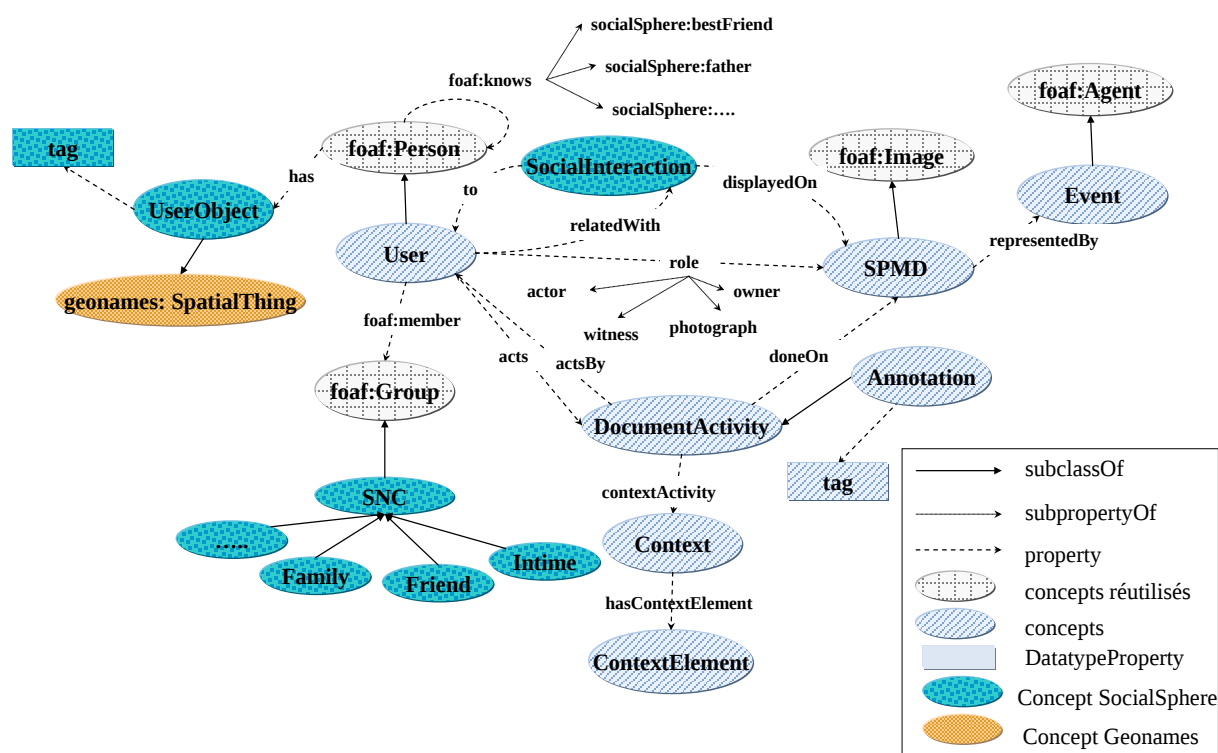
Figure 17. Les éléments de contexte et ses différentes couches

Dans la section suivante, nous décrivons notre modèle de données ayant pour objectif de structurer les métadonnées décrivant les documents multimédias socio-personnels. Le modèle proposé se base sur la capture du contexte lors de la prise de vue pour enrichir les métadonnées. Nous avons décrit le contexte de création d'une manière plus riche que les travaux précédents, les activités documentaires entre l'utilisateur et les documents multimédias socio-personnels, leurs rôles dans les photos et les interactions sociales entre acteurs.

5.3.2. Modèle centré utilisateur de représentation des métadonnées

Les personnes sont les repères les plus importants dans les photos sociales [92]. De ce fait, notre modèle de représentation est centré sur l'utilisateur et non sur la photo. Pour une gestion plus efficace des documents multimédias socio-personnels, nous avons besoin de modéliser la relation entre un utilisateur et les documents multimédias. Pour ce faire, nous avons construit une ontologie baptisée *SeMAT* (abréviation de *Semantic Model for self-Adaptive Tag*). L'ontologie *SeMAT* modélise la relation entre l'utilisateur et le document. Elle est représentée en format *RDF/RDFS/OWL*. Selon notre conceptualisation, l'association utilisateur-document

peut être distinguée en deux types : (i) association qui représente les rôles que les utilisateurs peuvent jouer dans une photo (i.e. représentée par la propriété *role* et ses sous-propriétés *actor*, *witness*, *photograph* et *owner* dans la Figure 18) et (ii) association qui représente des activités documentaires exercées sur le document dont le contexte qui peut être utile pour enrichir les métadonnées décrivant les photos (i.e. représentée par la classe *DocumentActivity* et sa sous classe *Annotation*). La classe *Annotation* possède un attribut (i.e. *DatatypeProperty* au sens *RDF*) appelé *tag* destiné à contenir les mots-clés donnés manuellement par les utilisateurs.



La Figure 18 représente le squelette des concepts de l'ontologie *SeMAT*. Les différents éléments seront détaillés dans la section 5.3.3. Les classes en hachures représentent celles propres à l'ontologie *SeMAT* (voir légende de la Figure 18). Dans cette conceptualisation nous avons employé le maximum de standards existants. Rappelons que, l'objectif fondamental du web sémantique est d'apporter la sémantique formelle nécessaire pour que des machines puissent consulter et interpréter les informations disponibles sur le web [22]. Pour cela, il est important de se resservir des ressources du web (i.e. des vocabulaires et des schémas d'ontologies existants). En essayant d'employer le maximum de standards, nous avons

réutilisé l'ontologie *foaf* (abréviation de Friend Of Friend) [48], *GeoNames*²¹[103] et *Time*²²[104]. L'ontologie *foaf*²³ a été proposée pour représenter le profil d'une personne sur le web. La modélisation du profil utilisateur selon *foaf*, inclut la représentation de son réseau de connaissance, de ses organisations, de ses documents, etc. Nous avons représenté une classe *User* propre à l'ontologie *SeMAT* comme une sous-classe (i.e. *RDF:subclass* au sens RDF) de la classe *Person* de *foaf*. Nous avons, aussi, représenté une classe *SPMD* (i.e. représentant un document multimédia socio-personnel, voir Définition 1) comme une sous-classe de la classe *Image* de *foaf*. La classe *Event* destinée à contenir une à plusieurs instances de *SPMD*, est représentée comme une sous-classe de la classe *Agent* de l'ontologie *foaf* (i.e. *property partOf*). Quant à l'ontologie *GeoNames*²⁴, elle a été utilisée pour décrire les entités spatiales du modèle *SeMAT*. Par exemple, pour décrire des éléments de contextes spatiaux (i.e. une instance de *semat: ContextElement*), nous avons représenté cette classe comme une sous-classe de *SpatialThing* de l'ontologie *GeoNames*. La classe *SpatialThing* est décrite via un ensemble d'attributs (i.e. *DatatypeProperty*) comme latitude (*lat*), longitude (*long*) et altitude (*alt*) qui sont les coordonnées géographiques d'une entité spatiale. L'ontologie *GeoNames* a été réutilisée, aussi, pour décrire les objets d'utilisateur dans le cadre de la spécification de l'ontologie *SocialSphere* décrite dans la section 5.3.3. Finalement, l'ontologie *Time* a été réutilisée pour décrire les entités temporelles du modèle. Ces instances sont principalement, une description plus détaillée du contexte temporel l'un des éléments de la classe *ContextElement*. Les exemples étalés sur les pages suivantes expliqueront la manière d'exploitation de ces trois ontologies en vue de représenter les métadonnées des documents multimédias socio-personnels ainsi que le profil social.

5.3.3. Modélisation du profil social pour enrichir les métadonnées

Parmi les enjeux de notre travail est de permettre une description personnalisée du contenu et du contexte des documents multimédias. Les enquêtes de Rodden et Wood [9] sur les exigences des utilisateurs dans un système de gestion de photos sociales en ligne ont prouvé que ces derniers préfèrent gérer les collections de photos selon des axes sémantiques et personnalisés, tels que les dimensions sociales caractérisant les acteurs des photos (e.g. ma

²¹ <http://www.geonames.org/>

²² <http://www.w3.org/TR/owl-time/>

²³ <http://xmlns.com/foaf/spec/>

²⁴ http://www.geonames.org/ontology/ontology_v2.1.1.rdf

mère, le cousin de mon mari) ou des indicateurs spatiaux caractérisant la localisation des photos (e.g. ma maison, le lieu de mon travail). Cet enjeu peut être réalisé en se servant d'un nombre d'informations qui se rattache à un utilisateur.

Pour atteindre cet objectif, nous proposons d'introduire la notion de profil social. Le profil social permet de représenter, pour chaque utilisateur, des informations personnelles, l'ensemble de son réseau de connaissance et des objets géo-localisés et tagués.

Définition 9. Profil social

Un profil social est associé à un utilisateur. Il comporte : (i) une description relevant des informations personnelles telles que son nom, son prénom, l'adresse Bluetooth relative à son dispositif mobile, etc.; (ii) une description de ses relations sociales et ses interactions sociales avec d'autres utilisateurs; (iii) et une description de ses objets qui sont d'ailleurs, des noms géo-localisés comme sa maison, son lieu de travail, etc.

Définition 10. Objet d'utilisateur :

Un objet d'utilisateur est un objet ayant des coordonnées géographiques. Les coordonnées géographiques de cet objet peuvent être taguées manuellement par les utilisateurs directement sur une carte géographique, à la manière de Viana [14] pour renseigner des étiquettes de places comme "*mon travail*", "*chez moi*", etc. Les coordonnées géographiques de ces objets ainsi que leurs tags peuvent être fouillées à partir de la fréquence de création des documents multimédias socio-personnels dans une localisation donnée.

Nous avons proposé, alors, l'ontologie *SocialSphere* pour modéliser l'ensemble des éléments du profil social. Les classes en pointillées dans la Figure 18 représentent les classes propres à l'ontologie *SocialSphere* (voir légende de la Figure 18). La modélisation de *SocialSphere* s'est focalisée essentiellement sur l'extension de la description des liens sociaux de *foaf*. Cette ontologie étend aussi la notion de groupe propre à *foaf* (i.e. *foaf:group*) et introduit la notion d'objet d'utilisateur, représenté par la classe *SocialSphere:UserObject*. Cette classe est dotée d'un attribut tag permettant de caractériser sémantiquement l'objet choisi par l'utilisateur et lui est relié par l'attribut *SocialSphere:has*. La classe *SocialSphere:UserObject* est modélisée en une sous-classe de la classe *GeoNames:SpatialThing*. Elle hérite, donc, la description des coordonnées géographiques.

Nous allons détailler dans la section suivante l'extension proposée pour décrire les liens sociaux et les catégories sociales.

5.3.3.1. Description des liens sociaux et des catégories sociales

Nous nous proposons de caractériser sémantiquement des liens sociaux avec un vocabulaire plus riche et des contraintes d'intégrité ainsi que des règles d'inférence permettant de raisonner et d'inférer d'autres liens sociaux non explicités par les utilisateurs. Le terme lien social désigne à la fois des interactions sociales et des relations sociales. Les interactions sociales sont introduites dans la Définition 3. Les métadonnées des documents multimédias socio-personnels sont enrichies non seulement par l'identification des personnes, du cadre spatial, temporel et spatio-temporel, mais aussi d'autres informations telles que les interactions inter- personnel. Pour modéliser ces dernières informations, nous avons proposé une classe *SocialSphere : SocialInteraction* ayant pour domaine de définition (en anglais, *domaine*) et domaine de valeur (en anglais, *range*) la classe utilisateur (*semat : User*). Une instance de la classe *SocialInteraction* peut être associée à la classe document multimédia socio-personnel (*semat : SPMD*) par le biais de la propriété *displayedOn* définie sur le domaine de valeur *SocialSphere : SocialInteraction* et sur domaine de valeur *semat : SPMD*.

En ce qui concerne les relations sociales, nous avons dédié un vocabulaire, des contraintes d'intégrité et des règles d'inférence aux relations sociales entre les personnes ainsi qu'aux groupes sociaux classiques. L'ontologie *SocialSphere*, mise en place, enrichit la propriété connaissance *knows* définie sur le domaine de définition et le domaine de valeur des personnes (*foaf : Person*) par un vocabulaire riche de relations sociales. L'ontologie *SocialSphere* accroît la classe groupe (*foaf : Group*) en définissant un concept plus spécifique attaché à la notion de catégorie sociale baptisée *Social Network Category (SNC)*. La classe *SocialSphere : SNC* contient cinq catégories sociales, à savoir : catégorie amis (*SocialSphere : Friendship*), catégorie famille (*SocialSphere : Familyship*), catégorie relation intime (*SocialSphere : Intime relationship*), catégorie relation professionnelle (*SocialSphere : Professional relationship*) et enfin catégorie de relation voisinage (*SocialSphere : Neighbourship*).

Le concept abstrait SNC contient un attribut stabilité (en anglais, *stability*) de la catégorie sociale (*SocialSphere : stability*). En effet, dans la vie réelle, seule la catégorie sociale famille peut-être considérée comme relation stable. Les autres catégories, comme celle de relations intimes ou encore celle de relations amicales, peuvent s'accentuer, se transformer d'une catégorie à une autre ou disparaître. Par exemple, un ami peut devenir un mari.

Nous avons défini un vocabulaire s'inspirant de l'ontologie *Relationship*²⁵. L'ontologie *SocialSphere* définit un vocabulaire de relations sociales plus riche, en y ajoutant un mécanisme d'inférence en plus de la notion des catégories sociales (SNC). Le Tableau 3 décrit les différentes spécifications de la propriété connaissance *knows* de l'ontologie *foaf* ainsi que leurs catégories sociales concernées.

Professional relationship	Family relationship	Neighborhood	Friendly relationship	Intime relationship
colleague	father	neighbor	friend	spouse
collaborator	mother	ex- neighbor	acquaintance	husband
employedBy	child		closeFriend	lifePartner
mentor	son		ex-closeFriend	girlfriend
apprenticeTo	daughter		ex-friend	boyFriend
supervisor	step-mother		bestFriend	fiance
ex-colleague	step-father		ex-bestFriend	ex-spouse
ex-collaborator	grandChild		KnowsInPassing	ex-husband
ex-employedBy	grandFather			ex-lifePartner
ex-mentor	grandMother			ex-girlfriend
ex-apprenticeTo	sibling			ex-boyFriend
ex-supervisor	uncle			ex-fiance
	cousin			ex-father-in-law
	nephew			ex-mother-in-law
	niece			ex-brother-in-law
	brother			ex-sister-in-law
	sister			father-in-law
	aunt			mother-in-law
	half-sister			brother-in-law
	half-brother			sister-in-law

Tableau 3: les différentes relations sociales ainsi que leurs différentes catégories

La Figure 19 décrit une instance et l'ontologie sous-jacente d'un réseau social spécifié en *SocialSphere*. Nous illustrons un exemple de représentation de relations sociales ainsi qu'un exemple d'inférence de relations via des règles qui seront détaillées ultérieurement. L'exemple présente Michel, déclaré comme étant l'enfant "*child*" de Jérôme, Michel est l'ami "*friend*" de Bob. Alice a une relation de fraternité "*sibling*" avec Jérôme et Carol est la sœur "*sister*" d'Alice.

²⁵ <http://vocab.org/relationship/rel-vocab-20050810.rdf> (spécialisation de FOAF)

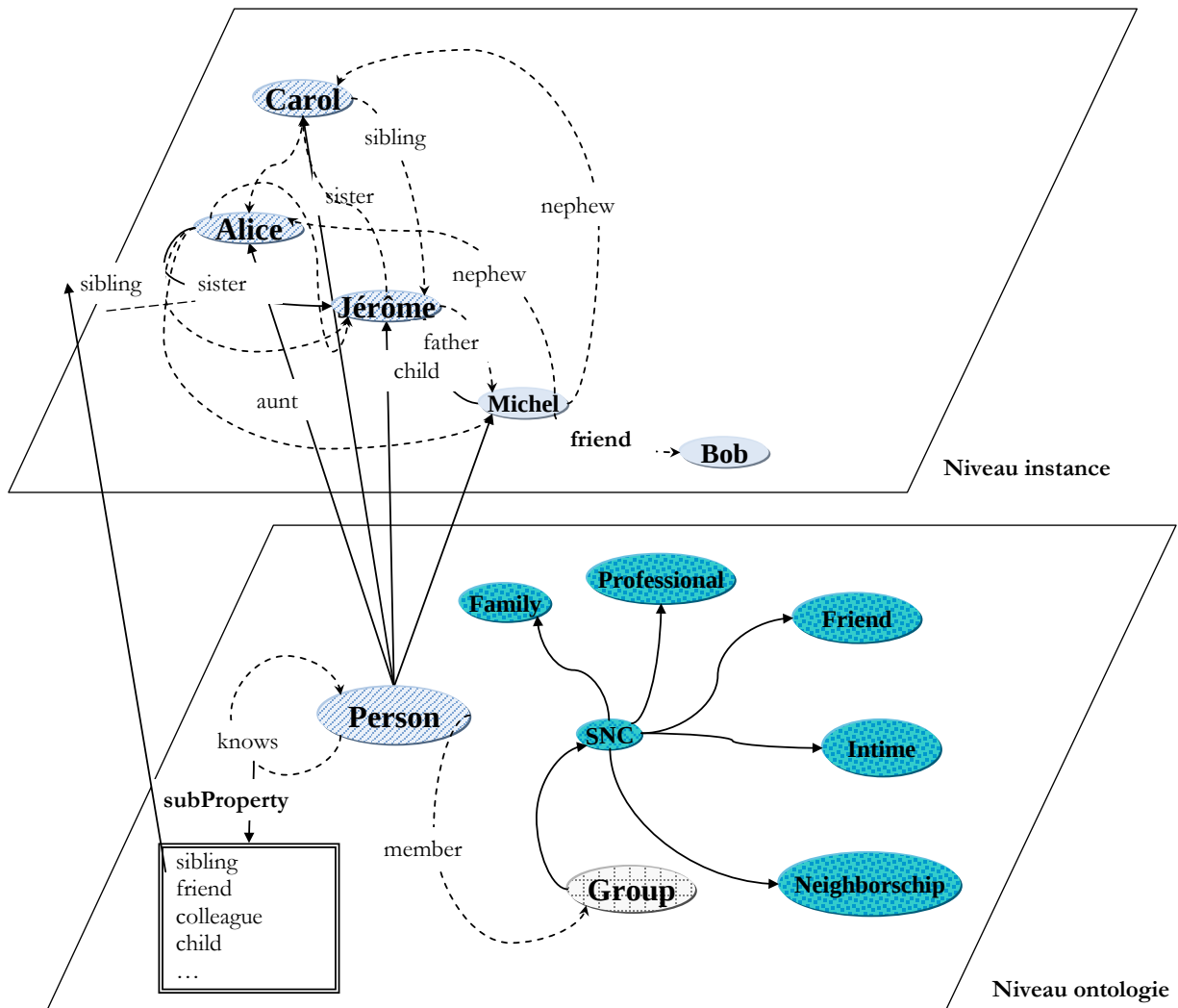


Figure 19. Réseau social enrichi avec l'ontologie SocialSphere

5.3.3.2. Règles d'inférence de sens commun

Les règles d'inférence jouent un rôle important dans les ontologies. En effet, elles permettent de substituer des relations à d'autres relations plus simples et d'inférer intelligemment des relations non explicitées.

Une règle d'inférence est, selon Wikipédia²⁶, une fonction qui prend des arguments appelés "les prémisses" et retourne une ou plusieurs valeurs appelées "conclusion". Les règles d'inférence peuvent également être vues comme des relations liant prémisses et conclusions par lesquelles une conclusion est dite "déductible" ou "dérivable" des prémisses. Si l'ensemble des prémisses est vide, alors la conclusion est appelée un "axiome" de la logique.

²⁶ http://fr.wikipedia.org/wiki/R%C3%A8gle_d%27inf%C3%A9rence

Nous avons mis en place des règles d'inférences de trois natures différentes qui sont à savoir : (i) implicites ; (ii) explicites ; et (iii) d'enrichissement. Dans ce qui suit, nous présentons les règles d'inférence que nous avons mises en place afin d'exploiter la puissance d'inférence de l'ontologie *SocialSphere* exprimée en OWL.

- **Règles implicites à l'ontologie exprimé en OWL**

Rappelons que, les ontologies exprimées en OWL se distinguent des autres représentations d'ontologies telles qu'UML, XML Schema, RDF Schema par son expressivité en termes sémantiques. OWL considère entre autres, les relations entre les classes et/ou instances, leur cardinalité, les caractéristiques des propriétés (symétrie, transitivité, inverse). Le langage OWL permet également d'exprimer des équivalences entre des classes et de créer de nouveaux concepts à l'aide des opérations ensemblistes telles que la disjonction, l'union, etc. OWL compte trois dialectes à l'expressivité croissante (OWL Lite, OWL DL et OWL Full). OWL DL, par exemple, repose sur la logique de description et peut être utilisé dans des processus d'inférence de connaissances. Ils existent des nombreux moteurs d'inférence (tels que Racer²⁷, Jess²⁸, Pellet²⁹ et Jena2 reasoner³⁰), pour raisonner sur la logique de description, qui acceptent des fichiers OWL DL comme entrée du processus. Ces moteurs d'inférence peuvent déduire de nouvelles connaissances (tels que des propriétés entre les instances de l'ontologie) à partir des "faits" décrits dans l'ontologie.

Nous avons exprimé des règles implicites via les caractéristiques de propriétés de transitivité, inverse et symétrie afin de déduire de nouvelles connaissances lors de l'expression de nouvelles relations. Par exemple, nous avons définie que la relation de fraternité «*sibling*» est une relation à la fois transitive et inverse de la manière suivante :

```
<owl:TransitiveProperty rdf:ID="sibling">
    <owl:inverseOf rdf:resource="#sibling"/>
</owl:TransitiveProperty>
```

Code 2. Règles implicites

²⁷ <http://www.sts.tu-harburg.de/~r.f.moeller/racer/>

²⁸ <http://www.jessrules.com/links/>

²⁹ <http://clarkparsia.com/pellet>

³⁰ <http://www.java2s.com/Open-Source/Java-Document/RSS-RDF/Jena-2.5.5/com.hp.hpl.jena.reasoner.htm>

- **Règles explicites**

La puissance d'expressivité d'OWL ne permet pas d'exprimer toutes les règles prévues par un concepteur. Pour cette raison OWL peut être associé à des règles d'inférence dites "*explicites*" exprimées à travers un langage à base de règles dédié comme le langage SWRL (Semantic Web Rule Language) ou le langage Jena Rules Language³¹. Les règles explicites sont prédéfinies par le concepteur. Nous avons utilisé Jena Rules Language comme langage d'expression de règles d'inférence explicites. L'expression de règles explicites permet d'étendre le raisonnement d'ontologies OWL à travers l'usage des variables et d'un moteur des règles d'inférence.

Le Code 3 présente un extrait des règles prédéfinies pour raisonner sur les relations sociales. Cet ensemble de règles permet, en le combinant aux règles implicites d'inférer d'autres connaissances. Nous illustrons l'interprétation des règles par l'exemple suivant. Selon la règle FamilyRule1 dans le Code 3, une relation de sœur "*sister*" est inférée entre la personne (P1) et la personne (P2) si on dispose de relation de fraternité "*sibling*" de (P1) à (P2) et (P2) possède une propriété sexe "*gender*" ayant la valeur femme "*women*".

Dans le cadre du même exemple, la relation de fraternité "*sibling*" est dotée à la fois de caractéristiques de propriété transitive et inverse. Ce qui signifie que si on dispose d'un fait (P2) *SocialSphere:sibling* (P3) alors en appliquant la propriété de transitivité une relation de fraternité est inférée de (P1) à (P3)

³¹ <http://jena.sourceforge.net/inference/>

```

# Jena Rules
@prefix SocialSphere:
<http://liris.cnrs.fr/~slajmi/ProfilFoaf/socialSphere.owl#>

[FamilyRule1:
  ( ?P1 SocialSphere:sibling ?P2 )
  (?P1 foaf:gender "women" )
  --> (?P1 SocialSphere:sister ?P2)
]

[FamilyRule2:
  (?P1 SocialSphere:sister ?P2)
  (?P3 SocialSphere:child ?P2)
  --> (?P1 aunt ?P3)
]

[FamilyRule3:
  (?P1 SocialSphere:aunt ?P2)
  (?P2 foaf:gender "man" )
  --> (?P2 SocialSphere:nephew ?P1)
]

[FamilyRule4:
  (?P1 SocialSphere:sibling ?P2)
  (?P2 SocialSphere:sibling ?P3)
  --> (?P1 SocialSphere:sibling ?P3)
]

```

Code 3. Exemple de règles d'inférences explicites exprimé en Jena Rule Language

Dans l'exemple représenté dans la Figure 19, La relation Jérôme est le père "*Father*" de Michel est obtenue grâce à l'utilisation de la caractéristique de propriété inverse "*inverseOf*" appliqué à la relation fils "*Child*" $\rightarrow \text{Inverse}(\text{child}) = \text{Father}$.

De même la relation tante "*aunt*" entre Alice et Michel est inféré grâce à la règle d'inférence (FamilyRule2). La relation neveu "*nephew*" entre Michel et Alice est inféré grâce à la règle (FamilyRule3) et puisque la relation de fraternité "*sibling*" est transitive nous obtenons que Carol est, aussi, "*sibling*" de Jérôme et que Michel et aussi le neveu "*nephew*" de Carol en appliquant la même règle (FamilyRule4).

L'ontologie "*SocialSphere*" définit, aussi, un ensemble de catégories sociales qui sont représenté en instance de la classe *Social Network Category* abrégé SNC. Il est possible, alors, de catégoriser les relations sociales définies par l'utilisateur en des catégories sociales. La déduction de l'appartenance à une catégorie sociale d'un utilisateur se fait automatiquement grâce aux règles d'inférences prédéfinies. Par exemple, un ensemble des relations sociales est classé en catégorie sociale *SocialSphere:Familyship* grâce à la mise en place de la règle illustré par le Code 4.

```

# Jena Rules
@prefix SocialSphere:
<http://liris.cnrs.fr/~slajmi/ProfilFoaf/socialSphere.owl#>
[FamilyshipRule:
  (?P1 SocialSphere:friend ?P2) or (?P1 SocialSphere:acquaintance ?P2)
or (?P1 SocialSphere:closeFriend ?P2) or (?P1 SocialSphere:ex-closeFriend
?P2) v (?P1 SocialSphere:ex-friend ?P2) or (?P1 SocialSphere:bestFriend
?P2) or (?P1 SocialSphere:ex-bestFriend?P2) or (?P1
SocialSphere:KnowsInPassing?P2)
  -->
  (?P2 SocialSphere:member SocialSphere:Familyship ) and
(SocialSphere:Familyship name ?C) and (?C SocialSphere:belongTo ?P1)
]

```

Code 4. Règle d'inférence de l'appartenance à des catégories sociales

Les règles qui prédefinisent l'appartenance de chaque relation sociale à une catégorie sociale sont mises en places au niveau de l'ontologie comme par exemple la règle (*FamilyshipRule*) définit dans le Code 4 qui infère l'appartenance à la catégorie sociale "*familyship*" en cas de l'existence de l'une des relations sociales définit dans le Code 4. De cette manière nous pourrions inférer, par exemple, que Michel appartient à la famille de Jérôme quant la relation sociale Jérôme le père de Michel existe. Le Code 5 illustre la définition de catégorie sociale "*familyship*" propre à Michel avec les membres qui y appartiennent.

```

<SocialSphere:Familyship>
  <foaf:name>ma Famille</foaf:name>
  <foaf:member>
    <foaf:Person rdf:resource="#Michel"/>
  </foaf:member>
  <SocialSphere:belongTo>
    <foaf:Person rdf:resource ="Jérôme"/>
  </SocialSphere:belongTo>
</SocialSphere:FamilyRelationship>
<foaf:Person rdf:resource ="Michel">
  <foaf:name>Michel Dupont</foaf:name>
  <foaf:knows>
    <SocialSphere:child rdf:ID="Jérôme"/>
  </foaf:knows>
</foaf:Person>

```

Code 5. Extrait du profil social de la personne Jérôme définissant l'appartenance de la personne Michel à sa catégorie sociale Familyship

Le Tableau 4 présente quelques exemples de règles implicites exprimées en RDF/RDFS/OWL.

Tableau 4. Exemple de règles et contraintes en OWL

	Règles implicites exprimées en OWL
1	<i>SocialSphere:parent</i> <i>rdfs:subPropertyOf</i> <i>foaf:knows</i> <i>SocialSphere:child</i> <i>rdfs:subPropertyOf</i> <i>foaf:knows</i> <i>SocialSphere:parent</i> <i>owl:inverseOf</i> <i>foaf:knows</i>
2	<i>semat:displayedOn</i> <i>rdfs:ObjectProperty</i> <i>rdfs:domain</i> = <i>semat:User</i> <i>rdfs:range</i> = <i>semat:Personal_Photo</i>
3	<i>SocialSphere:spouse</i> <i>rdfs:subPropertyOf</i> <i>foaf:knows</i> <i>SocialSphere:sister</i> <i>rdfs:subPropertyOf</i> <i>foaf:knows</i> <i>SocialSphere:sister</i> <i>owl:disjointWith</i> <i>SocialSphere:spouse</i>

Par exemple, nous définissons dans la première ligne du Tableau 4 que les propriétés parent "*SocialSphere:parent*" et enfant "*SocialSphere:child*", qui sont des sous propriétés de la propriété connaître "*foaf:knows*", sont des propriétés inverses. Dans la ligne 2 du Tableau 4, nous définissons une association (i.e. *rdfs:ObjectProperty* au sens RDF) au nom de "*semat:displayedOn*". Cet association part du domaine de définition des utilisateurs (i.e. *semat:User*) et va vers le domaine de valeur *semat:SPMD*. Enfin, dans la ligne 3 nous définissons qu'une personne définie comme sœur "*SocialSphere:sister*" ne peut pas être définie comme épouse "*SocialSphere:spouse*". Ceci à travers la définition de la relation de disjonction entre "*SocialSphere:sister*" et "*SocialSphere:spouse*".

• Règles d'enrichissement

Dans le monde des relations sociales, des relations complexes peuvent être simplifiées par des relations plus simples. Ainsi, une mise en place de règles de simplification ou de substitution s'annonce cruciale. Par exemple, la relation frère de mon père peut être substituée par la relation oncle. Si le fait existe : (P1) Father (P2) et (P3) Brother (P1) alors on peut inférer la relation (P3) uncle (P2) en appliquant la règle illustrée par le Code 6. L'ensemble de toutes les règles sont détaillé dans l'Annexe III.

```
# Jena Rules
@prefix SocialSphere:
<http://liris.cnrs.fr/~slajmi/ProfilFoaf/socialSphere.owl#>

[FamilyRule5:(?P1 SocialSphere:Father ?P2)
              (?P3 SocialSphere:Brother ?P1)
              --> (?P3 SocialSphere:uncle ?P2)
]
```

Code 6. Exemple de règles d'enrichissement

5.3.3.3. Contraintes d'intégrités

Les contraintes d'intégrités sont mise en place afin de garantir la cohérence du raisonnement ou d'interdire la mise en place de relations entre instances qui sont incohérent pour le concepteur. Pour ce faire, nous avons mis en place certaines contraintes afin d'interdire la mise en place de certaines relations. Par exemple, en mettant en place une restriction sur la cardinalité de la propriété *spouse* (*max cardinality=1*) nous garantissons l'interdiction de la polygamie. Ce qui signifie que si personne *x* déclaré comme épouse (*spouse*) d'une personne *y*. La même personne *x* ne peut pas établir la même relation *spouse* avec une autre personne *z* en même temps. De surcroît, la mise en place des contraintes de conjonction ou de disjonction comme la disjonction entre la propriété épouse "*property spouse*" et la propriété sœur "*property sister*" permet de garantir qu'une personne *x* déclaré comme sœur ne peut pas être en même temps épouse. La ligne 3 du Tableau 4 illustre la mise en place de cette contrainte en logique de description et en langage OWL.

5.3.4. Discussion

Nous avons présenté dans cette section un modèle de description des métadonnées basé sur la notion de contexte et centré sur l'utilisateur ainsi que des moyens de description du profil social. Nous avons montré l'importance du profil social pour l'enrichissement du contexte de prise de vue et son rôle dans l'enrichissement des annotations. Le profil social est décrit par trois types de descripteurs qui sont à savoir : (i) les informations personnelles spécifiées en "*foaf*" et en "*acronym*" (ii) les relations et les interactions sociales spécifiées en "*SocialSphere*" et (iii) les objets géo-localisés de l'utilisateur "*userObjects*" spécifiés en "*SocialSphere*".

Les relations sociales sont caractérisées par leurs instabilités et leurs évolutions. Par exemple, un simple ami peut devenir un mari. Vu que les documents multimédias socio-personnels sont sensé représenter toute une vie, il est important de distinguer, à un moment de

prise de vue d'une photo, la relation sociale exacte qui relie les tags identité des personnes à l'utilisateur. De ce fait, la représentation de métadonnées sur les relations sociales comme la date ou le lieu est cruciale comme par exemple associer la date de mariage à la relation épouse "*spouse*" ou mari "*husband*". Néanmoins, la représentation de ces métadonnées dans RDF/OWL est complexe. Il n'est pas possible de relier directement des classes à une propriété. La Figure 20 représente une comparaison de la représentation des métadonnées sur la relation apprenti "*apprentice*" entre Alice et Bob. Les métadonnées représentés sont le temps représenté par l'attribut TimeStamp=11 :40 et une instance de la classe Localisation "*Lyon*".

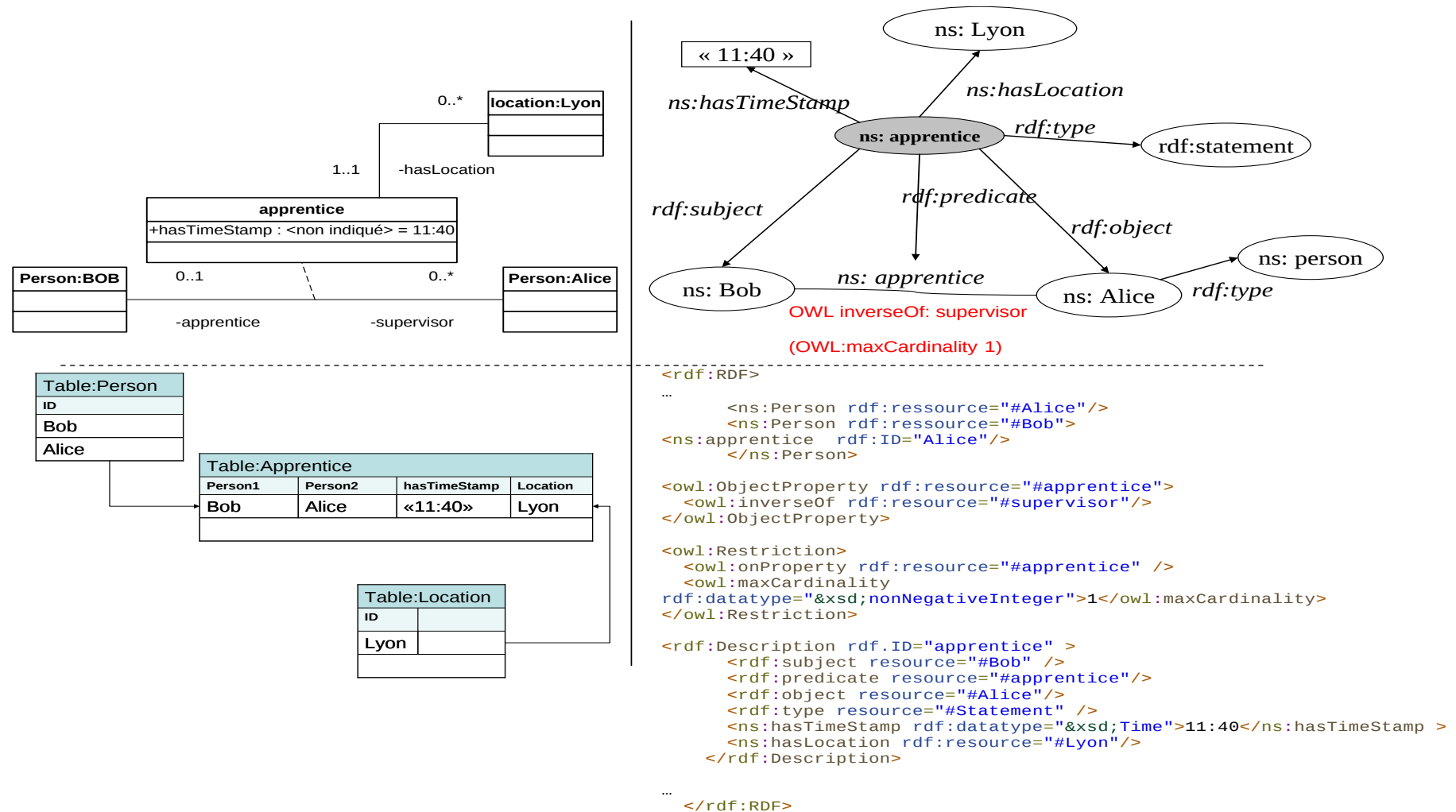


Figure 20: Comparaison entre les deux représentations des métadonnées sur les relations avec UML et avec RDF

5.4. Plateforme d'enrichissement des métadonnées

Dans les sections précédentes de ce chapitre, nous avons présenté plusieurs vocabulaires permettant une représentation formelle de quasiment toutes les annotations nécessaires pour décrire les documents multimédias socio-personnels. Dans cette section, nous proposons une plate-forme d'acquisition et de génération des annotations.

Nous souhaitons créer un processus d'annotation qui se base sur la capture d'un contexte de prise de vue initiale produit par un dispositif mobile doté de capteurs de contexte. Ces éléments de contexte sont tels que la date et l'heure, les coordonnées géographiques et les adresses Bluetooth doivent être enrichies et personnalisées afin de garantir leur exploitation lors de la phase de recherche.

L'objectif de la plateforme d'enrichissement des métadonnées est d'atteindre un plus haut niveau sémantique de métadonnées de documents multimédias socio-personnels à partir de ces informations initialement capturées par le dispositif mobile lors de la prise de vue.

Vue les importantes valeurs de capacités de calcul et d'accès aux sources de données exigées par les mécanismes d'enrichissement de métadonnées, les dispositifs mobiles d'aujourd'hui sont limités et une mise en œuvre de la plateforme d'enrichissement des annotations sur mobile sera moins efficace que son implémentation en version ordinateur de bureau (*desktop*) ou version sur serveur web. Pour cette raison, nous avons mis en œuvre la plateforme d'enrichissement sur le serveur web dont l'accès internet est supposé illimité et stable.

La plateforme d'enrichissement des annotations est utilisée sous réserve de l'existence de métadonnées initiales décrivant la couche physique du contexte.

5.4.1. Architecture de la plateforme d'enrichissement

Le mécanisme d'enrichissement de métadonnées (voir Figure 21) est effectué grâce à deux modules qui sont à savoir : (i) CCE (Contexte de Création Enrichi) ; et (ii) CCP (Contexte de Création Personnalisé).

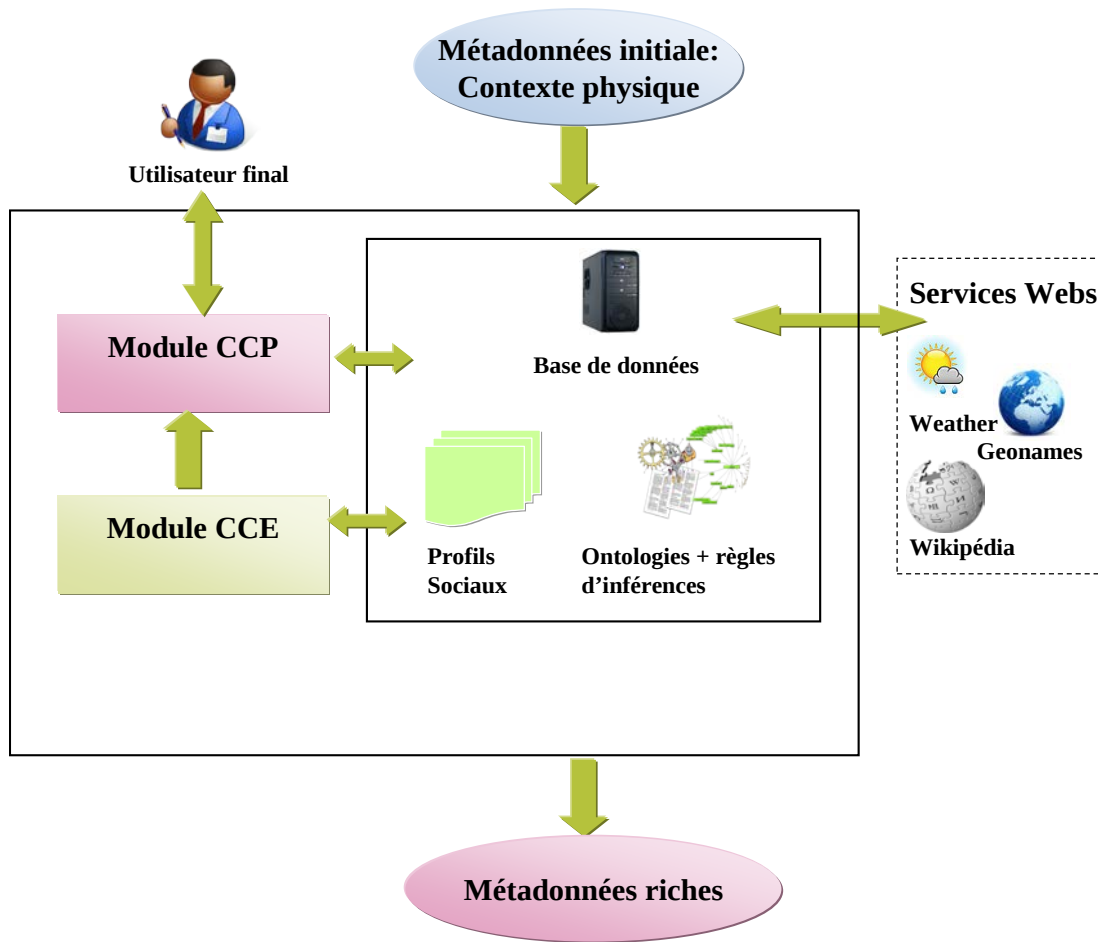


Figure 21. Architecture de la plateforme d'enrichissement semi-automatique des annotations

La plateforme fournit une ontologie (une instance de *SeMAT*) contenant les métadonnées d'une photo dont les métadonnées initiales sont passées en paramètre. Les deux modules sont composés d'un ensemble de services qui vont à leurs tour tenter d'augmenter ces métadonnées et représente le résultat sous-forme d'ontologie. Chaque service du module CCE utilise les métadonnées initiales (éléments du contexte physique) et génère d'autres informations (e.g. le service *NearbyObject* du module CCE génère des objets proches comme des monuments et des endroits célèbres) et les stocke dans le fichier correspondant à l'ontologie qui décrit les métadonnées de la photo en question.

Quant aux services du module CCP, ils sont destiné à proposer à l'utilisateur les informations générés afin qu'il les valide avant de les stocker.

5.4.2. Génération des éléments de la couche physique du contexte

Actuellement la plupart des téléphones mobiles sont dotés de captures d'informations sur la localisation comme le GPS ou le Cell-ID. Les principales motivations d'utiliser le mobile comme moyen de capture de photos sont:

- Premièrement, l'auto-synchronisation des horloges de téléphone avec l'opérateur réseau mondial. L'horloge élimine la nécessité pour eux d'être réglée manuellement.
- Deuxièmement, il est possible de capturer la localisation lors de la prise de vue. Le nombre des téléphones équipés de récepteurs GPS accroit. De plus d'autres moyens existent tels que le suivi cellulaire (Cell-ID) à proximité des tours ou dispositifs. Fondamentalement, les moyens de détections sous-jacentes de la localisation deviennent de plus en plus possible. Désormais, il est possible d'inclure dans son application mobile un service permettant à l'application d'accepter des sources de positionnement hétérogènes comme Skyhook³² qui combine plusieurs sources comme Wifi, GPS et autres pour donner les coordonnées géographiques les plus précises.

Ainsi, nous avons développé une application mobile dénommé PASMi (Photo Annotation and Sharing Middleware) permettant de capturer des éléments de contexte et de généré des métadonnées de la couche physique du contexte dont un exemple est illustré par la Figure 22. L'application PASMi permet aux utilisateurs mobiles de partager les photos capturer entre eux et garde les traces de partage de photos pour une utilisation dans le calcul de la proximité sociale³³.

³² <http://www.skyhookwireless.com/>

³³ La proximité sociale est un concept qui sera expliqué et étudié plus tard dans le chapitre suivant. Sa détermination sera utilisée dans le processus de recommandation.

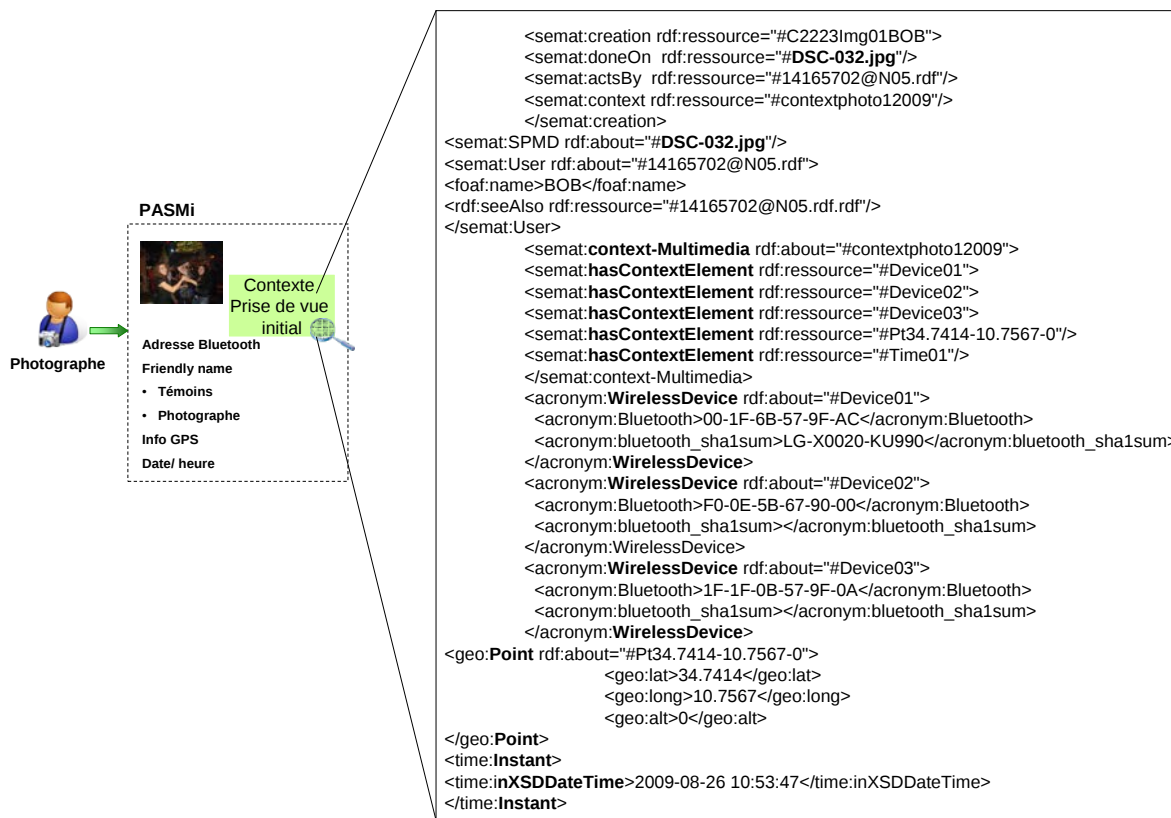


Figure 22. Exemple d'annotation automatique du contexte de prise de vue de la photo

5.4.2.1. Exemple de capture du contexte physique lors de la création

Nous illustrons un exemple de capture du contexte physique par un dispositif mobile. Bob crée une photo de Carol et Alice via son dispositif mobile. Lors de la prise de vue, son dispositif capture les adresses Bluetooth et le nom du dispositif (i.e. "*friendly name*") des dispositifs mobiles de David, Alice et Carol. Il capture aussi les coordonnées géographiques et le temps de prise de vue. La composante Mobile PASMi déployé sur le dispositif mobile de Bob crée, automatiquement, alors l'instance de contexte de la photo illustré par la Figure 23.

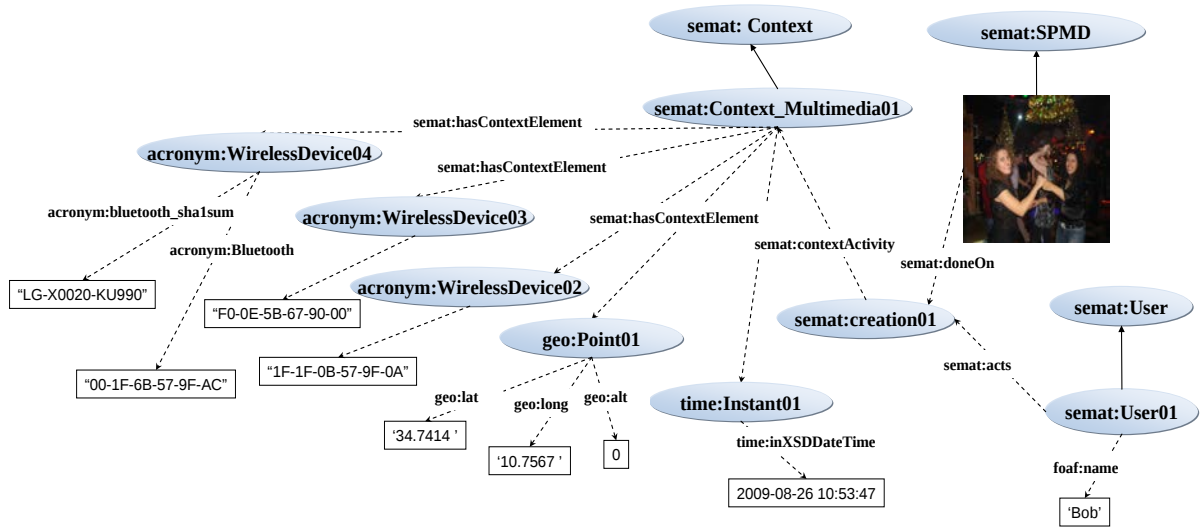


Figure 23. Exemple de création automatique de contexte de prise de vue

L'exemple décrit que lors de la prise de vue, le système a créé un contexte de création illustré par une instance de la classe *semat:creation* ayant un contexte de photo. Nous avons réutilisé des ontologies existantes telles que l'ontologie "acronym" développé par Monaghan et al. [92] et l'ontologie *foaf* qui a été initialement développé pour modéliser le profil de l'utilisateur sur le web. Sur cette étape nous avons ajouté que la classe *creation* qui permet de relier la photo (instance de la classe *foaf:Image*) à l'utilisateur qui est créateur de cette photo.

5.4.3. Contexte de création enrichi

Lors de la capture du contexte via des capteurs physiques ou logiques, le contexte est en état brut. Nous proposons de l'enrichir afin qu'il soit exploitable par l'utilisateur. En effet, comme expliqué dans le Chapitre 1, les études de l'utilisabilité de l'accès par contexte fournis dans [5] et [105] et [99] soutiennent qu'une partie importante des informations qui aident les personnes à se souvenir d'un document multimédia socio-personnel sont dérivés de contexte spatiale, temporelle et spatio-temporelle. Une partie importante des informations contextuelles que nous avons dérivées, grâce à l'utilisation des services web, ont été décrites dans l'étude de Naaman et al. [105]. Le module CCE : contexte de création enrichi, illustré par la Figure 24, a pour objectif d'enrichir les dimensions du contexte physique capturé au moment de la prise de vue. Ce module prend comme entrées : la longitude, latitude, altitude, timeStamp, des adresses Bluetooth et les enrichit en communiquant avec des services web dédiés comme le service web Wikipédia, le service

GeoNames et le service *sunrise and sunset Times*³⁴. Dans les sous-sections suivantes, nous avons décrit en détail l'utilisation des services web et le fonctionnement du module CCE.

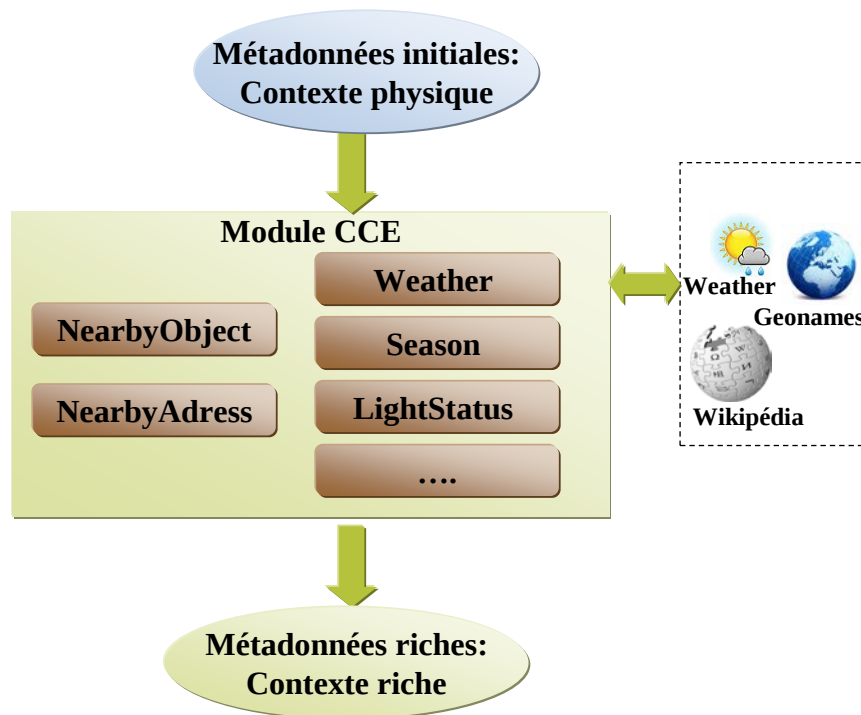


Figure 24. Entrées et sorties du module CCE

5.4.3.1. Enrichissement du contexte temporel

Au sein du module CCE, nous avons créé un service d'interprétation du temps permettant d'enrichir le temps de capture de la photo présentée sous la forme DD-MM-YY hh:mm:ss en des informations exploitables par l'utilisateur. Le service fournit, alors, divers comportements temporels tel que l'année, le mois et le jour de la semaine. Tous les attributs de la date calculés sont stockés dans l'ontologie comme des attributs "*dayOfWeek*", "*year*", "*month*".

5.4.3.2. Enrichissement du contexte spatial

Afin d'enrichir le contexte spatial, nous avons créé deux services *NearbyObject* et *NearbyAdress*. Le service *NearbyObject* a pour vocation de peupler l'ontologie avec des informations sur les dix plus proches points d'intérêts dans un rayon d'un kilomètre des

³⁴ <http://www.earthtools.org/>

coordonnées géographiques. En fait, ce service utilise le Wikipédia comme fonds documentaire de monuments touristiques. Wikipédia est un projet d'encyclopédie collective établie sur Internet, universelle, multilingue et fonctionnant sur le principe du wiki. Le nombre d'articles sur Wikipédia a dépassé les 11 millions. Ces descriptions sont utiles pour enrichir les annotations puisque dans certains cas, elles sont directement liées aux objets présents sur les documents multimédias socio-personnels créés par des utilisateurs. L'utilisation de ce service nous a permis d'enrichir les annotations avec des informations complémentaires relatives aux noms de monuments/ objets, les plus proches des coordonnées géographique contexte de prise de vue de la photo, indexés par la base Wikipédia.

Le Code 7 représente un exemple de réponse du service *findNearbyWikipedia*. En introduisant à ce service comme entrée les coordonnées géographiques, nous obtenons une réponse sous forme d'un ensemble d'entrées classées par la distance qui le sépare aux coordonnées passées en paramètre.

```
<GeoNames>
<entry>
<lang>fr</lang>
<title>Sfax</title>
<summary>
Sfax , deuxième ville et centre économique de Tunisie, est une ville
portuaire de l'est du pays située à environ 270 kilomètres de Tunis
[http://www.bab-el-web.com/carte/sfax.htm Distances à partir de Sfax].
Riche de ses industries et de son port, la ville joue un rôle économique
de premier plan avec l'exportation de l'huile d'olive et du poisson séché
(...)
</summary>
<feature>city</feature>
<countryCode>TN</countryCode>
<population>265131</population>
<elevation>0</elevation>
<lat>34.7414</lat>
<lng>10.7567</lng>
<wikipediaUrl>http://fr.wikipedia.org/wiki/Sfax</wikipediaUrl>
<thumbnailImg>
http://www.GeoNameess.org/img/wikipedia/85000/thumb-84863-100.jpg
</thumbnailImg>
<distance>1.0239</distance>
</entry>
</GeoNames>
```

Code 7. Exemple de réponse du service Wikipédia

Quant au service *NearbyAdress*, il permet de récupérer l'adresse avec ses variantes rue, ville et pays à partir des coordonnées géographiques de la prise de vue. En fait, ce service

repose sur le service *findNearestAddress* de GeoNames³⁵. GeoNames est une base géographique dont la description couvre tous les pays et les continents. Elle comporte plus de huit millions de noms de lieux qui sont disponible gratuitement. Nous avons utilisé la ressource sémantique Wikipédia en service web pour enrichir les coordonnées géographiques capturés lors de la prise de vue. Grâce à l'utilisation des services web de GeoNames, qui repose sur REST³⁶, nous avons enrichit les coordonnées géographiques par l'adresse et le pays.

Exemple : La récupération de l'adresse la plus proche d'une photo capturée au coordonnées géographiques suivantes *latitude*=34.734009 et *longitude*=10.763194 est illustré dans le Code 8. Elle correspond à un simple appel du service *findNearestAddress* qui sera formulé à l'issue de l'instanciation de la classe *Service_GeoNames*. La réponse sera sous format XML dont un exemple est illustré par le Code 9.

```
<?php
$geo=new Services_GeoNames();
$resultat=$geo-> findNearestAddress(array(
    'lat'=>37.76834106,
    'lng'=>-122.39418793)) ;
?>
```

Code 8. Utilisation des services de GeoNames depuis php

```
<GeoNames>
  <address>
    <street>R.Abou el Kacem Chabbi</street>
    <streetNumber>24</streetNumber>
    <lat>34.73400904100</lat>
    <lng>10.7631947853214000</lng>
    <distance>0.07</distance>
    <postalcode>94107</postalcode>
    <placename>Sfax, Tunisie</placename>
    <adminName2>Sfax</adminName2>
    <countryCode>TN</countryCode>
  </address>
</GeoNames>
```

Code 9. Exemple de réponse du service findNearestAddress

³⁵ <http://www.geonames.org/>

³⁶ <http://www.clever-age.com/veille/clever-link/soap-vs.-rest-choisir-la-bonne-architecture-web-services.html>

Les attributs récupérés utilisent la même division administrative de GeoNames. En effet, nous stockons dans l'ontologie une classe Adresse avec les attributs "street", "noStreet", "placename", etc.

5.4.3.3. Création du contexte spatio-temporel

La création du contexte spatio-temporel prend comme paramètre le contexte physique capturé. Nous avons créé un service appelé *Weather* permettant de récupérer des informations relatives au climat depuis la base GeoNames partant des coordonnées géographiques. Pour cela, nous utilisons le service *findNearByWeather* de GeoNames. Un exemple de réponse de ce service retournant des données relatives à la température, l'humidité et des conditions météorologiques en générale est illustré par le Code 10.

```
<GeoNames>
<observation>
<observation>DTTX 020800Z 23008KT CAVOK 26/15 Q1012 NOSIG</observation>
<observationTime>2010-06-02 08:00:00</observationTime>
<stationName>Sfax El-Maou</stationName>
<ICAO>DTTX</ICAO>
<countryCode>TN</countryCode>
<elevation>23</elevation>
<lat>34.7166666666667</lat>
<lng>10.6666666666667</lng>
<temperature>26</temperature>
<dewPoint>15</dewPoint>
<humidity>50</humidity>
<clouds>n/a</clouds>
<weatherCondition>n/a</weatherCondition>
<hectoPascAltimeter>1012.0</hectoPascAltimeter>
<windDirection>230.0</windDirection>
<windSpeed>08</windSpeed>
</observation>
</GeoNames>
```

Code 10. Exemple de réponse du service findNearByWeather

Ainsi, le service *Weather* peuple l'instance de l'ontologie relative à la photo capturée avec le contexte physique en paramètre au module avec les données du Climat récupéré depuis l'utilisation de ce service. Un exemple d'instance d'ontologie relative à l'annotation de la photo est illustré par la Figure 25.

Le service *LightStatus* permet de dériver une valeur décrivant le moment de la journée en prenant en compte l'heure de création de la photo. Nous considérons quatre moments de la journée qui sont à savoir : le matin (6 :00 à 12 :00) ; l'après midi (12 :00 à 17h :00) ; le soir (17 :00 à 21 :00) et la nuit (21 :00 à 6 :00).

Le service *Season* prend comme entrée la date de création de la photo, il permet de dériver une des quatre saisons : automne (21 septembre au 20 décembre), hiver (21 décembre au 20 mars), printemps (21 mars au 20 juin) et enfin été (21 juin au 20 septembre).

5.4.3.4. Exemple d'enrichissement automatique du contexte de prise de vue

Reprenons le même exemple illustré dans la section 5.4.2, dans lequel Bernard a créé une photo de Carol et Alice et le système a généré automatiquement un contexte physique. Pour ce contexte capturé, le module CCE génère des informations permettant l'enrichissement comme les objets les plus proches depuis la base Wikipédia, le jour de la semaine, l'état de la lumière, etc. illustré par la Figure 25.

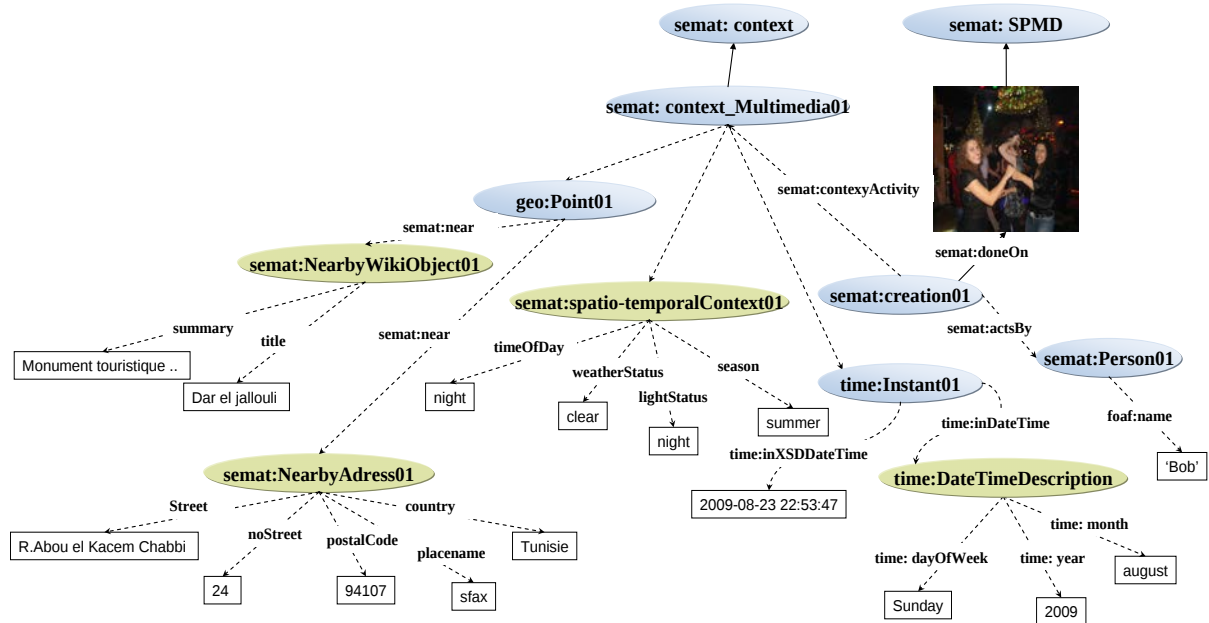


Figure 25. Enrichissement de contexte physique capturé

Pour l'enrichissement du *Time* nous avons utilisé l'ontologie *OWL Time* [106], [107]. Ensuite, pour enrichir les coordonnées géographiques, nous avons défini les classes *NearbyAddress* et *NearbyWikiObject* caractérisées chacune par un ensemble d'attributs. Nous avons défini pour la classe *NearbyAddress* les attributs *Street*, *noStreet*, *postalCode*, *placename* et *country*. Nous avons suivi la division administrative de *GeoNames*. Il est de même pour la classe *semat:NearbyWikiObject* ou la division administrative de *GeoNames* a été adopté pour stocker les dix entrées les plus proches du point en question. Ces entrées

correspondent à des endroits célèbres ou des monuments touristiques. Nous stockons pour chaque entrée Wikipédia *semat :NearbyWikiObject* un titre et un résumé.

5.4.4. Contexte de création personnalisé

Nous nous proposons d'enrichir le contexte physique par des données personnalisés et de plus haut niveau sémantique que les données produites par le module CCE. En effet, des études empiriques [9] ont prouvés que l'utilisateur préfère organiser, annoter et exploiter (rechercher) les photos en utilisant des axes sémantiques et personnalisés comme des événements (e.g. "*mon anniversaire*", "*mariage de mon voisin*", etc.), des personnes (e.g. "*moi*", "*mon fils*", "*ma mère*", etc.), des indicateurs temporelles (e.g. "*mois passé*", "*cette année*", etc.) et des indicateurs spatiales (e.g. "*ma maison*", "*lieu de travail de mon frère*", etc.). Ces mots-clés ne font sens que pour l'annotateur ou les participants dans la prise de vue. L'interprétation du contexte physique par des données consistantes (valable pour tous les utilisateurs) est guidée par le module CCE expliqué dans la section 5.4.3.

L'objectif du module CCP (contexte de création personnalisé) est de produire des données ayant un très haut niveau sémantique pour les participants à la prise de vue. Pour cette raison, les méthodes de ce module ne stockent pas directement l'information produite mais la propose à l'annotateur et/ou au consultateur en une liste de recommandation dans l'objectif de la valider. Le module CCP est composé de trois composants qui sont à savoir : (i) *PathEngine* ; (ii) *EventSearch* ; et (iii) *userObjectSearch*.

Pour cette fin, le module CCP prend comme entrées la photo et ses métadonnées initiales (contexte physique). La Figure 26 représente les composants les plus importants du module CCP. Ce module se sert du profil social pour produire des objets utilisateurs géo-localisés, des événements pouvant être en rapport avec l'utilisateur ou un des membres de son réseau social, et des liens sociales caractérisant mieux l'identité des personnes dans les photos. Le modèle correspondant au profil social est décrit dans la section 5.3.3. La suggestion des événements utilise les événements passés avec leurs contextes afin de prédire les événements associés aux nouvelles photos. Dans cette section nous nous contentons de montrer des exemples de stockage de l'information produite par chacune des composants (*PathEngine*, *EventSearch*) après validation de l'utilisateur. Le processus de recommandation ainsi que les algorithmes proposés et mis en œuvre pour ce faire font l'objet de Chapitre 6.

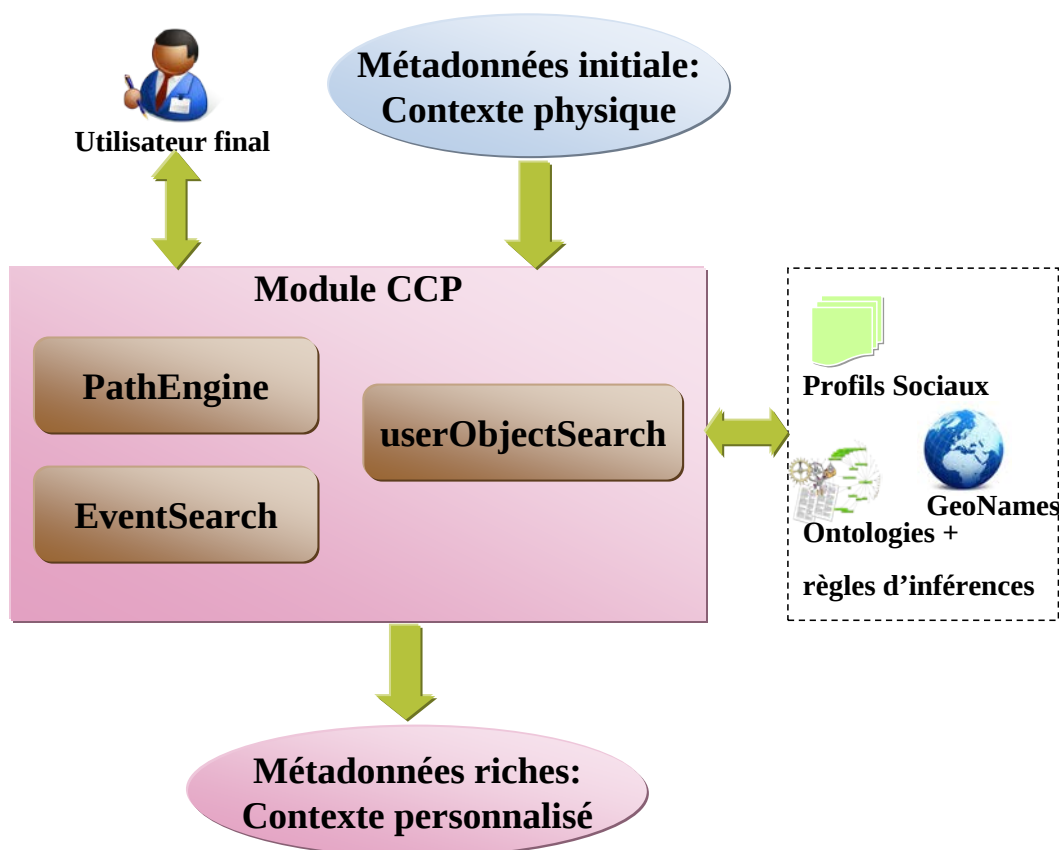


Figure 26. Composant du module CCP

5.4.4.1. Enrichissement par des objets utilisateurs

Un composant du module CCP appelé *userObjectSearch* se charge de retrouver les objets tagués par les participants à la prise de vue des coordonnées géographiques de la prise de vue. La recherche de ces objets se base sur un calcul de distance euclidienne [108] entre les coordonnées géographiques de la photo en cours d'annotation ou de consultation et les coordonnées géographiques de chaque objet d'utilisateur appartenant aux participants à la prise de vue. Les tags décrivant les objets utilisateurs les plus proches (à un seuil paramétré) sont proposés à l'utilisateur en vue de l'assister à l'annotation de noms de localisations personnalisées et proches à la photo.

Dans le cadre de cette thèse, nous supposons la préexistence de ces informations dans le profil social de chaque utilisateur.

Dans l'exemple cité au début du chapitre, nous supposons la préexistence dans le profil social de May un objet ayant des coordonnées géographiques proche (à un seuil près) des coordonnées de la photo en cours d'annotation. Cet objet a comme tag "*home*". Le système en découvrant le profil de May à travers le profil de Bob propose à Bob le tag "*home May*"

pour annoter la photo en cours. Si Bob accepte la proposition du système, le nœud représentant l'objet d'utilisateur en question sera relié à la photo en cours d'annotation comme illustré dans Figure 27.

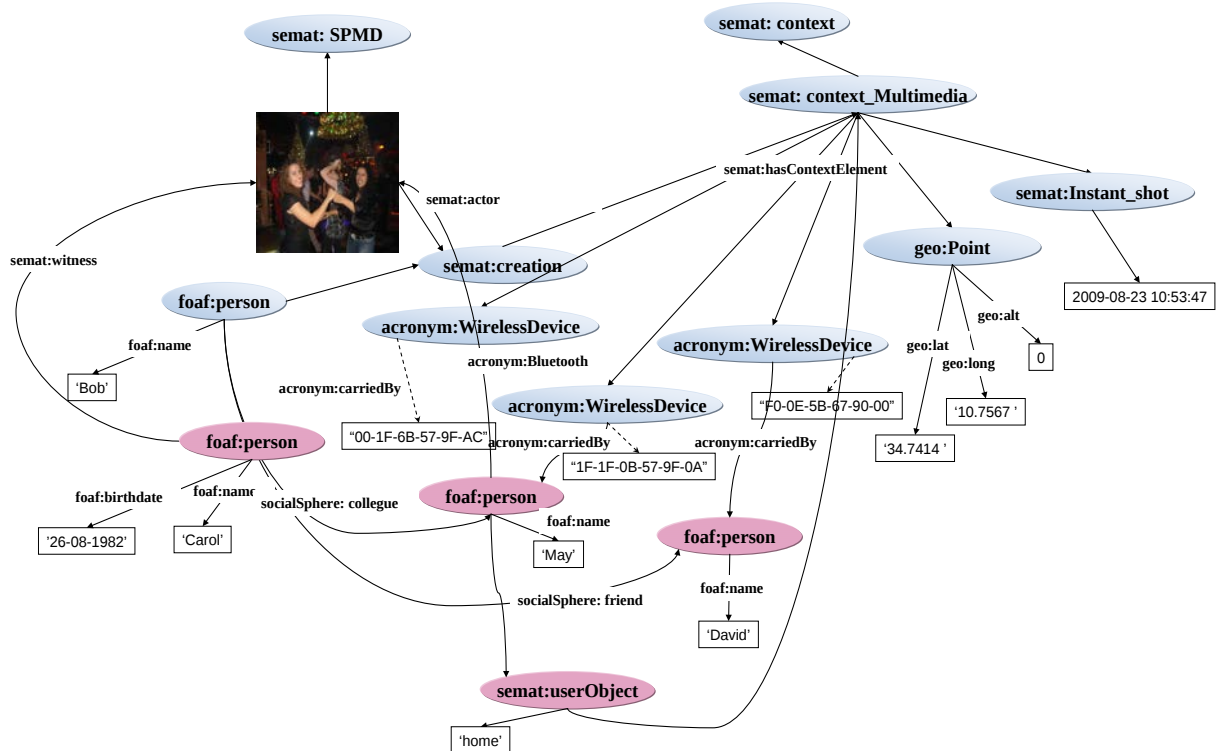


Figure 27. Exemple d'enrichissement personnalisé se servant du profil social: l'objet d'utilisateur de May, une participante à la prise de vue est rattachée à la photo

5.4.4.2. Enrichissement par événement

Les événements peuvent être des tags annotés manuellement par les utilisateurs. Ils correspondent à des instances de la classe *semat:Event*. Pour assister l'utilisateur dans la spécification des événements, nous pouvons nous baser sur les travaux Lim et al [62] pourtant sur la modélisation de taxonomie d'événements de la vie quotidienne. La taxonomie a comme racine 'events' et les catégorise en 'gathering', 'family activity' et 'visits places'. Pour le haut niveau sémantique qu'elles présentent, nous avons choisie d'adopter une approche semi-automatique pour assister l'utilisateur dans l'annotation des événements liées aux collections de photos. Une approche basée sur des règles d'associations générées par groupe social est présenté dans la section 6.3 Chapitre 6. L'algorithme de génération des recommandations d'annotations des événements se sert de règles d'associations préparé pour l'annotateur et des personnes décrit dans son profil social ayant un haut niveau de proximité sociale. L'algorithme proposé est exécuté au sein du composant *EventSearch*.

5.4.4.3. Enrichir par l'identité des personnes proches

Un composant du module CCP baptisé *PathEngine* s'en charge de parcourir le réseau social du photographe pour retrouver l'identité des adresses Bluetooth capturées. Ainsi, les nœuds instance de *foaf:person* ayant les adresses Bluetooth capturées seront reliées à la photo par la propriété témoin (*Witness*). *PathEngine* permet, aussi, de suggérer des dimensions sociales correspondant à des relations ou des interactions sociales entre personnes lors de la visualisation des photos. L'algorithme SQO (Social Query Optimisation) est exécuté au sein du composant *PathEngine*. Il permet de : (i) retrouver les nœuds personnes et leurs profils social à partir des adresses Bluetooth capturées ; (ii) retrouver des liens sociaux qui relient des consultateurs des photos aux tags identités des personnes dans les photos. Le principe de cet algorithme sera détaillé dans le Chapitre 6.

5.4.4.4. Exemple d'enrichissement personnalisé pour Bob du contexte physique

L'enrichissement personnalisé est une étape qui s'exécute lorsque l'utilisateur charge les photos dans le composant web du système. La composante *EventSearch* permet de proposer à l'utilisateur des événements pour assister son annotation des événements liés aux documents multimédias socio-personnels.

En reprenant le même exemple cité au début du chapitre, lorsque Bob charge les photos dans le composant web du système. Le composant *PathEngine* retrouve l'identité des propriétaires des dispositifs mobiles capturés ayant les adresses Bluetooth suivantes : 1F-1F-0B-57-9F-0A, F0-0E-5B-67-90-00 et 00-1F-68-57-9F-AC. Les adresses Bluetooth ainsi que les potentiels noms de dispositifs ont été décrit grâce au module de capture et génération de contexte physique décrit dans la section 5.4.2. Le module de capture instancie l'ontologie ACRONYM [92] qui permet de décrire un dispositif mobile dans un profil social.

Dans l'exemple décrit par la Figure 27, le composant *PathEngine* a créé un lien entre l'instance de *WirelessDevice* ayant comme valeur d'attribut *acronym* : Bluetooth : 1F-1F-0B-57-9F-0A et l'instance de *foaf : person* ayant comme valeur d'attribut *foaf : name* : David. Le composant *PathEngine* procède de la même manière en ce qui concerne le dispositif mobile *WirelessDevice* ayant l'adresse Bluetooth = F0-0E-5B-67-90-00, pour relier le nœud représentant May partant de l'adresse Bluetooth correspondant. Les deux

personnes David et May sont retrouvés directement à partir du profil social de Bob. Quant à la personne Carol, elle est retrouvée via le profil social de May.

Quant au composant `userObjectSearch`, il permet à travers un calcul de distance entre les coordonnées géographiques de la photo en question et les objets utilisateurs (i.e. instance de *semat:userObject*) de retrouver les objets utilisateurs proches et de les relier à la photo dans l'ontologie. La recherche des objets utilisateur est effectuée sur le profil social des participants à la prise de vue. Le résultat de la recherche des objets utilisateur est présenté à l'annotateur en une liste de recommandation. L'annotateur est, alors, libre de choisir ou refuser les propositions. Après validation de l'annotateur, l'instance de *semat:userObject* choisie sera relié à la photo et sera présenté dans la suite en tant que métadonnées de la photo en question.

Dans l'exemple de la Figure 27, l'instance de la classe *semat:userObject* ayant un attribut `tag=Home` décrit dans le profil social de la personne May est relié à la photo après validation de l'annotateur.

5.4.5. Enrichissement manuel des métadonnées du contenu

L'enrichissement des métadonnées qui concernent le contenu comme, par exemple, le marquage des acteurs dans la photo est effectué manuellement par l'utilisateur.

L'utilisateur pourra annoter librement en utilisant une liste de mot-clés. L'utilisateur pourra aussi lier des acteurs entre eux via des interactions sociales.

La Figure 28 illustre une instance de l'ontologie *SeMAT* sous-jacente d'une description manuelle de deux acteurs May et Carol dans la photo et une interaction sociale entre eux. Une instance de la classe *SocialSphere:SocialInteraction* (*SocialSphere:dancing*) est créée et a été relié de la personne May (avec l'attribut *relatedWith*) à la personne Carol (avec l'attribut *to*). Cette interaction sociale a été, aussi, rattachée à la photo grâce à l'attribut *displayedOn*.

La saisie des interactions sociales est assistée par un vocabulaire définis dans l'ontologie *SocialSphere*.

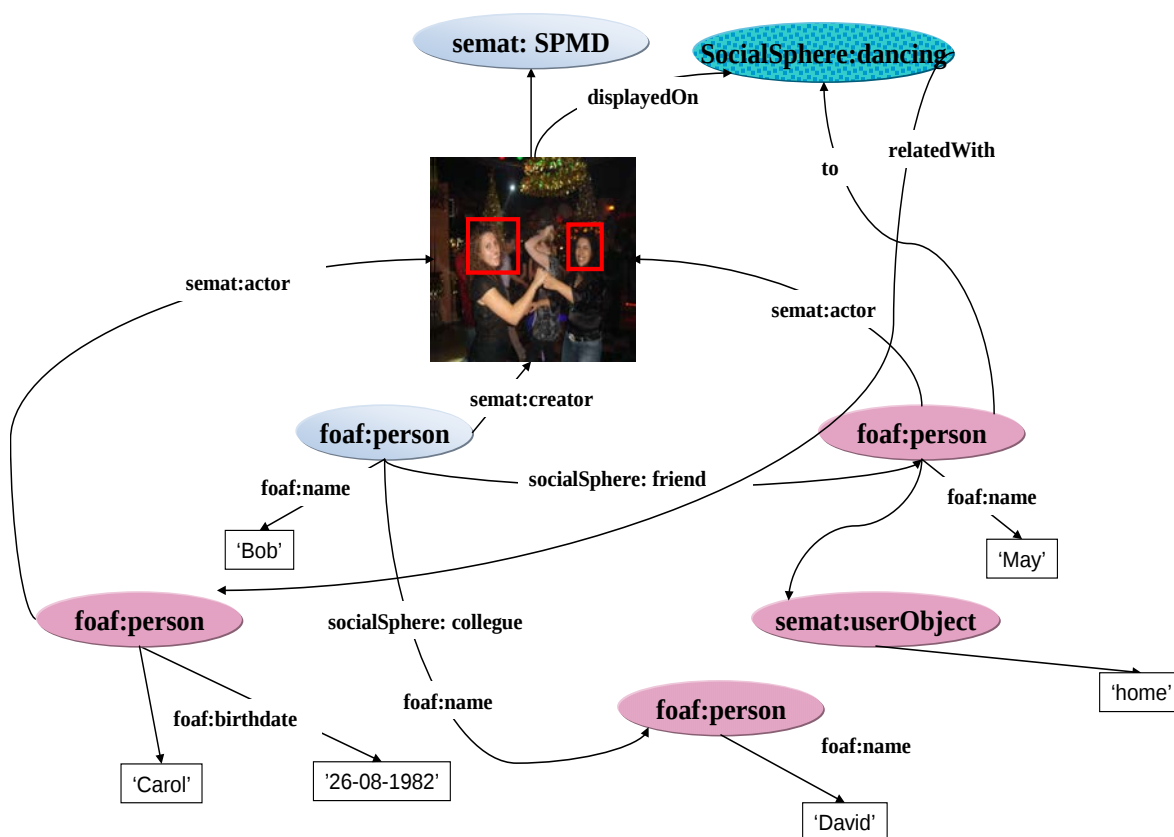


Figure 28. Insertion de métadonnées sur le contenu par l'utilisateur :actor et socialInteraction

5.5. Résumé et discussion

Dans ce chapitre, nous avons présenté un modèle baptisé "SeMAT" (Sematic Model for Self-Adaptive Tag) permettant aux annotations décrivant les documents multimédias de s'adapter selon les profils sociaux des utilisateurs et de leurs groupes sociaux respectifs. Pour obtenir des annotations ayant un haut niveau sémantique, nous partons de la capture du contexte physique. Le modèle conceptuel "SeMAT" met en évidence la relation entre l'utilisateur et le document en modélisant son rôle dans ce dernier et le contexte de ses activités documentaires. Une attention spécifique est accordée au contexte pour le rôle qu'il joue à apporter un enrichissement des annotations. Nous avons divisés le contexte en (i) temporel, (ii) spatial, (iii) social et (iv) environnemental. En ce qui concerne le module Contexte de Création Enrichi (CCE), la valeur ajoutée par notre plateforme réside dans le module Contexte de Création Personnalisé (CCP) qui permet de proposer à l'utilisateur des annotations sensible à son contexte social. Ceci est assuré en se servant du profil social. Une modélisation de quelques éléments du profil social était nécessaire : nous avons réutilisé la description des informations personnelles proposées dans l'ontologie "foaf" et

"acronym", nous avons modélisé les liens sociaux et les objets utilisateurs en proposant une nouvelle ontologie baptisée "*SocialSphere*".

Le chapitre suivant présente en détail les algorithmes proposés pour suggérer des annotations personnalisés à l'utilisateur (en particulier les tags événement et les tags identité des personnes dans les photos).

Chapitre 6. Sensibilité sociale des recommandations d'annotations

6.1.	INTRODUCTION.....	100
6.2.	MOTIVATIONS	101
6.3.	RECOMMANDATION DES EVENEMENTS POUR ANNOTER DES PHOTOS.....	102
6.3.1.	<i>Nature des données considérées</i>	<i>104</i>
6.3.2.	<i>Découverte des motifs fréquents</i>	<i>106</i>
6.3.3.	<i>Découverte des règles d'association.....</i>	<i>107</i>
6.3.4.	<i>Suggestion de tag événement basée sur la proximité sociale.....</i>	<i>109</i>
6.3.5.	<i>Discussion.....</i>	<i>113</i>
6.4.	ENRICHISSEMENT DES TAGS IDENTITE LORS DE LA CONSULTATION DES PHOTOS.....	114
6.4.1.	<i>Heuristiques considérées</i>	<i>116</i>
6.4.2.	<i>Algorithme SQO.....</i>	<i>118</i>
6.4.3.	<i>Analyse de la complexité de l'algorithme SQO</i>	<i>121</i>
6.4.4.	<i>Discussion.....</i>	<i>122</i>
6.5.	RESUME.....	122

6.1. Introduction

Les recherches dans le domaine de l'annotation de documents multimédias socio-personnels deviennent de plus en plus attractives en raison du progrès des moyens de capture de contexte et de caméras dans des dispositifs mobiles intelligents. Aujourd'hui, les utilisateurs chargent un nombre important de photos dans des sites de réseau social et les partagent avec leurs réseaux de connaissances.

Des études empiriques [99] ont soutenu qu'une partie importante des informations qui aident les personnes à se souvenir d'un document multimédia socio-personnel fait référence au *contexte*. Les événements en font partie et représentent l'élément le plus difficile à déterminer vu leur niveau sémantique élevé (i.e. on ne peut pas le capturer automatiquement via l'environnement ambiant de l'utilisateur tel que d'autres éléments de contexte comme la localisation, le temps et les dispositifs présents lors de la prise de vue). Des exemples d'événements peuvent correspondre à des occurrences de scènes vécues réellement comme '*voyager à l'étranger*', '*assister à une conférence*', '*regarder un film au cinéma*', '*assister à un dîner gala*', etc.

Dans ce chapitre, nous nous préoccupons de l'assistance de l'utilisateur lors de l'annotation de documents multimédia. Plus spécifiquement, nous nous intéressons aux tags qui font référence à l'événement de la photo, à l'enrichissement des informations sur l'identité des personnes selon le consultateur et de l'enrichissement des coordonnées géographies des documents multimédias par des noms d'objets proches

Quelques stades d'évolution des propositions détaillés dans ce chapitre ont été rapportés dans [109], [110] et [111].

Ce chapitre est organisé comme suit. Nous commençons, tout d'abord, par une présentation des motivations qui nous ont amenées à bâtir un système de recommandation sensible au contexte social (section 6.2). La section suivante (section 6.3) a pour objectif de présenter notre approche de suggestion des événements sous jacents à la corrélation entre le cadre spatial, temporel, social et les événements. Ensuite, nous présentons une approche visant à enrichir les tags identités des personnes lors de la consultation dans la section 6.4. De ce fait, les noms de personnes présents dans les photos sont enrichis par des dimensions sociales. Ces dimensions s'adaptent selon le consultateur.

6.2. Motivations

La vie d'un individu peut être vue comme un cycle d'événements qui se répètent. La fréquence de ces événements est variable. Par exemple, pour certains, la fréquence de l'événement '*aller au cinéma*' est de l'ordre d'une fois par semaine, le lieu est : '*Gaumont Pathé de Cordeliers, Lyon*', la fréquence de l'événement '*faire des vacances au camping*' est de l'ordre d'une fois par an, la durée étant d'un mois pendant le mois d'août. Cette fréquence de cooccurrence des événements avec le contexte spatio-temporel peut être exploitée pour suggérer les tags qui font référence à l'événement dans les documents multimédias socio-personnels.

Les documents multimédias socio-personnels de chaque utilisateur sont annotés par des événements qui ocurrent dans son monde réel avec des personnes appartenant à son réseau social. En fait, les membres d'un même réseau social sont, soit des membres d'une même famille, des amis, des voisins ou des collègues, *etc.* En effet, une personne est présente sur ses photos avec des membres de son réseau social lors des événements qui peuvent, parfois se répéter plusieurs fois avec les mêmes personnes. Par exemple, deux utilisateurs amis, peuvent participer à des événements ensemble; ils peuvent faire du "*shopping*" ensemble au "*centre commercial de la Part-Dieu*". Pour l'utilisateur "*Carol*" le "*shopping*" et l'utilisateur "*May*" et le "*centre commercial de la Part-Dieu*" vont ensemble. En outre, l'événement "*shopping*" à la localisation "*centre commercial de la Part-Dieu*" peut se reproduire plusieurs fois. Par conséquent, une certaine cooccurrence existe entre un événement et une catégorie sociale³⁷ comme dans l'exemple de l'utilisateur "*Carol*" l'événement "*shopping*" coïncide avec la catégorie sociale "*Friend*". Aussi, nous pouvons associer le contexte fréquent de déroulement de l'événement "*shopping*" (i.e. qui est ici "*centre commercial de la Part-dieu*") avec la catégorie sociale "*Friend*". Cette cooccurrence entre contexte et catégorie sociale et la cooccurrence entre contexte et événement peut être utilisée pour suggérer des événements pour des photos non annotées ou pour suggérer les identités de personnes dans les photos. Pour un utilisateur ayant peu de photos, une photo non annotée peut le devenir en analysant l'historique des photos d'autres personnes appartenant à son réseau de connaissance.

De surcroît, les utilisateurs sont intéressés, en visualisant les photos des membres de leur réseau social, de connaître l'identité des personnes dans les photos. Dans la pratique, les

³⁷ Les catégories sociales sont définies dans le cadre de l'ontologie *SocialSphere*

utilisateurs se rappellent des personnes avec lesquelles ils ont des relations directes ou indirectes par une dimension sociale (i.e. chaîne de liens sociaux qui peut être par exemple l'amie de ma mère, la femme de mon professeur, etc.) définie dans la section 5.3.3 Chapitre 5) et non par leurs noms et prénoms. Comme ces dimensions sociales dépendent du consultateur de la photo, les tags identités de photos doivent être adaptés en fonction du consultateur. Dans ce cadre, la cooccurrence entre le contexte (localisation et temps) et les catégories sociales peut être utilisée pour faciliter le processus d'enrichissement des tags identités.

Les règles d'associations sont utilisées pour corréler le contexte de prise de vue avec l'événement et avec les catégories sociales et fournir des suggestions d'annotations contextuelles basées sur les liens sociaux. Les règles d'association peuvent être de deux catégories différentes: (i) explicites si elles sont spécifiées manuellement par l'utilisateur ; (ii) implicites si elles sont extraites à partir de données. Nous utilisons les règles de la première catégorie pour associer un contexte aux catégories sociales. Quant à la deuxième catégorie, nous l'utilisons pour associer le contexte à un événement.

Au cours d'un processus d'annotation de photos, nous appliquons des règles d'association produites à partir des collections de documents multimédias socio-personnels de l'annotateur pour suggérer une liste de tags événement. Pour un annotateur qui n'a pas assez de collections de photos, nous proposons un algorithme pour sélectionner les règles d'association préparées par la ou les personnes ayant la valeur la plus élevée de proximité sociale avec l'annotateur. Nous proposons également un algorithme qui adapte la visualisation des tags identité des personnes dans une photo en fonction de la relation du consultateur avec les personnes qui apparaissent sur la photo.

6.3. Recommandation des événements pour annoter des photos

Les événements vécus par une personne sont souvent récurrents. Cette répétition implique une corrélation de divers éléments de contexte. La détection de la corrélation d'un événement fréquent avec un contexte donné peut être effectuée au moyen d'outils de fouille de données comme la fouille des règles d'associations. Les règles d'association produites ont pour vocation de suggérer à l'annotateur des annotations de type événement.

L'extraction de règles d'association, étant une méthode d'apprentissage non supervisé, permet de découvrir à partir d'une base de données (i.e. l'ensemble des lignes de données), des règles d'association [112]. Une règle d'association est une implication conditionnelle

entre ensembles d'éléments appelés items. Une ligne de données est une succession d'items (i.e. un élément de la ligne) exprimée selon un ordre donné.

En général, une règle d'association comporte deux parties: (i) une partie prémisse (appelée aussi antécédent) et (ii) une partie conséquence (appelé aussi conclusion). La partie prémisses d'une règle décrit la condition qui doit être accomplie pour que la partie conséquence soit valide. Les règles d'association sont attachées à deux valeurs : la confiance et le support. Le support d'un ensemble X d'items (noté $Supp(X)$) est la fréquence d'apparition simultanée de l'ensemble d'items X considérant tout le nombre de lignes de données. L'ensemble d'items considéré doit être supérieur à un seuil ($MinSup$). Cet ensemble est appelé itemset fréquent. Soit :

$$Supp(X) = \frac{|X|}{|B|} \geq MinSup$$

Tel que $|X|$ représente le nombre de fois où X apparaît dans une transaction et $|B|$ est le nombre de lignes de données de la base considérée. Le $MinSup$ désigne le seuil l'acceptation de $Supp(X)$.

Quant à la confiance de la règle, elle décrit la probabilité qu'un item ayant la propriété spécifiée par la partie conséquence de la règle fait que l'item correspond à la partie prémisses. Soit :

$$Conf(X \Rightarrow Y) = \frac{Supp(X, Y)}{Supp(X)} \geq MinConf$$

Tel que le $Supp(X \cup Y)$ est le support de X et de Y et le $Supp(Y)$ est le support de Y . Pour accepter une règle, il faut que la valeur de la confiance accordée à cette règle soit supérieure ou égale à une valeur $MinConf$ appelée valeur minimale de confiance.

L'application d'algorithmes de production de règles d'association permet d'extraire des motifs récurrents dans les annotations de photos. Nous considérons les annotations des photos comme la base de données sur laquelle nous effectuons la fouille de règles d'association. Ce processus général d'extraction, illustré par la Figure 29 permet de générer des règles d'association pour chaque utilisateur à partir des annotations. Comme illustré sur la Figure 29, le processus d'extraction des règles d'association peut être décomposé en trois étapes: (i) la sélection et la préparation des données. (ii) la découverte des itemsets fréquents, (iii) la génération des règles d'associations.

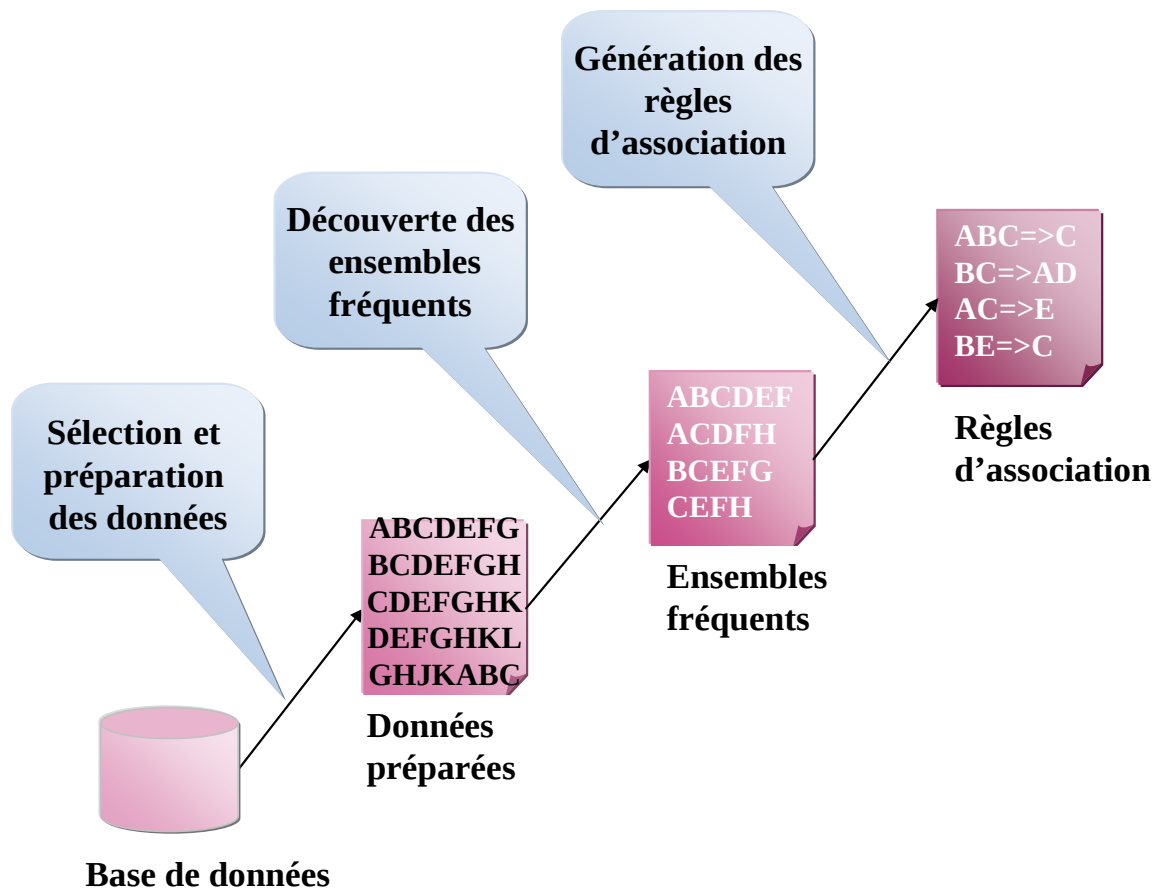


Figure 29. Étapes d'extraction de règles à partir de données

Dans notre cadre d'application, nous procédons à une construction de règles d'association à partir d'une base où chaque transaction représente un événement et le contexte de prise de vue de cet événement.

6.3.1. Nature des données considérées

Dans cette section nous utilisons comme type d'items (appelé aussi propriété) objets, le lieu, le temps et les événements des photos pour produire des règles d'association afin de suggérer des événements. Une ligne de données est décrite par les propriétés: objet, lieu, temps, événement. Chaque propriété est décrite par un ou plusieurs items. Le lieu, par exemple, est décrit par : le nom de l'arrondissement ou ville et le pays. Quant à l'objet, il est décrit par:

- des noms de monuments comme "tour Eiffel", l'"Arc de triomphe" (i.e. donnée annotée automatiquement ou manuellement),

- les noms d'objets identifiés par l'utilisateur (e.g. maison personnelle, maison des parents, travail, etc.). Ces objets sont représentés par le tag (i.e. le nom de l'objet) et l'identifiant de l'utilisateur qui l'a tagué.

Pour des raisons de clarté, les identifiants des utilisateurs sont remplacés par des noms dans les exemples d'instances représentés dans le Tableau 5. La troisième propriété de ces données représente le temps. Ce dernier décrit par trois attributs: moment de la journée "*timeOfDay*", jour de la semaine "*dayOfWeek*" et le mois "*month*". La dernière propriété représente l'événement. Il est décrit par le nom d'un événement générique "*Event*" comme : "*aller au cinéma*", "*visiter un lieu*", etc. ou un événement spécifique à une personne comme l'"*anniversaire de Carol*".

Nous rappelons que le moment de la journée "*timeOfDay*" correspond à une valeur appartenant à l'ensemble suivant {"*Afternoon*", "*Early morning*", "*Evening*", "*Late night*", "*Morning*" et "*Night*"}. Pour les définir nous nous sommes basés sur les travaux de Naaman et al. [113]. De même, le jour de la semaine "*dayOfWeek*" représente les 7 jours de la semaine soit une valeur appartenant à l'ensemble suivant {"*Sunday*", "*Monday*", "*Tuesday*", "*Wednesday*", "*Thursday*", "*Friday*", "*Saturday*"}. Quant au mois "*month*" il correspond à une valeur parmi les 12 mois de l'année soit une valeur parmi l'ensemble suivant {"*January*", "*February*", "*March*", "*April*", "*May*", "*June*", "*July*", "*August*", "*September*", "*October*", "*November*", "*December*"}.

Le Tableau 5 représente quelques lignes de données de la base de données considérées.

Tableau 5. Exemples de lignes des données considérées pour générer les règles d'associations

	Object	Location	Time	Event
1	Tour Eiffel	Paris, France	Evening, Saturday, July	trip
2	parent's home Salma	Sfax, Tunisie	Late night, Friday, March	celebrating birthday
3	City hall	Lyon, France	Afternoon, Friday, April	wedding
4	souk Rachel	Elkantawi, Tunisie	Evening, Friday, July	trip
5	Amphitheatre of El Jem	el jem, Mehdiya, Tunisie	Evening, Friday, July	trip
6	parent's home Salma	Sfax, Tunisie	Late night, Tuesday, August	celebrating birthday
7	Roman Theatre of Carthage	Carthage, tunis, Tunisie	Night, Saturday, August	concert
8	Job Marco	Villeurbanne,	Morning, Monday,	professional

		Lyon, France	January	meeting
9	home Sihem	Villeurbanne, Lyon, France	Night, Saturday, August	engagement
10	Place Bellcour	Lyon1, France	Night, Saturday, August	festival of light
11	Place jean Jaurès	Lyon 7, France	Night, Saturday, August	music festival

Nous avons besoin d'un outil de fouille de données qui permet d'extraire des règles d'association à partir d'un grand ensemble de données produisant des règles de forme booléenne. La forme booléenne donne seulement des règles valides ou invalides. Pour l'exemple présenté dans la section 6.2, un événement '*shopping*' est valide pour le contexte '*centre commercial de la Part-Dieu*'.

L'algorithme classique APRIORI [114] et [115] répond à ces enjeux du fait qu'il permet d'extraire les ensembles et les sous ensembles fréquents (appelé aussi itemsSet fréquents) à partir d'une grande base de données et il permet d'extraire des règles d'association booléennes.

Nous avons appliqué l'algorithme APRIORI [114] et [115] pour la découverte des règles d'association. Cet algorithme procède en deux phases (i) La première étape consiste à découvrir les itemsSet fréquents en considérant un seuil *MinSup* prédéfini. Cette étape est effectuée en plusieurs itérations. En effet, les 1-itemSet sont extraits en premier lieu. Ce sont ceux qui contiennent un seul item. Par exemple, ceux qui se réfèrent à la propriété localisations fréquentes dans la base. Seuls les 1-itemSets qui satisfont un seuil de support minimum (*MinSup*) sont considérés. Ils sont appelés, par exemple, L1. Puis, une jointure est appliquée sur les L1 afin de trouver les 2-itemSets. Les 2-itemSets sont ceux qui contiennent à la fois 2 items. Par exemple, ceux qui se réfèrent aux propriétés localisations et les temps fréquents. Ce processus est itéré jusqu'à atteindre des k-itemSets c'est-à-dire ceux qui contiennent des items événements (ii) La deuxième étape consiste à générer des règles en considérant un seuil minimal de support *MinSup* et un seuil minimal de confiance *MinConf* prédéfini. Nous considérons uniquement un ensemble prédéfini de combinaison de règles qui sera expliqué dans la section suivante.

6.3.2. Découverte des motifs fréquents

Le calcul de fréquence est effectué sur les propriétés *Object*, *Time*, *Location* et *Event* afin d'extraire les itemSets fréquents.

L'algorithme d'extraction de caractéristiques calcule le support de chaque ensemble fréquent, le support et la confiance d'un ensemble des items pour retrouver des itemSets fréquents. Le support est un indicateur permettant d'indiquer si un item se référant à l'une des propriétés précédemment citées est fréquent ou pas. L'algorithme récupère, dans une première étape, tous les items de *ContextEvent* de chaque groupe social où la mesure est supérieure à un seuil minimal.

$$FreqSet = \{ \cup item, Supp(item) \geq MinSupp \}$$

Le modèle de *contexteEvent* dont une instance est illustrée dans le Tableau 5, comporte une localisation exprimée en *NearbyUserObject*, *NearbyObject*, *NamePlace* et *Country* co-occurente avec un temps (Time) exprimé en *timeOfDay*, *dayOfWeek* et *month* avec un événement (Event) exprimé en *nameEvent*.

Les itemSets fréquents sont, par exemple : *evening*, *july* et *visiting place*. Lorsque ces données co-occurrent fréquemment (à un support égal à 0.2, par exemple), ils peuvent construire une règle. Par exemple, les itemSets fréquents dans le Tableau 5 supérieur ou égal à $MinSupp=0.2$ sont :

$FreqSet=\{evening, july, trip\}$ support: 0.27 (ligne 4 et ligne 5)

$FreqSet=\{parent's home Salma, Sfax, Tunisie, Late night, Celebrating Birthday\}$ support: 0.2 (ligne 2 et ligne 6)

6.3.3. Découverte des règles d'association

A partir des itemSets fréquents avec un support $\geq MinSup$, la génération des règles d'association pour un seuil de confiance *MinConf* est un problème qui dépend exponentiellement de la taille de l'ensemble des itemSets. Il faut, donc, combiner les itemSets fréquents *FreqSets* et générer des règles d'association.

L'extraction des règles d'association de deux ou plusieurs itemSets fréquents *FreqSets* crée des règles avec une confiance élevée (supérieure à un seuil *MinConf*). La confiance peut être interprétée comme une estimation de la probabilité conditionnelle de la partie conséquence sachant la partie prémisses.

Soit la confiance : $Conf(X \rightarrow Y) = Supp(X,Y) / Supp(X)$

Nous utilisons les symboles suivant : *O* désigne un Objet, *T* désigne un temps, *L* désigne une Localisation et *E* désigne un Événement.

Nous avons extrait des règles d'association ayant l'événement (*E*) comme partie conséquence. La partie prémisse peut correspondre à plusieurs propriétés qui sont :

localisation (L), à un objet (O), à un temps (T) ou n'importe quelle combinaison de 2 ou 3 paramètres.

Plus formellement, soit $FreRules$ l'ensemble des règles d'association extraites.

$FreRules =$

$FreRules_T \cup FreRules_L \cup FreRules_O \cup FreRules_{T \cup L} \cup FreRules_{T \cup O} \cup FreRules_{L \cup O} \cup FreRules_{S_{L \cup O \cup T}}$

Tel que :

L'ensemble des règles mettant en jeu le temps est défini comme:

$FreRules_T = \{ \cup (T \Rightarrow E), \text{Supp}(T \cup E) \geq \text{MinSupp}, \text{Supp}(T \cup E) / \text{Supp}(E) \geq \text{MinConf} \} ;$

Où

Par exemple:

$FreRules_L = \{ \cup (L \Rightarrow E), \text{Supp}(L \cup E), \text{Supp}(L \cup E) / \text{Supp}(E) \geq \text{MinConf} \} ;$

Représentent les règles dont la partie prémisse contient seulement des propriétés Localisation (L).

$FreRules_{T \cup L} = \{ \cup (T \cup L \Rightarrow E), \text{Supp}(T \cup L \cup E), \text{Supp}(T \cup L \cup E) / \text{Supp}(E) \geq \text{MinConf} \} ;$

Représentent les règles dont la partie prémisse contient à la fois des propriétés Temps (T) et Localisation (L)

$FreRules_{T \cup O} = \{ \cup (T \cup O \Rightarrow E), \text{Supp}(T \cup O \cup E), \text{Supp}(T \cup O \cup E) / \text{Supp}(E) \geq \text{MinConf} \} ;$

Représentent les règles dont la partie prémisse contient à la fois des propriétés Temps (T) et Objets (O).

$FreRules_{L \cup O} = \{ \cup (L \cup O \Rightarrow E), \text{Supp}(L \cup O \cup E), \text{Supp}(L \cup O \cup E) / \text{Supp}(E) \geq \text{MinConf} \} ;$

Représentent les règles dont la partie prémisse contient à la fois des propriétés Localisation (L) et Objets (O).

et

$FreRules_{L \cup O \cup T} = \{ \cup (L \cup O \cup T \Rightarrow E), \text{Supp}(L \cup O \cup T \cup E), \text{Supp}(L \cup O \cup T \cup E) / \text{Supp}(E) \geq \text{MinConf} \}$

Représentent les règles dont la partie prémisse contient à la fois des propriétés Localisation (L), Objets (O) et Temps (T).

Exemple:

{evening, july} => {trip} support: 0.55 confiance: 1.0

{my parent's home Salma, Sfax, Tunisie, Late night} => {Celebrating Birthday} support: 0.44 confiance: 1.0

Les règles produites sont rejetées si la valeur du support est inférieure à une constante prédéfinie appelée support minimal (*MinSupp*) et si, la valeur de la confiance est inférieure à une constante prédéfinie mais paramétrable appelée confiance minimale (*MinConf*).

6.3.4. Suggestion de tag événement basée sur la proximité sociale

Les règles d'association sont construites à partir de collections de photos de chaque utilisateur. Ces règles d'association peuvent être utilisées pour suggérer des tags 'événement' à un utilisateur en situation d'annotation ou de consultation. Cependant, il est possible que l'utilisateur ne dispose que de très peu de photos qui coïncident avec le contexte de la photo en cours d'annotation. Par conséquent, on peut avoir peu ou pas de règles qui correspondent à un contexte donné avec un support et une confiance nécessaires. Ce problème est connu sous le nom de démarrage à froid (cold start problem) dans les systèmes de recommandation [116]. Pour remédier à ce problème, nous proposons d'utiliser les règles produites pour des utilisateurs ayant une haute valeur de proximité sociale avec l'annotateur.

La proximité sociale noté *ProxSocial* (u_1, u_2) est, comme présenté dans la Définition 5, une fonction qui retourne une valeur entre 0 et 1 : $[0,1]$. La proximité sociale décrit le degré de connaissance entre deux utilisateurs u_1 et u_2 . Si deux utilisateurs sont proches socialement (i.e. ont une haute valeur de proximité sociale), il y a une forte probabilité qu'ils partagent (ou repartagent) le même événement. La proximité sociale n'est pas commutative i.e. $proxSocial(u_1, u_2) \neq proxSocial(u_2, u_1)$.

La valeur de proximité sociale peut être déterminée via trois paramètres :

- Le nombre de cooccurrence sur les photos (le nombre de fois où les deux utilisateurs sont tagués ensemble sur une même photo) ;
- Le nombre de leurs collaborations dans l'annotation manuelle des photos (i.e. le nombre de fois où l'utilisateur u_1 a commenté ou annoté les photos de u_2) ; et
- Le nombre de fois où l'utilisateur u_1 a partagé avec u_2 des photos.
- Les liens explicites entre les nœuds du réseau social.

Donc, pour calculer la valeur de proximité sociale, nous proposons la fonction multicritères ci-dessous :

$$proxSocial(u_1, u_2) = \alpha. freqPart(u_1, u_2) + \beta. freqAnn(u_1, u_2) + \delta. freqCoocPh(u_1, u_2) + \lambda. N(u_1, u_2)$$

Tel que:

- $freqPart(u_1, u_2) = \frac{NbPart(u_1, u_2)}{NbTPart(u_1)}$ avec $NbPart(u_1, u_2)$ est le nombre de photos

partagées par l'utilisateur u_1 avec l'utilisateur u_2 et $NbTPart(u_1)$ est le nombre total de partage de photos de u_1 avec tous les utilisateurs³⁸.

- $freqAnn(u_1, u_2) = \frac{NbAnn(u_1, u_2)}{NbTAnn(u_1)}$ avec $NbAnn(u_1, u_2)$ est le nombre de fois où u_1 a

annoté/commenté les photos de u_2 et $NbTAnn(u_1, p)$ est le nombre totale où u_1 a annoté/commenté les photos des autres utilisateurs. Le nombre de commentaires/annotations est considéré et pas le nombre de photos annotées. Ce qui veut dire que si un utilisateur u_1 a commenté deux fois une photo de l'utilisateur u_2 nous considérons que son nombre d'annotation sur cette photo est égal à 2.

- $freqCoocPh(u_1, u_2) = \frac{NbCoocPh(u_1, u_2)}{NbCoocPh(u_1)}$ avec $NbCoocPh(u_1, u_2)$ est le nombre de

fois où u_1 et u_2 sont tagués ensemble dans la même photo et $NbCoocPh(u_1)$ est le nombre total où u_1 a été tagué avec d'autres utilisateurs.

- $Nb(p)$: une fonction qui retourne l'ensemble des voisins directs à une personne p

- $N(u_1, u_2) = 1$ si une relation explicite est exprimée de u_1 à u_2 et 0 sinon.
- α, β, δ et λ sont des coefficients de pondération.
- $\alpha + \beta + \delta + \lambda = 1$

Dans le Tableau 6, nous décrivons des exemples de valeurs de proximité sociale. Les cellules foncées sur le Tableau 6 montrent la valeur de la proximité sociale la plus élevée pour chaque utilisateur représenté par une colonne. Par exemple, pour Carol, May est la plus proche socialement d'elle avec une valeur de $proxSocial(Carol, May) = 0.9$. Pour Alice, Carol est la plus proche socialement d'elle avec une valeur de $proxSocial(Alice, Carol) = 0.5$.

³⁸ Ici on entend par "partage" l'utilisation de l'application mobile pour le partage avec les pairs au moyen de Bluetooth. Cette information est tracée envoyée à l'application web avec les informations contextuelles en même temps que le chargement de la photo.

Tableau 6. Proximité sociale entre Alice, Carol, Bob et May

	Carol	Alice	Bob	May
Carol		0.5	0.4	0.6
Alice	0.4		0.001	0.03
Bob	0.3	0.008		0.55
May	0.9	0.08	0.5	

Nous proposons un algorithme de recommandation de tag événement, appelé RSP (Recommendation via Social Proximity), qui prend en considération la proximité sociale. L'algorithme se base sur le principe suivant :

Supposons qu'un utilisateur u tente d'annoter la photo ph . Le contexte de prise de vue de la photo ph est c . Un seuil *seuil* de nombre de recommandations est prédéfini. Supposons, aussi, que P est l'ensemble de personnes du réseau social de u ayant la proximité sociale supérieur à une valeur *PrSeuil*.

Définition 11. Règles d'association spécifiques à un utilisateur

Les règles d'association correspondant au contexte c sont recherchées dans l'ensemble de règles préparées pour l'utilisateur u . En effet, l'extraction des règles d'association est appliquée pour chaque utilisateur dans la base ; ce qui signifie que l'ensemble des données considéré lors de l'extraction d'un ensemble de règles d'association ne considère qu'un seul utilisateur à la fois. Si le nombre de recommandations résultant des règles préparées pour l'utilisateur u est inférieur à *seuil*, des règles d'associations produites pour les personnes appartenant à l'ensemble P sont utilisées afin de produire une liste de recommandations dont le nombre est suffisant.

L'algorithme RSP de recommandation de tags événement, utilise un ensemble de fonctions définies comme suit :

- $R(p, c)$: une fonction qui retourne les règles préparées pour la personne p correspondant à un contexte c . En effet, la comparaison entre les éléments de contexte (i.e. temps, localisation et objets) est effectuée par correspondance exacte.

- $SelectSuggestion(R, k)$: une fonction qui retourne une liste S de suggestions correspondant aux parties conséquence de l'ensemble des règles R tel que $|S| \leq k$. La liste de suggestion est triée par maximum de support et/ou confiance.

- $Sort(P)$: une fonction qui trie, par ordre décroissant de la valeur de proximité sociale, un ensemble P d'utilisateurs.

Algorithme 1: Algorithme RSP de recommandation de tag événement via la proximité sociale

Algorithme RSP(u, c)
//recommandation via social proximity

Entrées
 u : utilisateur en cours d'annotation
 c : le contexte de prise de vue de la photo en cours d'annotation
Variables Locales
 $RSeuil$: seuil de recommandation considéré
 $Pfile$: une file des utilisateurs ayant une proximité sociale avec u
 $PsSeuil$: seuil de considération de la proximité sociale
Sorties
 $Recom$: liste de suggestion de tags événements
Début
1. $P = \{proxSocial(u, p) \geq PsSeuil\}$
2. $Pfile \leftarrow Sort(P)$
3. $Recom \leftarrow SelectSuggestion(R(u, c), RSeuil)$
4. **Tant que** ($|Recom| < RSeuil \ \& \ Pfile \neq \emptyset$) **faire**
5. $p \leftarrow defiler(Pfile)$
6. $Recom \leftarrow SelectSuggestion(R(p, c), RSeuil - |Recom|)$
7. **Fin Tant que**
8. **Retourner** $Recom$
Fin algorithme

Le déroulement de l'Algorithme 1 sur un exemple est comme suit : partant du contexte, supposons qu'il correspond à $c = \langle evening, Part \ Dieu \rangle$ de l'utilisateur $u = Carol$, l'algorithme *RSP* (*Recommendation via social proximity*) illustré par Algorithme 1 crée une file d'utilisateurs ayant une proximité sociale supérieure à un seuil $PsSeuil$. Supposons que $PsSeuil = 0.2$. La file ($Pfile$) est triée par ordre décroissant de valeur de proximité sociale. Supposons que la file ($Pfile$) contiendra après l'exécution de la *ligne 1* de l'algorithme *May*, *Alice* et *Bob* dans cet ordre. La liste de recommandation d'événements ($Recom$) (*ligne 2*) prend initialement la partie conséquent des règles produites à l'avance pour l'utilisateur $u = Carol$ qui correspondent au contexte $c = \langle evening, Part \ Dieu \rangle$. Le nombre de recommandations considéré ne doit pas dépasser un seuil ($RSeuil$). Supposons que $RSeuil = 5$ et que l'exécution de la *ligne 2* de l'algorithme a produit un seul élément dans la liste $Recom = \{shopping\}$. Dans ce cas, la partie de l'algorithme (*ligne 3* à la *ligne 6*) va s'exécuter tant que le nombre d'éléments dans la liste $Recom$ n'atteint pas $RSeuil = 5$ ou la file $Pfile$ d'utilisateurs soit vide.

De *ligne 3* à la *ligne 6*, à chaque itération on sélectionne la partie conséquence de règles pré-préparées pour l'utilisateur à l'entête de la file *Pfile*. Les règles sélectionnées correspondent au contexte *c*.

L'algorithme retourne la liste $Recom=\{shopping, going\ to\ the\ cinéma, party, celebrate\ Carol's\ birthday\}$

La liste de recommandation retournée par l'Algorithme 1 permet d'assister l'annotation des événements des documents multimédias socio-personnels.

L'algorithme proposé est de complexité $O(|P|. Avg(R(u, c)))$ avec $|P|$ est le nombre d'éléments dans la file *Pfile* et $Avg(R(u, c))$ est la moyenne de temps d'exécution nécessaire pour retrouver les règles d'associations correspondant à un contexte *c*.

Néanmoins, pour un contexte donné, on trouve peu de règles par rapport à l'ensemble de règles de chaque utilisateur donc la moyenne $Avg(R(u, c))$ est négligeable.

6.3.5. Discussion

Il existe différents types de règles d'association [117]. La forme simple celui qui montre que l'association est valide ou invalide. Ce type est de nature booléenne et est appelé règles d'association booléenne. Toutes les règles que nous avons produites appartiennent à ce type de règles.

Le deuxième type de règles est celui qui permet d'agréger plusieurs règles d'association dans une seule règle. Ce type est appelé '*règles d'association à plusieurs niveaux*' ou '*règles d'association générales*'. Ces règles impliquent généralement une hiérarchie et la fouille ramène un concept de niveau supérieur à partir d'un concept plus spécifique. Par exemple, si nous considérons la règle 'Night→concert', l'événement concert est issue d'une hiérarchie (e.g. WordNet) d'événements ayant comme concept plus générique 'social event'. Dans d'autres cadres d'applications comme le commerce électronique l'extraction des règles d'association génériques peuvent être très utiles (e.g. extraire une règle en rapport avec une marque spécifique de lait peut être généralisé sur tous les marques de lait). Dans notre cadre d'application, nous avons besoin de recommander des événements spécifiques. Pour cette raison une extraction de règles à plusieurs niveaux n'est pas utile pour nous ce qui explique notre choix d'extraction de règles booléennes.

La comparaison entre les éléments de contexte dans les règles d'association et le contexte actuel de la photo en cours d'annotation ou de consultation mérite plus d'attention. En effet, pour le moment nous nous basons sur la correspondance exacte entre les éléments de contexte. Une comparaison considérant la similarité sémantique selon

WordNet pour les objets par exemple ou des localisations selon GeoNames pourra considérablement améliorer la pertinence des suggestions retournées par l'algorithme *RSP*.

6.4. Enrichissement des tags identité lors de la consultation des photos

L'enrichissement des annotations selon le consultateur a pour vocation de caractériser sémantiquement les annotations que nous appelons "*dimension sociale*". En effet, plusieurs études ont démontré que les utilisateurs se souviennent des photos en se rappelant principalement de l'identité des personnes qui y figurent [118]. Etant donné que les photos sont sensées raconter une histoire, celle-ci commence toujours par les noms des personnes qui y sont impliquées. Comme : *[personne]* et *[personne]* font *[quelque chose]*. L'histoire est racontée d'un point de vue personnel, c'est-à-dire une *[personne]* est toujours exprimée en *[nom Personne]*, mon *[dimension sociale]* etc. Une *dimension sociale* (Définition 6) correspond, donc, à une suite de relations sociales entre le consultateur d'un document multimédia socio-personnel et les acteurs de la photo (une instance de la classe *foaf : person* ayant une relation acteur avec la photo en cours de consultation).

Les dimensions sociales ajoutées pour chaque consultateur donnent plus de sens aux annotations pour le consultateur. En effet, selon le consultateur, nous avons proposé une approche qui permet de retrouver des dimensions sociales potentiellement intéressantes. Ces dimensions correspondent aux plus courts chemins qui relient le consultateur aux personnes identifiées sur la photo (tag identité).

Le processus d'enrichissement des tags "*identité de personnes*" exploite la puissance des ontologies pour inférer intelligemment des relations sociales. En effet, lors de l'ajout d'une nouvelle relation sociale, le système applique les règles d'inférences afin d'inférer et stocker dans la base d'autres relations. Le principe de l'inférence en appliquant des règles d'inférences implicites (i.e. propres à l'ontologie comme la transitivité, la symétrie, etc.) ou explicites (i.e. que nous avons mis en œuvre pour permettre l'inférence de relation de sens commun) est expliqué plus en détail dans la section 5.3.3.2 Chapitre 5. La recherche des dimensions sociales pour chaque tag décrivant l'identité d'une personne est effectuée dans les profils sociaux. Plus spécifiquement dans les liens sociaux de chaque personne spécifiée à l'aide de l'ontologie "*SocialSphere*" définie dans la section 5.3.3 Chapitre 5. Le consultateur de la photo consulte les différentes propositions du système pour chaque

identité de personne décrite dans la photo. Il est libre d'accepter ou de refuser les suggestions. Nous appelons "*tag identité*" une annotation précisant le nom d'une personne présente sur la photo et sa dimension sociale (relation à plusieurs niveaux) avec celle qui annote (*[Dimension sociale] [nomActeur]*). Le tag identité est calculé en temps réel lors de l'annotation et/ou la consultation des photos.

La recherche de dimensions sociales est effectuée dans les profils sociaux partant du profil du consultateur de la photo. Le réseau parcouru est représenté d'une manière distribuée c'est-à-dire que l'ensemble des liens sociaux de chaque utilisateur est représenté dans son profil social. En parcourant l'ensemble des liens sociaux dans chaque profil social nous accédons à un graphe : c'est le réseau social. Le réseau social est, donc, un graphe composé de nœuds qui représentent les utilisateurs et d'arcs qui représentent les liens sociaux.

La Figure 30 illustre deux cas d'enrichissement de tags identité de personnes dans la photo. Le cas de *Bob* dont l'interface est représentée dans la partie gauche de la Figure 30 et le cas de *May* dont l'interface est représentée dans la partie droite de la Figure 30. Les dimensions sociales proposées à l'utilisateur *Bob* sont différentes des dimensions sociales proposées à l'utilisateur *May*. La dimension sociale dépend, en effet, des chemins entre le consultateur et les personnes identifiées sur la photo dans le réseau social.

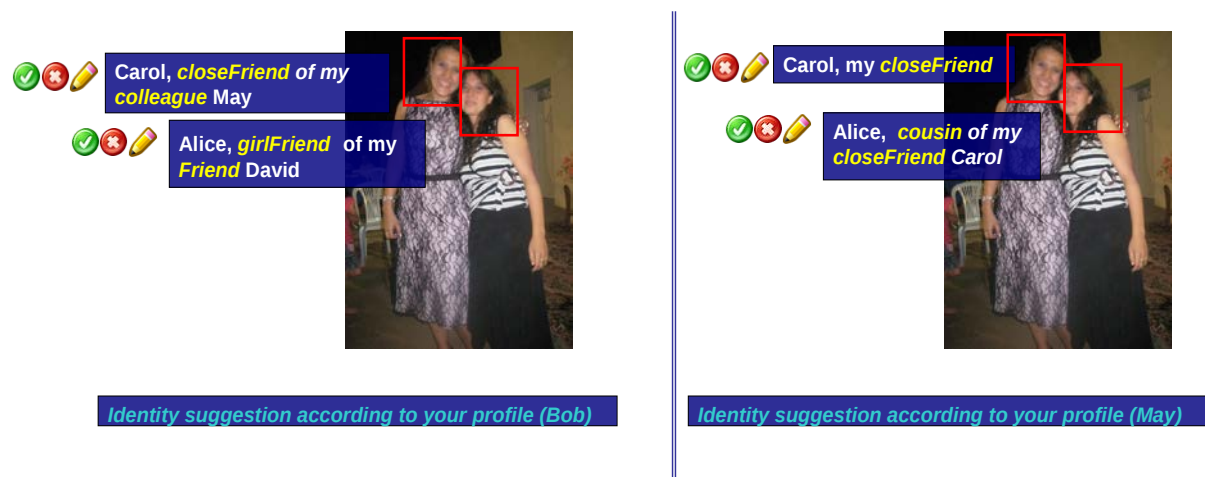


Figure 30. Enrichissement des annotations avec dimension sémantique pour deux utilisateurs

Si nous considérons le cas de l'enrichissement des tags identités des personnes dans les photos par des dimensions sociales pour le consultateur *Bob*, la recherche de ces dimensions est effectuée en partant du profil social de *Bob*. La Figure 31 représente le réseau social parcouru partant du nœud qui représente *Bob*. Le nœud *Bob* est le nœud

initial noté v_i . Les nœuds *Nader*, *Salwa*, *May*, *David* et *Ahmed* sont les nœuds qui représentent les $N(v_i)$ c'est-à-dire les liens sociaux ayant un lien direct avec Bob.

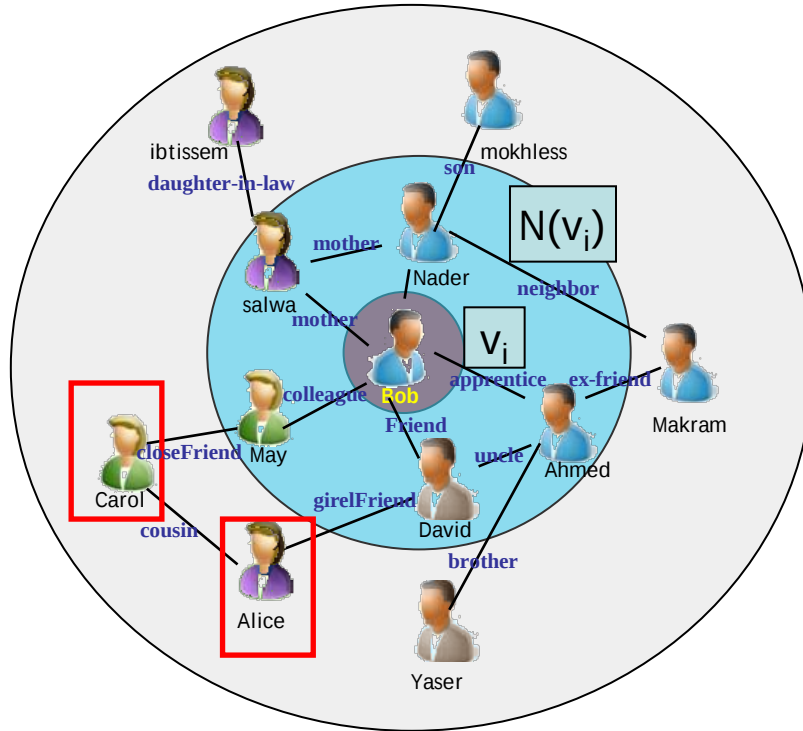


Figure 31. Exemple de réseau social parcouru depuis le profil social de Bob

Le réseau social risque d'atteindre une très grande taille vue qu'il contient un grand nombre de nœuds et de relations. Par exemple, le graphe du réseau social Facebook a atteint en 2010 plus de 500 millions de nœuds. Partant d'un nœud à la recherche d'un autre, nous courons le risque de parcourir grand nombre de nœuds. La recherche dans le réseau social peut alors être très coûteuse en terme de temps d'exécution si, par exemple, le nœud source v_i et le nœud destination (i.e. représentant l'identité de la personne dans la photo) noté v_d sont très éloignés. Afin d'accélérer le processus de suggestion des dimensions sociales associées aux tags identités de personnes dans les photos, nous proposons d'accélérer le temps de réponse de chemins dans le réseau social.

6.4.1. Heuristiques considérées

Afin de réduire le temps d'exécution d'une requête de recherche de chemins dans le réseau social, nous utilisons trois heuristiques basées sur les hypothèses suivantes :

- Algorithme de recherche en largeur plutôt qu'une recherche en profondeur: le choix de l'algorithme de recherche en largeur est justifié par le fait qu'il est plus probable de trouver des acteurs d'une photo appartenant à des contacts directement liées. Notre heuristique se base alors sur l'algorithme classique de recherche en largeur BFS (Breadth First Search) [119].
- La recherche dans le réseau ne dépasse pas deux niveaux: ce choix est justifié par une étude empirique menée par notre collègue Omar Hasan en 2009 [120]. L'étude qui a été menée sur le site *Advogato*³⁹ a constaté sur 40 000 de nœuds, 36 578 nœuds reliées à d'autres par deux niveaux de séparation seulement et 6 739 nœuds sont reliées à d'autres par 3 niveaux de séparation.
- Prise en considération du contexte pour le filtrage de la recherche: Le contexte joue un rôle important dans la réduction de l'espace de recherche. En effet, ayant un rôle '*témoin d'une photo*', nous pourrions estimer les potentiels catégories sociales avec lesquelles il est probable qu'il se retrouve dans une localisation L ou un temps T . Cette heuristique n'est appliquée donc que dans le cas où le consultateur est décrit comme participant dans la prise de vue c'est-à-dire relié à la photo objet d'observation avec la relation témoin "*Witness*". La description de la relation de type "*rôle*" entre l'utilisateur et la photo est décrite dans la section 5.3.2 Chapitre 5. En appliquant cette heuristique, la recherche se centre sur un ensemble limité de catégories sociales. Pour cela, nous nous basons sur des règles d'associations qui associent un contexte spatio-temporel à une catégorie sociale. Ces règles sont spécifiées manuellement par l'utilisateur mais ils peuvent être fouillées à partir de l'historique de chaque utilisateur. Les règles d'association considérées ont la forme suivante:

$$L \vee T \Rightarrow socialCategory, T \Rightarrow socialCategory, \text{ ou } L \Rightarrow socialCategory$$

C'est-à-dire la partie prémisses est Localisation (L) et Time (T) ou Localisation (L) ou Time (T) et la partie conséquence est une catégorie sociale (*socialCategory*).

³⁹ www.advogato.org/.

Par exemple, la règle $\text{'Evening'} \Rightarrow \text{'Friendly relationship'} \wedge \text{'Professional relationship'}$ spécifiée pour un consultateur signifie que si le contexte de la photo sujet d'observation est $T = \text{'Evening'}$ alors la recherche va être limitée à ses connaissances qui appartiennent à ses catégories sociales amies et familles qui sont respectivement $\text{socialCategory} = \text{'Friendly relationship'}$ et $\text{'Professional relationship'}$ et tous les liens directs de ces catégories (en considérant la deuxième heuristique).

6.4.2. Algorithme SQO

Avant de détailler l'heuristique d'optimisation, nous définissons quelques fonctions, le réseau social G_S ⁴⁰ défini par l'ensemble des nœuds N_S et l'ensemble des liens sociaux R_S : $G_S (N_S, R_S)$:

On désigne deux fonctions sur R_S :

- $C:(v_I, v_D) \rightarrow \text{RelationCategories}$, qui définit la catégorie de la relation entre v_I et v_D .

Par exemple : $C(\text{Bob}, \text{May}) = \text{'Professional relationship'} \in \text{RelationCategories}$.

- $R:(v_I, v_D) \rightarrow \text{SocialRelations}$, qui définit la relation qui relie v_I à v_D . Par exemple : $R(\text{Bob}, \text{May}) = \text{colleague} \in \text{SocialRelations}$.

- $\text{GetConcernedCategory}(v_I, \text{Contexte})$ une fonction heuristique renvoie la/les catégorie(s) sociale(s) concernée(s) par le nœud v_I à partir du contexte spatio-temporel. Exemple : $\text{GetConcernedCategory}(\text{'Bob'}, \text{'Evening'}) = \text{'Friendly relationship'}$ ou $\text{'Professional relationship'}$.

L'exécution de l'algorithme de recherche en largeur avec heuristique baptisé SQO (*Social Query Optimization*) qui utilise les fonctions et les heuristiques mentionnées plus haut est expliquée par le pseudo code présenté dans l'Algorithme 2.

Partant du réseau social G_S et deux nœuds v_I et v_D (initiale et destination), l'algorithme SQO retourne un ensemble de propositions qui correspondent à des chemins potentiels qui relient v_I et v_D . Par exemple : pour Bob : Carol peut être la *cousine (cousin)* de la *petite amie (girl friend)* de son ami (*Friend*) David. Comme elle peut être la *meilleure amie (closeFriend)* de sa collègue (*colleague*) May. La variable *SocialDimension* représente les

⁴⁰ La définition sera présentée formellement dans Définition 14

différentes solutions. Quelques conditions sont nécessaires pour l'extraction des dimensions sociales :

- $C(v_I, v_D)$ donne des valeurs non nuls à chaque arête $(v_I, v_D) \in R_S$.
- Il existe une relation de dépendance entre R et C : si par exemple $R = \textit{SiblingOf}$ alors $C = \textit{Family relationship}$. La dépendance entre les deux fonctions est expliquée dans la section 5.3.3 Chapitre 5).
- Chaque nœud v_i possède une fonction $N(v_i)$ renvoyant tous les nœuds voisins partant de v_i . Le nombre maximum retourné par cette fonction est noté $\mathbf{b}$. Au pire des cas, $\mathbf{b} = n - 1$ (n est le nombre de nœuds du réseau social), ce qui signifie que chaque nœud est relié à tous les autres.

Algorithme SGO (Social Query Optimization)**Entrées** $G_S(N_S, R_S)$: le réseau social v_I : nœud initial // l'observateur de la photo v_D : nœud destination // un tag qui référence une personne identifiée dans la photo

Context: liste d'éléments de contextes // les éléments de contexte de la photo

Sorties: socialDimension liste de listes // tous les chemins possibles entre v_I et v_D **Variable locales**

nodesCovered liste, parentState liste, q file, concernedCategory list, state nœud

Initialization $q \leftarrow \emptyset$, state $\leftarrow v_I$, socialDimension $\leftarrow \emptyset$, $j \leftarrow 0$, parentState $\leftarrow \emptyset$, nodesCovered $\leftarrow \emptyset$ **Début**// recherche du nœud v_D 1. **Répéter**2. **Si** $v_I \in N(state)$ || state = v_I **alors**3. concernedCategory $\leftarrow$ **GetConcernedCategory**(state, Context)4. **Pour chaque** c dans concernedCategory **faire**5. enfile v dans q telque $v \notin q \ \&\& \ \exists c(v, state)$ 6. **Fin Pour**7. **Fin Si**

8. enfile state dans NodesCovered

9. state $\leftarrow$ q.defiler()10. **Jusqu'à** state = v_D || q.vide()// si le nœud v_D n'a pas été retrouvé11. **Si** state $\neq v_D$ **alors**12. retourner v_D 13. **Fin Si**// construire socialDimension pour v_I 14. ParentState $\leftarrow \{ \forall v / v \in N(state) \ \&\& \ v \in nodesCovered \}$ 15. **Pour chaque** v dans ParentState **faire**16. **Répéter**

17. enfile state dans SocialDimension[j] // j est le nombre de chemins

18. **Si** $v_I \in N(state)$ **Alors**

19. enfile 'My' dans SocialDimension[j]

20. enfile R(state, v_I) dans SocialDimension[j]21. state = v_I 22. **Sinon**

23. enfile R(state, v) dans SocialDimension[j]

24. state = v

25. **Fin Si**26. **Jusqu'à** state = v_I

27. incrementer j

28. state = v_D 29. **Fin Pour****Fin Algorithme****Algorithme 2. Algorithme SGO de recherche en largeur avec heuristique d'optimisation**

La première partie de l'algorithme est consacrée à la recherche du nœud v_D sur deux niveaux partant du nœud $state = v_I$ ou $state = N(state)$: les nœuds adjacents au nœud $state$ (ligne 2). La variable *concernedCategory* prend les nœuds adjacents au nœud $state$ ⁴¹ retourné par la fonction heuristique *GetConcernetCategory* (ligne 3). Tous les nœuds à examiner sont placés dans la file q (ligne 5). Les nœuds sont examinés partant de la tête de la file (ligne 8 à ligne 10) jusqu'à ce que la variable $state$ soit égale au nœud recherché (v_D). Dans le cas contraire, (ligne 11 à ligne 13) l'algorithme retourne v_D .

La deuxième partie de l'algorithme est consacrée à la construction de tous les chemins entre les deux nœuds v_D et v_I une fois que v_D retrouvé. La construction de ces chemins est effectuée par l'intermédiaire d'une liste de listes appelée *socialDimension*. La variable *socialDimension* (ligne 17 à ligne 28) prend tous les chemins possibles⁴² de v_D à v_I en remontant aux parents de chaque nœud.

6.4.3. Analyse de la complexité de l'algorithme SQO

Considérons b le nombre maximal de voisins directs à un nœud v , m la profondeur maximale du réseau social c'est-à-dire la distance (en termes de nombre de liens sociaux) qui relie les nœuds les plus éloignés du réseau et n le nombre de nœuds dans le réseau social.

Nous analysons la complexité de l'algorithme "Social Query Optimisation" (SQO) dans les deux cas suivants : le meilleur cas c'est-à-dire lorsque le nœud v_D se trouve dans les liens directs de v_I : $v_D \in N(v_I)$ et le pire cas c'est-à-dire lorsque le nombre de nœuds adjacents à chaque nœud v du réseau est égale à $n-2$ ($|N(v)|=n-2$) et que le nœud $v_D \notin N(v_I)$.

- **Le meilleur cas**

Dans le meilleur cas, la complexité de la recherche des dimensions sociales sera égale de l'ordre de $O(b)$. Si b tend vers $n-1$ c'est-à-dire le nœud initial v_I est adjacent à tous les nœuds du réseau, alors on aura une complexité d'ordre $O(n)$ ou plus précisément $O(n-1)$. Dans le cas général, la complexité de l'algorithme BFS est de l'ordre de $O(b^m)$. Dans ce

⁴¹ $state$ prend le nœud en cours d'exploration est égal, initialement, au nœud v_I

⁴² Les chemins possibles les plus courts ne dépassant pas le niveau dont lequel nous avons retrouvé v_D dans l'étape de la recherche de v_D .

meilleur cas, la complexité de l'algorithme SQO est égale à la complexité de l'algorithme BFS car le premier nœud v_I a une distance=1 avec le nœud le plus éloigné du réseau.

- **Le pire cas**

Dans le pire cas, v_D ne fait pas partie des nœuds adjacents à v_I (i.e. $v_D \notin N(v_I)$) et le nombre de nœuds adjacents à chaque nœud v du réseau est égal à $n-2$ ($|N(v)|=n-2$). Dans ce cas, la complexité de l'algorithme SQO est égale à $O(n^2)$ car le parcours des nœuds s'arrête à deux niveaux. Quant à la complexité de l'algorithme BFS, elle est égale à $O(n^m)$.

En conclusion, la complexité de la recherche des dimensions sociales dans les deux cas est raisonnable ; elle est égale au pire des cas $O(n^2)$. Cela est très encourageant puisque nous ne pouvons affronter un très gros ensemble de données ce qui est le cas généralement dans les réseaux sociaux.

6.4.4. Discussion

L'heuristique de recherche des dimensions sociales dans le réseau social est basée sur trois hypothèses (i) favoriser la recherche en largeur plutôt qu'une recherche en profondeur ; (ii) ne pas dépasser une distance de recherche égale à deux ; (iii) prise en considération du contexte pour filtrer l'espace de recherche en largeur. Cette dernière heuristique n'a un sens que si le consultateur représenté par le nœud v_I a un rôle témoin "Witness" dans la photo c'est-à-dire que nous avons capturé l'identifiant de son dispositif mobile lors de la prise de vue de la photo objet de consultation. En pratique, cette dernière heuristique n'apporte pas un gain considérable puisqu'elle ne fait sens que pour filtrer les nœuds adjacents à v_I .

6.5. Résumé

Dans ce chapitre, nous avons présenté des contributions en rapport avec l'assistance de l'utilisateur au cours de la consultation ou au cours de l'annotation. Nous avons présenté trois approches; une première qui consiste à assister l'utilisateur à l'annotation des tags événements; une deuxième vise à apporter plus de sémantique aux personnes identifiées sur les photos en proposant au consultateur une liste de dimensions sociales. Ces objets sont issus de son profil social et celui des participants à la prise de vue.

Le chapitre suivant présente d'autres moyens permettant d'exploiter le graphe social pour la recherche de documents multimédia

Chapitre 7. Exploitation du graphe social

7.1.	SCENARIO INTRODUCTIF.....	125
7.2.	UNE CONCEPTUALISATION ABSTRAITE DU GRAPHE SOCIAL	126
7.3.	GRAPHE SOCIAL	128
7.3.1.	Arcs.....	129
7.3.2.	Nœuds	129
7.3.3.	Attributs d'un nœud	129
7.3.4.	La liste des attributs pour chaque nœud du graphe social	131
7.4.	GRAPHE MOTIF POUR L'EXPRESSION DES REQUETES	132
7.4.1.	Nœuds du graphe motif.....	133
7.4.2.	Attributs des nœuds du graphe motif.....	133
7.4.3.	Fonctions de comparaison du graphe motif.....	133
7.4.4.	Les arcs du graphe motif.....	134
7.5.	INSTANCIER LES GRAPHE MOTIFS.....	135
7.5.1.	Isomorphisme de sous-graphe partiel.....	136
7.5.2.	Notations et définitions	138
7.5.3.	Algorithme d'isomorphisme de sous-graphe partiel.....	140
7.5.3.1.	Construction du domaine D_x initial.....	142
	• Heuristique de réduction de l'espace de recherche.....	143
	• Construction du domaine D_x initial	145
7.5.3.2.	Propagation des contraintes de différence FC_{Diff}	146
7.5.3.3.	Propagation des contraintes d'arrêt FC_{Edges}	147
7.5.3.4.	Propagation des contraintes d'arrêt AC_{Edges}	148
7.5.3.5.	Considération de la similarité entre nœuds pour filtrer les domaines de valeurs.....	150
7.5.3.6.	Analyse de la complexité de notre algorithme h-Pruning	150
7.5.4.	Exemple complet.....	151
7.6.	RESUME ET DISCUSSION	157

7.1. Scénario introductif

Nous décrivons, dans cette section, un scénario motivant la proposition de notre approche d'exploitation de graphe social : Un groupe anglophone prépare une surprise à *Carol* à l'occasion de son anniversaire à la maison de son amie *May*. Lors de la soirée, *Bob*, un des collègues de *May* prend des photos. *Bob* annote l'ensemble des photos prises durant la soirée par "*celebrate birthday*". Supposons que lors de la prise de vue, *Bob* utilise une application déployée sur son mobile qui permet de décrire et d'enrichir automatiquement des métadonnées contextuelles basées sur notre modèle de données décrit dans la section 5.4 Chapitre 5. Le module GPS du téléphone mobile de *Bob* retourne les coordonnées géographiques précises du dispositif au moment de la prise de vue de la photo. Au cours de la soirée, le groupe rencontre une personne : *David* l'ami de *Bob*. À l'issue de cette rencontre, *David* partage via Bluetooth⁴³ quelques photos avec *Bob* et les publie par la suite, sur son site de réseau social préféré.

Supposons qu'après un certain temps, *David* souhaite revoir ces photos. Il lance alors une requête : "*May home*", "*dance party*". Nous considérons, alors, trois possibilités : 1) la recherche de photos est textuelle, elle utilise seulement les mots-clés ajoutés manuellement, 2) le moteur de recherche utilise l'approche de MediaAssist [91], [90] (présentée dans la section 4.6.1) et, 3) le moteur de recherche utilise l'approche de PhotoMap [121], [122] et [123]. Dans le premier cas, le système retourne comme réponse à la requête soumise, des photos annotées par les mots-clés "*May*" ainsi que par "*home*" et/ou "*party*". Dans l'approche de MediaAssist, qui se base sur le contexte, les résultats manquent de précisions puisque le système fournit également des photos prises lors du mois de Mai (en anglais, "*May*"), une confusion apparaît, puisque le système considère que "*May*" est non seulement un prénom mais aussi un indicateur temporel. PhotoMap propose une méthode d'annotation en considérant des catégories de métadonnées : *David*, annote alors avec précision les photos, autrement dit, pour chaque objet à annoter, il précise le type des mots-clés (i.e. l'endroit "*May home.where*" ou l'événement "*party.what*"). Selon cette approche, les résultats ne sont toujours pas précis, car le système PhotoMap n'est pas capable d'interpréter "*May home*" comme indicateur spatial si cet indicateur n'est pas annoté

⁴³ www.bluetooth.com/

manuellement par *David* sur la carte géographique. De plus, les photos que *David* souhaite récupérer, sont annotées par *Bob*, et ceci par l'événement "*celebrate birthday*". L'indicateur de l'événement (i.e., *what*) ne correspond pas à celui de la requête simplement, parcequ'il ne s'agit pas du même utilisateur.

Supposons maintenant que le système incorpore pendant le processus d'annotation des métadonnées qui décrivent à la fois le contenu et le contexte, et ce par l'intermédiaire d'une méthode semi-automatique (le processus est décrit dans le Chapitre 6). Ces métadonnées permettent de relier les informations des métadonnées des photos avec d'autres métadonnées en se basant sur les informations disponibles sur le réseau social. Le système crée, lors de l'annotation, une instance du *graphe social*. Lorsque *David* formule dans sa requête un indicateur d'événement, qui est dans notre cas "*Dance party*", le système compare ce mot-clé à d'autres décrivant également d'autres photos. Dans ce contexte, le système utilise une similarité sémantique de thesaurus (par exemple, le thesaurus WorldNet). Ainsi, selon la similarité de ce thesaurus, "*birthday*" et "*party*" ont une similarité non nulle. L'indicateur spatial de la requête est spécifié par un champ indicateur spatial (*home*) relié à une personne (*May*) pour déterminer que l'indicateur spatial de la requête correspond à un marqué (tagué) par *May*.

7.2. Une conceptualisation abstraite du graphe social

Commençons par rappeler que, parmi les défis soulevés dans le premier chapitre, on distingue la capacité du système à retrouver des informations relatives au contenus, notamment, les acteurs et les liens sociaux (e.g. *P1* offre un cadeau à *P2*, *P1* félicite *P2*, etc.). En effet, des études psychologiques ont montré que les interactions sociales sont un des principaux canaux par lesquels nous comprenons la réalité [8]. De plus, on espère par la modélisation du graphe social représenter des informations relatives au contexte de création dont l'importance pour la recherche des documents multimédias socio-personnels a été prouvée dans les études empiriques de Naaman et al. [113] et [99]. D'autres études menées par Rodden et Wood [9] ont prouvé l'importance du contexte, exprimé de manière personnalisée, pour la recherche des photos personnelles. Ces éléments de contexte concernent des événements (e.g. 'mon anniversaire', 'mariage de mon voisin', etc.), des personnes (e.g. 'moi', 'mon fils',

‘ma mère’, etc.) et des places/localisations (e.g. ‘ma maison’, ‘lieu de travail de mon frère’).

Pour prendre en considération ces différents éléments, nous proposons un modèle abstrait appelé « Graphe social ». Ce modèle est une version simplifiée de *SeMAT* et *SocialSphere* présentés respectivement dans la section 5.3.2 et la section 5.3.3 du Chapitre 5. Cette version appelée modèle de graphe social contient les nœuds et les relations qui seront exploitables par l’approche d’appariement de graphes détaillée dans la section 7.5. Les concepts clés définis dans cette conceptualisation abstraite du graphe social sont (voir niveau modèle de la Figure 32):

SPMD: un concept permettant de représenter un document multimédia socio-personnel manipulé par l’utilisateur.

Person: un concept central du modèle. Ses instances permettent de représenter, d’une part, le contexte social des documents multimédias socio-personnels qui est leur repère de rappel le plus important [92] ; d’autres parts, le représenter en tant qu’utilisateur du document multimédia socio-personnel.

Role: un concept permettant de représenter le rôle d’une personne dans un document multimédia socio-personnel. Les spécifications du rôle peuvent être extraites à partir d’un vocabulaire dédié. Dans notre modèle, nous avons distingué notamment trois rôles à savoir : acteur, témoin, créateur.

SocialLink: un concept permettant d’exprimer les liens sociaux entre des personnes. Ces liens sont classés en deux catégories : les relations sociales et les interactions sociales. Les instances de relations sociales sont issues de l’ontologie *SocialSphere* (présentée dans la section 5.3.3 Chapitre 5). Par exemple, en se référant à la Figure 32, on spécifie la relation sociale *mother* entre Alice et Carol.

DocumentActivity: un concept permettant de donner un sens à l’interaction de la personne avec les documents multimédias socio-personnels. Les instances de ce concept peuvent être : *annotation*, *visualization*, *search*, etc. Quant il s’agit d’annotation, l’attribut *tag* est obligatoire. Le *tag* est renseigné librement par l’utilisateur (e.g. à la manière de Flickr). Cependant, nous pouvons identifier la catégorie du tag en utilisant, par exemple, des catégories Wordnet (voir la section 8.3.4 Chapitre 8).

Context: un concept permettant de représenter les éléments du contexte de la prise des photos. Nous avons distingué trois dimensions du contexte, à savoir : le contexte spatial, le contexte temporel et le contexte spatio-temporel. Quand un nœud de type *Context* est défini, des spécifications '*Spec*' TemporalElement, SpatialElement ou SpatioTemporalElement sont à renseigner.

Group : un concept dont les instances sont issues d'un thésaurus de termes. Le nom d'un groupe est à définir par un utilisateur. Plusieurs utilisateurs peuvent appartenir à un groupe.

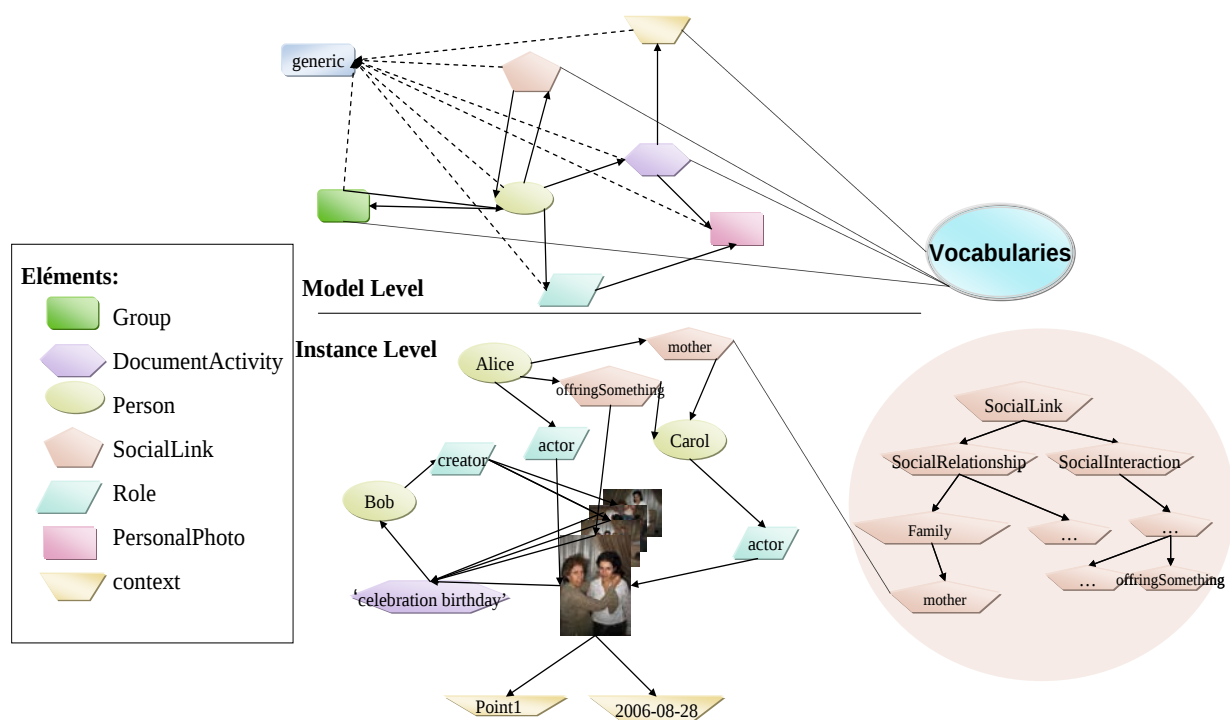


Figure 32: Modèle du graphe social à deux niveaux

Le niveau instance de la Figure 32 représente des instances concrètes du modèle de graphe social. Le niveau instance est appelé " graphe social".

7.3. Graphe social

Définition 12.Graphe Social

Un graphe social noté G_S est un graphe contenant des instances du modèle conceptuel du contenu social. G_S est un graphe orienté et connexe contenant différents types de nœuds. Pour assurer la connexité du graphe, nous avons représenté un nœud

générique permettant de relier tous les nœuds de G_S . Formellement, un graphe social est défini par:

$$G_S = \langle N_S, R_S \rangle$$

Tel que :

- N_S est l'ensemble des nœuds dans G_S
- R_S est l'ensemble des arcs r qui sont des relations binaires, orientées et non étiquetées.

7.3.1. Arcs

Formellement, un arc $r \in R_S$. R_S est l'ensemble d'arcs possibles selon le type de nœud qui sont :

(Group-Person), (Person-Group), (SocialLink-Person), (Person-SocialLink), (Person-DocumentActivity), (DocumentActivity-SPMD), (Person-Role), (Role-SPMD)..

7.3.2. Nœuds

Formellement, un nœud $n_S \in N_S$, est défini par le tuple :

$$n_S = \langle ID_S, Att_S \rangle$$

Tel que :

- ID_S correspond à un identifiant d'un nœud dans le graphe social G_S .
 - Att_S est un ensemble d'attributs
- L'ensemble Att_S comporte des attributs obligatoires et uniques: $Type_S$ et $Spec_S$

7.3.3. Attributs d'un nœud

Formellement, un attribut d'un nœud noté $ar \in Att_S$, est défini par :

$$ar = (AttributName, value, Multiplicity)$$

- <i>AttributeName</i> le nom de l'attribut - <i>Value</i> la valeur de l'attribut - <i>Multiplicity</i> est la multiplicité des attributs qui peut être unique (1) ou multiple (* (zéro ou plusieurs)). Exemple d'<i>AttributeName</i> : - $Type_S$ correspond au type du nœud dans le graphe social. Nous avons distingué l'ensemble de types suivants {Person, SPMD, SocialInteraction, SocialRelationship, DocumentActivity, Role, Group, Context} - $Spec_n$ correspond à la spécification du nœud dans G_S. Exemple de <i>AttributeName=Value</i> : - $Type_n = 'SocialInteraction'$ $Spec_n = 'offerSomething'$. Un nœud peut avoir des attributs non obligatoires et multivalués (voir Tableau 7)

L'exemple de la Figure 32, représente une instance du graphe social G_S . Dans cet exemple, nous représentons une photo créée et annotée par Bob par les annotations '*celebrate*' et '*birthday*'. La création de la photo est : (1) représentée par sa dimension temporelle, illustrée par le nœud de type contexte avec l'étiquette '2006-08-28' ayant les attributs : $Type = 'Context'$, $Spec = 'TemporalElement'$, $TimeStamp = '2006-08-28'$ $ID_S = 1$ (2) par sa dimension spatiale, illustrée par le nœud de type contexte avec l'étiquette = 'Point1' et les attributs : $Type = 'Context'$, $Spec = 'SpacialElement'$, $ID = 'Point1'$. Nous représentons, également, dans cette photo, Alice et Carol identifiées comme actrices avec une interaction sociale: *OfferingSomething*. Alice est la mère de Carol.

La Figure 33 représente quelques nœuds du graphe social avec les détails relatifs à leurs attributs. Par exemple, le nœud de $ID_S = 'http://liris.cnrs.fr/annotating1234'$ possède les attributs $a_1 = (Type_S, 'Documentactivity')$, $a_2 = (Spec_S, 'annotation')$, $a_3 = (Tag, 'celebrate')$ et $a_4 = (Tag, 'birthday')$.

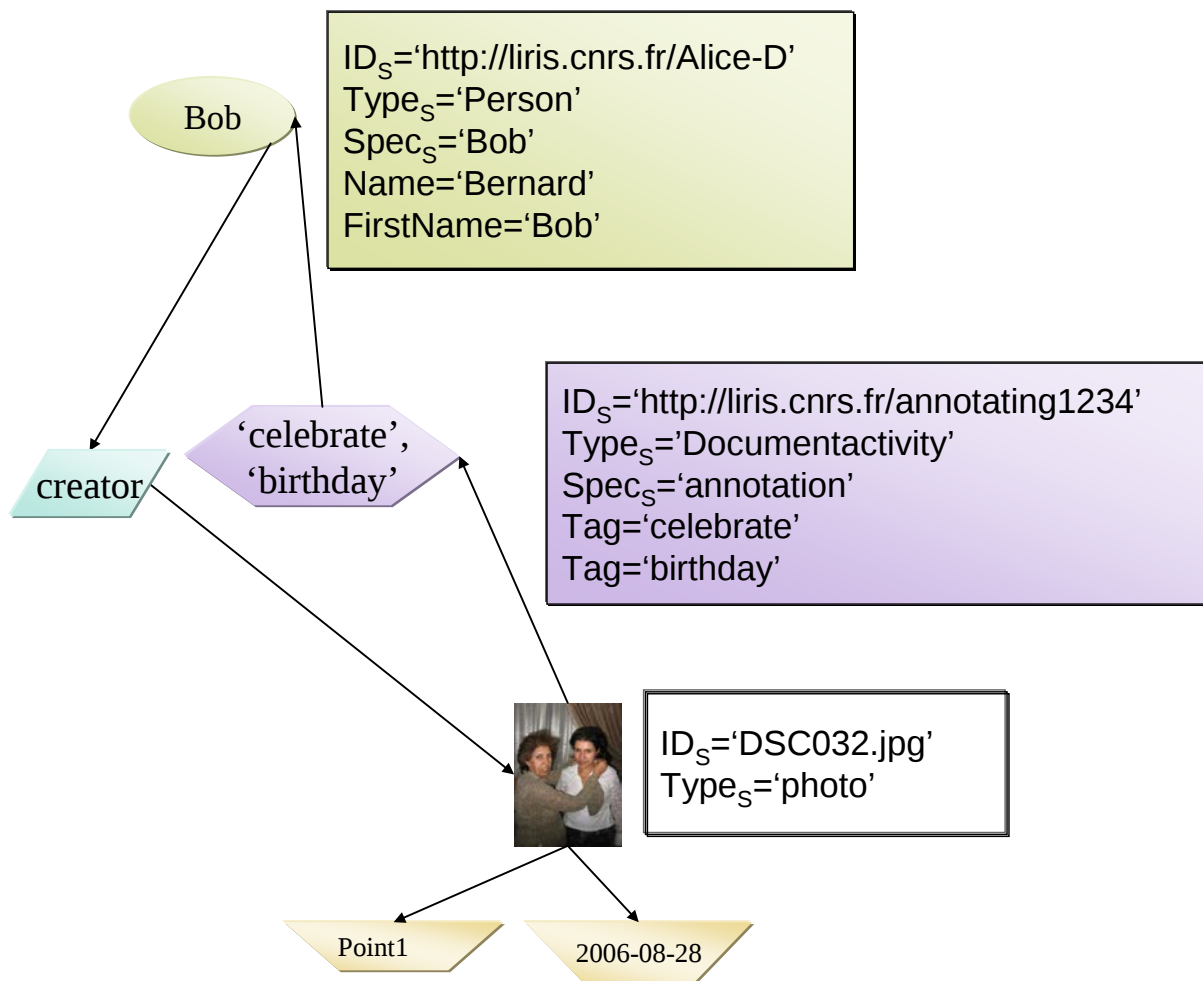


Figure 33. Définition de nœuds avec un ensemble d'attributs et d'instances

7.3.4. La liste des attributs pour chaque nœud du graphe social

La liste de noms des attributs est définie pour chaque type. Cependant, la liste que nous précisons est flexible : l'utilisateur n'est pas obligé de renseigner tous les attributs. La liste peut aussi être étendue. Le Tableau 7 représente la liste des attributs Att_n pour chaque type de nœud.

Tableau 7. Types Att_n (et leur multiplicité $*$ =0 ou n) pour chaque spécification de nœuds

Spécification du nœud	Liste des noms (et multiplicité) d'attributs associés
Person	FirstName(*), Name(1)
SPMD	URI(1)
SocialInteraction	/
SocialRelationship	/
DocumentActivity	Tag(*)
Role	/
Group	Tag(*)
SpatialElement	Tag(*), Lat(1), Long(1), Alt(1), Town(1), Country(1)
TemporalElement	TimeStamp(1), Day(1), Month(1), Year(1), DayOfWeek(1), DayOfMonth(1), PeriodOfDay(1)
Spatio-temporalElement	Season(1), Weather(1), LightStatus(1)

7.4. Graphe Motif pour l'expression des requêtes

Nous présentons, dans cette section, notre outil d'expression de requêtes appelé graphe motif G_M . Cet outil permet d'interroger le graphe social.

Définition 13. Graphe Motif

Un graphe motif exprime une requête. Il doit être instancié dans un graphe social. La notion est une extension de celle du graphe potentiel définis par Egyed et Prié [124], [125] dans laquelle des nœuds et des arcs peuvent être apparentés à des nœuds d'un graphe global. Les auteurs ont défini le concept de nœud générique "*un nœud pouvant être instancié par des nœuds du graphe global*". Un graphe motif s'instancie, ainsi, dans un graphe social G_S , fournissant les résultats à la requête qu'il représente.

Formellement, un graphe motif est défini par un tuple:

$$G_M = \langle N_M, R_M, F \rangle$$

Tel que :

- N_M est l'ensemble des nœuds génériques dans G_M
- R_M est l'ensemble des arcs r qui sont des relations binaires $\subset N_M \times N_M$, orientées et non étiquetées
- F un ensemble de fonctions de comparaison

7.4.1. Nœuds du graphe motif

Formellement, un nœud générique $n_M \in N_M$, est défini par le triplet :

$$n_M = \langle ID, Att, RefID_n \rangle$$

Tel que :

- ID est l'identifiant qui caractérise le nœud d'une manière distincte des autres nœuds du graphe motif G_M .
- Att est un ensemble d'attributs
- $RefID_n$ correspond à un identifiant d'un nœud concret dans le graphe social G_S .

7.4.2. Attributs des nœuds du graphe motif

Formellement, un attribut $at \in Att$, est défini par :

$$at = (AttributeName, value)$$

- $AttributeName$ le nom de l'attribut
- $Value$ la valeur de l'attribut

Pour les attributs multi-valués et des fonctions comparant 2 valeurs, la comparaison s'effectue entre l'attribut du nœud G_M et tous les attributs du même type de G_S . Suite à cette comparaison, les résultats s'additionnent (ou logique).

7.4.3. Fonctions de comparaison du graphe motif

Formellement, $fc_x \in F$ ou $x \in \{\text{exact, IncludeIn, thesaurus}\}$, est définie par :

$$fc_x : at \times ar \rightarrow [0,1]$$

Par exemple si:

$$a_1 = (\text{FirstName}, \text{'Alice'})$$

$$a_2 = (\text{FirstName}, \text{'Alice'})$$

$$a_3 = (\text{Tag}, \text{'partly Cloudy'})$$

$$a_4 = (\text{Tag}, \text{'partly'})$$

$$a_5 = (\text{Tag}, \text{'event'})$$

$$a_6 = (\text{Tag}, \text{'birthday'})$$

$fc_{\text{exact}}(a_1, a_2) = 1$ signifie que les valeurs des attributs $a_1 \in Att$ et $a_2 \in Att_S$, sont identiques.

$f_{includeIn}(a_4, a_3)=1$ signifie que la valeur de l'attribut $a_4 \in Att$ est incluse dans la valeur de l'attribut $a_3 \in Att_S$.

$f_{thesaurus}(a_5, a_6)=0,53$ signifie que la valeur de l'attribut $a_5 \in Att$ est similaire à la valeur de l'attribut $a_6 \in Att_S$ de 0,53. $f_{thesaurus}$ dénote les mesures de dimilarités sémantiques comme celles présentées dans le Chapitre 3.

7.4.4. Les arcs du graphe motif

Les nœuds génériques définis auparavant sont reliés par un ensemble d'arcs pour former un graphe connexe⁴⁴.

Tout arc $r \in R_M$ représente une relation à la fois binaire et orientée.

Notation. Projection du graphe motif : la projection du graphe motif noté

$\pi(G_M)=S$, tel que S est un sous-graphe du graphe motif ($S \subseteq G_M$).

Définition 14. Réseau Social

Le réseau social est une projection du graphe social sur les nœuds de type personne et des liens sociaux. La Figure 34 représente le graphe motif dont l'instance concrète est un réseau social.

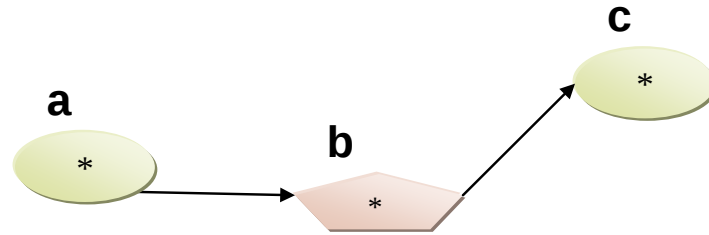


Figure 34. Graphe motif ramenant un réseau social

Exemple

Reprenons le scénario initial : la requête de David pourra être reformulée via le graphe motif G_M de la manière suivante :

Lister toutes les photos photographiées à la maison de May, tel que sur ces photos sélectionnées, figure, la photo de Carol avec un de ses connaissances.

⁴⁴ le graphe social est un graphe obligatoirement connexe.

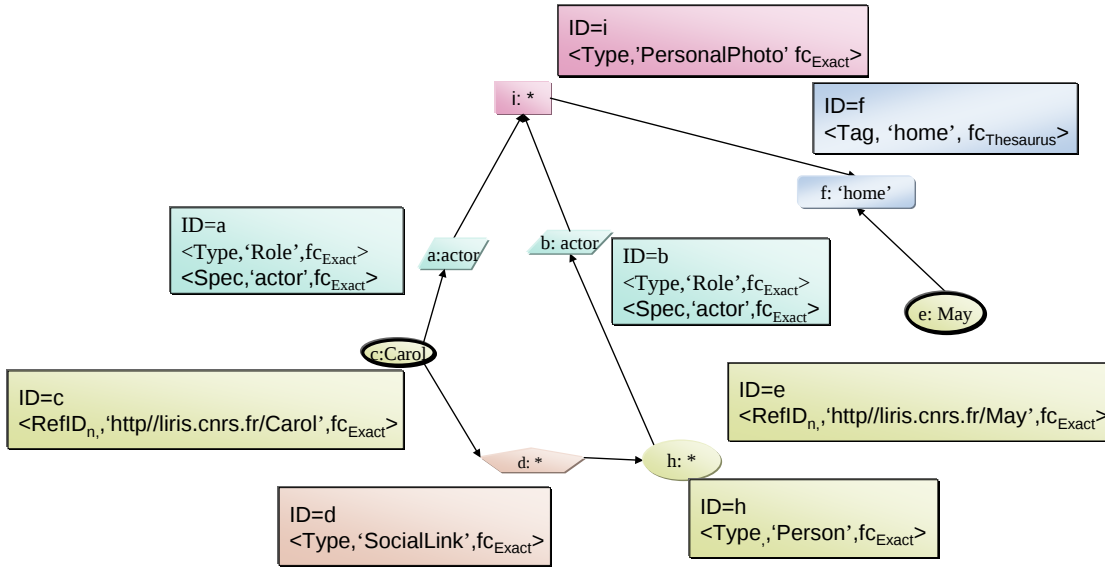


Figure 35. Exemple de graphe motif représentant la requête de David

La Figure 35 illustre la requête de David formalisée en un graphe motif G_M . Les nœuds encadrés en noir représentent des nœuds connus c'est-à-dire dont l'attribut $RefID_n$ est renseigné (i.e. on connaît leur instance unique dans le graphe social G_S). Les rectangles représentent des attributs spécifiés pour quelques nœuds. Par exemple, le nœud ayant l' $ID=f$ et possédant l'attribut $Tag='home'$ doit être comparé avec les nœuds du graphe G_S en utilisant les fonctions de comparaison de thésaurus (notées $fc_{thesaurus}$) définies dans le Chapitre 3.

Afin de répondre à la requête, on instancie le graphe G_M dans G_S en considérant les fonctions de comparaison pour les nœuds de G_M dont les attributs sont renseignés. Pour ce faire, nous avons proposé un algorithme qui permet la résolution de la requête. Cet algorithme est détaillé dans la section 7.5.

7.5. Instancier les graphes motifs

Nous nous intéressons, dans cette section, à l'instanciation des graphes motifs qui représentent des requêtes. Les graphes motifs doivent retrouver leurs instances dans le graphe social G_S . Nous avons ramené le problème d'instanciation du graphe motif à un problème de recherche d'isomorphisme de sous-graphe partiel [126]. La modélisation en graphe et la résolution du problème en recherche d'isomorphisme sont justifiées par la nécessité de prendre en considération la relation sémantique entre les descripteurs et les photos personnelles. Dans un cas général, la complexité théorique des algorithmes de recherche d'isomorphisme de sous-graphe partiel n'est

pas précisément établie : le problème apparaît clairement dans la catégorie NP ; néanmoins, il n'est pas détecté de manière explicite dans les catégories P ou NP-complet mais on ne sait pas encore s'il est dans la catégorie P ou s'il est NP-complet [126].

7.5.1. Isomorphisme de sous-graphe partiel

Le problème d'isomorphisme est fondamental en informatique. On peut le formaliser de la façon suivante :

Définition. Isomorphisme de sous-graphe partiel:

Un graphe G_S est défini par un couple (N_S, R_S) , N_S étant l'ensemble fini de nœuds de G_S et $R_S \subseteq N_S \times N_S$ l'ensemble des arcs de G_S .

Le problème d'isomorphisme de sous-graphe entre un graphe motif $G_M=(N_M, R_M)$ et le graphe cible G_S consiste à décider si G_M est isomorphe à un sous-graphe de G_S .

Il s'agit de trouver une fonction **injective** $f_b : N_M \rightarrow N_S$, qui associe un nœud cible différent à chaque nœud motif et qui préserve les arcs du motif : à savoir, $\forall (x, x') \in R_M$, $(f_b(x), f_b(x')) \in R_S$. La fonction f_b est appelée *fonction de sous-isomorphisme*.

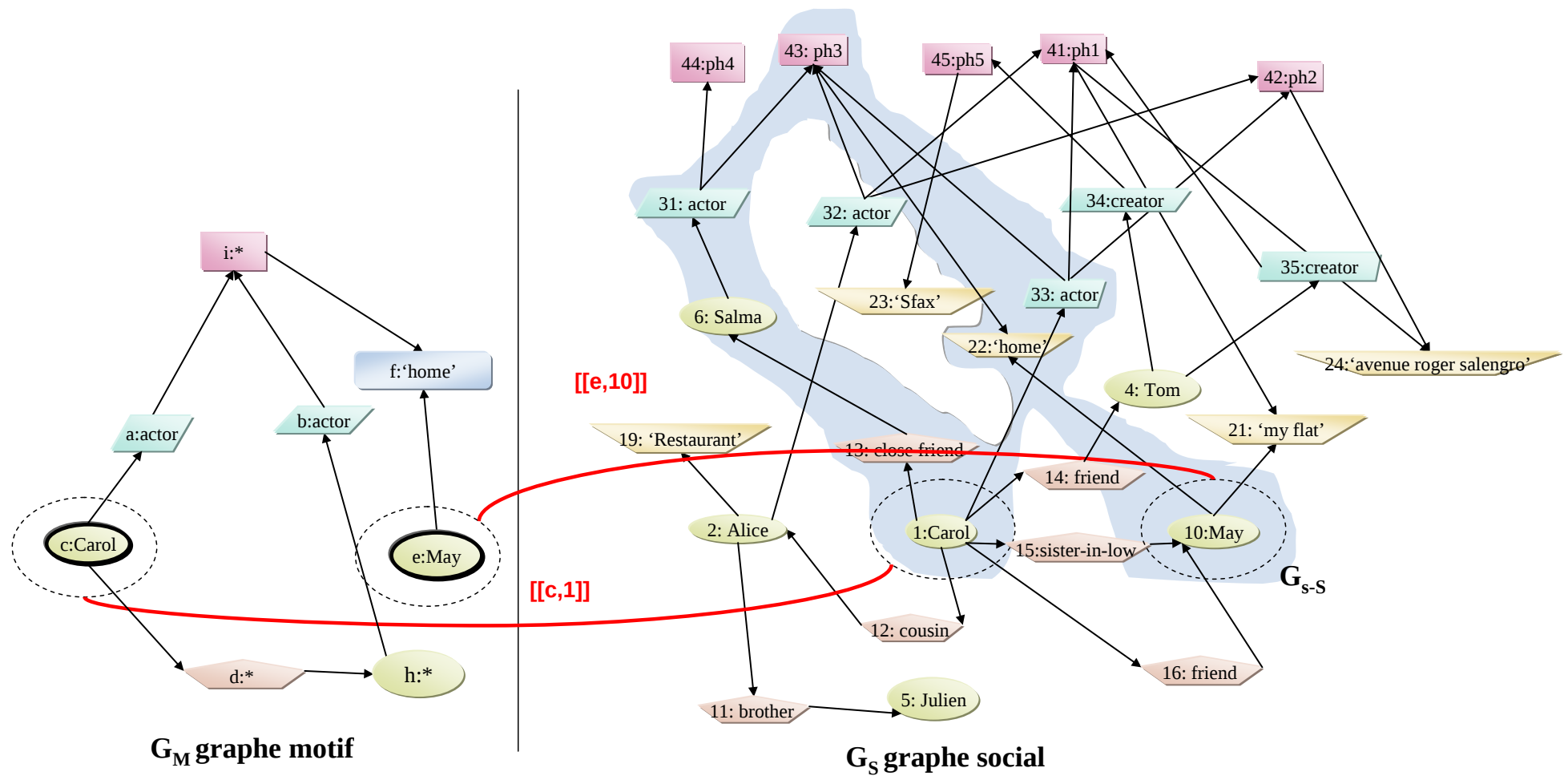


Figure 36. Isomorphisme de sous graphe partiel entre GM et GS

La Figure 36 représente un exemple d'isomorphisme de sous-graphe partiel où le graphe motif G_M est instancié du graphe G_{s-s} . G_{s-s} est un sous-graphe du graphe social G_S . Ainsi, G_M est défini par:

$$N_M = \{ab, c, d, e, f, i, h\}$$

et

$$R_M = \{(i, f), (e, f), (a, i), (b, i), (c, a), (c, d), (d, h), (h, b)\}$$

et G_S est défini par :

$$N_S = \{1, 2, 3, 4, 5, 6, 11, 12, 13, 14, 15, 16, 19, 21, 22, 23, 24, 31, 32, 33, 34, 35, 41, 42, 43, 44, 45\}$$

et

$$R_S = \{(1, 13), (1, 33), (1, 14), (1, 15), (1, 16), (1, 12), (2, 19), (2, 32), (2, 11), (3, 22), (3, 21), (4, 34), (4, 35), (6, 31), (11, 5), (12, 2), (13, 6), (14, 4), (15, 3), (16, 3), (31, 44), (31, 43), (32, 41), (32, 42), (32, 43), (33, 41), (33, 42), (33, 43), (34, 45), (35, 41), (41, 21), (41, 24), (42, 24), (43, 22), (45, 23)\}$$

L'instance de G_M dans G_S est appelé G_{s-s} et est défini par :

$$N_{s-s} = \{1, 3, 6, 13, 22, 31, 33, 43\}$$

et

$$R_{s-s} = \{(1, 13), (1, 33), (3, 22), (6, 31), (13, 6), (31, 43), (33, 43), (43, 22)\}$$

Entre G_M et G_{s-s} il existe une fonction **bijective** $f : N_M \rightarrow N_{s-s}$ telle que

$$f_b(c)=1, f_b(d)=12, f_b(e)=10, f_b(f)=9, f_b(i)=3, f_b(j)=2, f_b(a)=6, f_b(b)=5 \text{ et } f_b(h)=11$$

Pour la résolution du problème d'isomorphisme de sous-graphe partiel, nous avons adapté l'algorithme proposé par C. Solnon [127]. Nous avons amélioré l'algorithme au niveau de la construction des domaines de valeurs initiaux.

7.5.2. Notations et définitions

Nous introduisons, dans cette partie, quelques définitions et nous précisons quelques notations nécessaires pour la description de notre algorithme

Définition 15. Correspondance

Un couple de nœuds $[x, y] \rightarrow G_M \times G_S$ en relation par f_b est appelé correspondance. Pour une correspondance $c=[x, y]$, le nœud x est l'extrémité initiale ou *source* (dans G_M) de c , y est l'extrémité finale ou *destination* (dans G_S) de c .

En se référant à l'exemple de la Figure 36, G_M est isomorphe au sous-graphe G_{s-s} de G_S . Une bijection solution $f_b : G_M \rightarrow G_{s-s}$ est totalement caractérisée par l'ensemble des correspondances :

$\{[c,1], [d,12], [e,10], [f,9], [i,3],[j,2], [a,6], [b,5], [h,11]\}$

Définition 16. Domaine de Valeurs D_x

Un domaine de valeurs noté D_x spécifie pour chaque nœud $x \in N_M$ l'ensemble de ses correspondances.

Exemple :

Initialement, $D_f = \{[f, 21], [f, 20], [f, 9], [f, 8]\}$

Lorsqu'un domaine de valeurs se réduit en une seule correspondance. Cette correspondance est considérée comme étant un appariement.

Définition 17. Appariement

Un appariement est une correspondance valide notée $[[x, y]] \rightarrow G_M \times G_S$ donnée (correspondance de départ). Comme pour la correspondance, on appelle x la *source* (notée source($[[x,y]]$)) et y la *destination* de l'appariement (notée destination($[[x,y]]$)). La Figure 36 illustre deux appariements de départ : ils sont encadrés en noir.

Exemple : $[[c, 1]]$ est un appariement de départ où c =source($[[c,1]]$) et 1 =destination($[[c,1]]$).

L'ensemble des appariements de départ est noté Appariement_{G_M} . Dans la Figure 36, l'ensemble $\text{Appariement}_{G_M} = \{[[c,1]], [[e,10]]\}$, tel que $e = \text{source}([e,10]) = \text{source}(\text{Appariement}_{G_M}(2))$.

Le nombre d'appariement, noté $|\text{Appariement}_{G_M}|$, dépend du nombre de nœud $n_M = |N_M|$. $|\text{Appariement}_{G_M}| = n_M$ au maximum.

Définition 18. Adjacent $\text{adj}(x)$

La fonction $\text{adj}(x)$ avec $x \in N_M$ ou $x \in N_S$ retourne tous les voisins directs à x .

Exemple : $\text{adj}(f) = \{e, i\}$

Définition 19. Distance entre deux nœuds $d(x, x')$

La distance entre deux nœuds (en négligeant l'orientation des arcs) est la longueur du plus court chemin (forcément simple : sans cycle) de x vers x' , si celui-ci existe sinon elle est égale à infini.

Définition 20. Excentricité de x

L'excentricité, est la distance maximale entre le nœud $x \in N_M$ et les autres nœuds $\in N_M$ sans considérer les cycles. L'excentricité de x dans G_M est notée $excent(x, G_M)$.

G_M est caractérisé par une valeur appelée excentricité maximale, noté em . L'excentricité maximale est le diamètre du graphe c'est-à dire la distance maximale qui sépare les nœuds les plus lointains du graphe G_M .

Définition 21. Degré de x $deg(x)$

Le degré de x , noté $deg(x)$, est le nombre des nœuds adjacents à x . $deg(x) = |adj(x)|$.

G_M est caractérisé par une valeur appelée degré maximal du graphe, noté d_M . Elle représente le degré maximal de tous les nœuds de G_M . Dans l'exemple de la Figure 36, la valeur d_M associée au graphe G_M est celle du nœud d' $ID=i$, cette valeur $d_M=4$.

Définition 22. Restriction sur N_S

Une restriction notée σ sur N_S considérant une condition C est noté $\sigma_C(N_S)$ tel que

$$\sigma_C(N_S) = \{n_S \in N_S, n_S \text{ satisfait la condition } C\}$$

La condition C est une conjonction de conditions liées par des opérateurs logiques comme l'opérateur AND.

Les conditions qui composent C peuvent être des comparaisons basées sur des opérateurs de comparaison simples (i.e. $=, !=, >, <, \geq, \leq$).

Définition 23. Projection d'un nœud selon un attribut

La projection notée π d'un nœud n selon un de ses attributs (*attribut*) est notée $\pi_{attribut}(ID_n) = valeur$ tel que *valeur* est la valeur de l'attribut de n et ID_n est l'identifiant du nœud n .

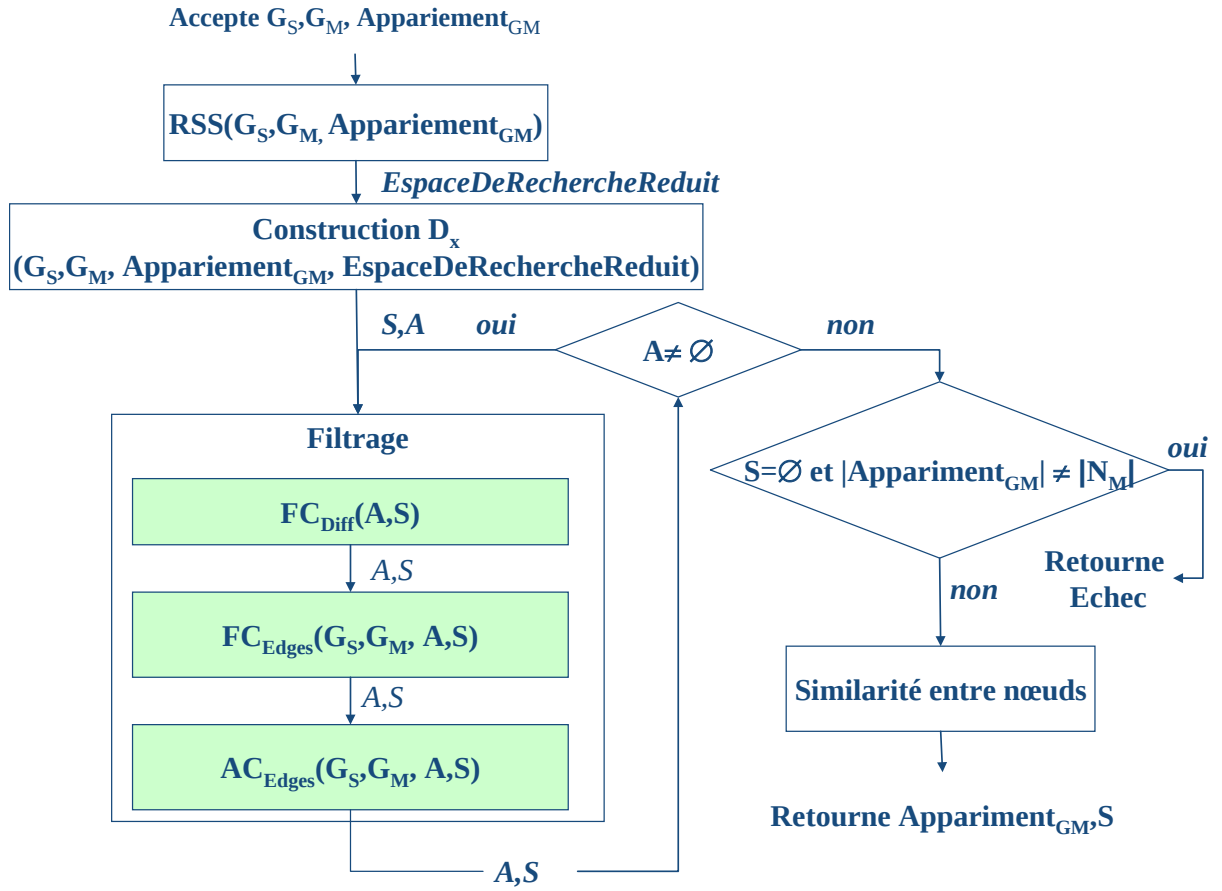
Par exemple, $\pi_{Type}(i) = 'PersonalPhoto'$.

7.5.3. Algorithme d'isomorphisme de sous-graphe partiel

Nous avons proposé un algorithme, appelé *h-Pruning*, qui résout l'isomorphisme de sous-graphe partiel en quatre étapes essentielles ; à savoir : (1) La construction des appariements de départ et les domaines de valeurs, (2) la propagation de contrainte de différence FC_{Diff} , (3) la propagation de la contrainte d'arcs FC_{Edges} et enfin, (4) la propagation de contrainte d'arcs AC_{Edges} .

Les étapes de l'exécution de l'algorithme *h-Pruning* sont illustrées par la Figure 37. L'algorithme *h-Pruning* accepte les graphes G_S , G_M et l'ensemble d'appariements initial

Appariement_{GM} . L'algorithme *RSS* est appelé pour retourner un ensemble baptisé *EspaceDeRechercheReduit*. Cet algorithme sera présenté par l'Algorithme 4 et détaillé dans la section 7.5.3.1. Ensuite, l'Algorithme 5, baptisé construction D_x , sera appelé afin de construire un ensemble S contenant une liste de domaines de valeurs D_x et un ensemble A d'appariements. Ce dernier prend initialement Appariement_{GM} . L'algorithme *construction D_x* est détaillé dans la deuxième partie de la section 7.5.3.1 baptisée *Construction du domaine D_x initial*. Puis, les algorithmes de filtrages FC_{Diff} , FC_{Edges} et AC_{Edges} (présenté respectivement dans les sections 7.5.3.2, 7.5.3.3 et 7.5.3.4) seront exécutés successivement. Le dernier algorithme de filtrage AC_{Edges} remet l'ensemble A à vide si aucun domaine de valeur D_x dans S ne se réduit à un seul élément. Les étapes de filtrages seront ré-exécutées, alors, tant que l'ensemble A est différent de vide. Si la condition $A = \emptyset$ est vérifiée, nous testons si l'ensemble $S = \emptyset$ et le nombre d'appariements Appariement_{GM} est différent du nombre de nœuds motifs N_M . Si ces deux dernières conditions sont vérifiées, l'algorithme *h-Pruning* ne retourne pas de solutions. Sinon, l'algorithme retourne S et Appariement_{GM} comme solutions possibles à l'isomorphisme de sous- graphe partiel après avoir filtrer S en considérant la similarité entre nœuds présentée dans la section 7.5.3.5.


 Figure 37. Etapes d'exécution de l'algorithme *h-Pruning*

7.5.3.1. Construction du domaine D_x initial

Dans l'algorithme proposé par C.Solnon [127], le domaine D_x est réduit à l'ensemble des nœuds cibles dont le degré est supérieur ou égal au degré de x car un nœud x ne peut être apparié à un nœud y que si $|adj(x)| \leq |adj(y)|$.

L'application de cette condition n'est pas suffisante pour réduire le nombre de candidats initiaux car nous sommes obligés de parcourir tous les nœuds du graphe social pour rechercher les candidats initiaux. Comme nous l'avons déjà indiqué, la taille du graphe social risque d'atteindre un nombre gigantesque. Par exemple, si nous considérons seulement le nombre d'utilisateurs, il dépasse sur Facebook les 400 millions. Nous proposons de réduire l'espace de recherche. L'heuristique proposée pour réduire l'espace de recherche prend en considération les appariements initiaux. Dans la suite, nous présentons l'heuristique permettant de réduire l'espace de recherche de N_S à un ensemble appelé *EspaceDeRechercheReduit* tel que $\text{EspaceDeRechercheReduit} \subseteq N_S$.

- **Heuristique de réduction de l'espace de recherche**

L'idée de base de l'heuristique est de centrer la recherche sur un espace qui comporte uniquement les voisins des nœuds ayant un appariement. Soit $y \in N_S$ une destination d'un appariement tel que $[[x, y]] \in N_M \times N_S$. Les voisins sont considérés à une distance du nœud y égale à l'excentricité de x dans G_M .

Ainsi, l'espace de recherche des correspondances de chaque nœud $x' \in N_M$, n'ayant pas un appariement, sera réduit à un sous-ensemble de N_S de taille proportionnelle aux nombres de nœuds du graphe motif $|N_M|$.

Soit $[[x, y]]$ un appariement avec x (la source) $\in N_M$ et y (la destination) $\in N_S$. L'Algorithme 3, baptisé *MarquerNœuds*, permet de sélectionner dans un ensemble appelé *NœudsMarqués* tous les nœuds à une distance de y inférieure ou égale à la valeur de l'excentricité de x . L'algorithme *MarquerNœuds* est basé sur une recherche en profondeur (Depth First Search DFS). Il procède par remplissage récursif de l'ensemble *NœudsMarqués* partant des nœuds voisins directs de y .

Cet algorithme (i.e., *MarquerNœuds*) est ré-exécuté tant que des appariements initiaux sont détectés. L'Algorithme 4, baptisé RSS (Reducing Search Space) permet d'appeler la procédure *MarquerNœuds* tant que des appariements sont détectés. Un vecteur $V[i]$ de *NœudsMarqués* est construit par appariements (ligne 5). L'algorithme RSS retourne un ensemble de nœuds, appelé *EspaceDeRechercheReducit*, correspondant à l'intersection de tous les vecteurs construits $V[i]$ (ligne 7).

Algorithme MarquerNœuds(G_S , NœudsMarqués, *excent*, s , s_{ini} , distance)
 //procédure récursive basé sur DFS

Entrées

Le graphe social $G_S = \langle N_S, R_S \rangle$
excent : l'excentricité d'un nœud de N_M
 s un nœud du G_M
 s_{ini} la destination d'un appariement
 NœudsMarqués : liste contenant les nœuds marqués

Début

1. **Si** distance > *excent* **alors**
2. **retourner** NœudsMarqués
3. **Sinon**
4. ajouter s_{ini} à NœudsMarqués
5. **Pour chaque** $sfils \in adj(s)$ **faire**
6. **Si** $sfils \notin$ NœudsMarqués **alors**
7. **MarquerNœuds**(G_S , NœudsMarqués, *excent*, $sfils$, s_{ini} , distance+1)
8. **Fin Si**
9. **Fin Pour**

Fin algorithme

Algorithme 3. Algorithme MarquerNœuds de sélection des les nœuds potentiellement intéressants

Algorithme RSS(G_S, G_M , Appariement $_{GM}$) //Reducing Search Space

Entrées

Le graphe social $G_S = \langle N_S, R_S \rangle$
 Le graphe motif $G_M = \langle N_M, R_M \rangle$

Variables locales

$V[]$: ensemble de vecteurs contenant les nœuds potentiellement intéressants

Sorties

EspaceDeRechercheReducit : ensemble de nœuds constituant l'espace de recherche réduit

Début

1. **Pour** i de 1 à $|Appariements_{GM}|$ **faire**
2. $s = source(Appariement_{GM}(i))$;
3. $e = excent(s, G_M)$;
4. $d = destination(Appariement_{GM}(i))$;
5. **MarquerNœuds**(G_S , $V[i]$, e , s , d , 0);
6. **Fin Pour**
7. EspaceDeRechercheReducit = intersection des $V[i]$, $i : 1$ à $|Appariements_{GM}|$

Fin algorithme

Algorithme 4. Algorithme RSS de réduction de l'espace de recherche

L'heuristique proposée ajoute un coût à la complexité de l'algorithme d'isomorphisme de sous-graphe original proposé par C. Solnon. En effet, la complexité de l'algorithme RSS est de l'ordre de $O(n_M \times d^{em})$ tel que d représente le degré maximal d'un nœud dans le

graphe G_S , em représente la valeur de l'excentricité maximale du nœud dans G_M et n_M est le nombre de nœuds dans G_M .

- **Construction du domaine D_x initial**

Suite à la construction de l'ensemble V des nœuds concernés par la recherche, la deuxième étape consiste à construire des domaines de valeurs pour chaque nœud $x' \in G_M$ n'ayant pas d'appariement. Le domaine de valeurs d'un nœud x' , appelé $D_{x'}$, est un ensemble contenant les correspondances à x' . La recherche des correspondants initiaux (i.e. éléments d'un domaine de valeurs) prend en considération le nombre de nœuds adjacents à $x' \in G_M$ au nombre de nœuds adjacents à $y \in EspaceDeRechercheReduit$. Soit n_{h-S} le nombre de nœuds considérés après avoir appliqué l'heuristique RSS ; et soit n_M le nombre de nœuds dans G_M . La complexité de l'Algorithme 5 décrivant la construction des domaines de valeurs de chaque nœud de G_M est de l'ordre de $O(n_{h-S} \cdot n_M)$. Toutefois, pour construire les domaines de valeurs, l'application de l'heuristique RSS réduit considérablement l'espace de recherche ce qui représente un gain en termes de temps d'exécution (voir expérimentation section 8.3.2 Chapitre 8).

Pour construire le domaine de valeurs pour un nœud $x \in N_M$, une identification des appariements est effectuée en premier lieu. Soit k le nombre d'appariements initiaux. Après avoir appliqué l'algorithme RSS pour la construction de l'ensemble $EspaceDeRechercheReduit$ des nœuds concernés par la recherche, l'Algorithme 5 est appliqué pour déterminer un ensemble de correspondances initiales pour chaque nœud $x \in N_M$ n'ayant pas d'appariements. Pour construire la correspondance $[x, y]$ dans D_x , on vérifie si $|adj(x)| \leq |adj(y)|$.

L'Algorithme 5 illustre le processus de construction des domaines de valeurs.

Algorithme 5. Algorithme Construction D_x de construction des domaines de valeurs initiaux

Algorithme Construction D_x (G_M , G_S , Appariement_{G_M} ,
EspaceDeRechercheReduit)

Entrées

Le graphe motif $G_M = \langle N_M, R_M \rangle$
 Le graphe social $G_S = \langle N_S, R_S \rangle$
 L'ensemble d'appariements initial Appariement_{G_M}

Sorties

S Liste d'ensembles de correspondance D_x
 A ensemble d'appariements

Initialisation

$D_x \leftarrow \emptyset$, $S \leftarrow \emptyset$, $A \leftarrow \emptyset$

Début

1. **Pour chaque** $x \in N_M$ **et** $\exists y' \in N_S$ **tel que** $[[x, y']] \in \text{Appariement}_{G_M}$ **faire**
2. **Si** $\exists Y \subseteq \text{EspaceDeRechercheReduit}$ **tel que** $|\text{adj}(x)| \leq |\text{adj}(y)|$ **alors**
3. **Pour chaque** $y \in Y$ **faire**
4. $D_x \leftarrow [x, y]$
5. **Fin Pour**
6. **Fin Si**
7. Ajouter D_x à S
8. **Fin Pour**
9. $A \leftarrow \text{Appariement}_{G_M}$

Fin algorithme

En considérant le même exemple de la Figure 36, le domaine de valeurs du nœud f , soit D_f , est sélectionné comme suit :

$$D_f = \{[i, 21], [i, 22], [i, 19]\}$$

Les domaines de valeurs initiaux sont réduits en utilisant des techniques de filtrage qui propagent les contraintes pour supprimer des correspondances non valides des domaines.

Nous avons choisi trois types de filtrage parmi ceux proposés dans [127]. Ces filtrages sont basés sur le principe de la propagation (forward-checking) et sont baptisés FC_{Diff} , FC_{Edges} et AC_{Edges} .

7.5.3.2. Propagation des contraintes de différence FC_{Diff}

Une propagation par vérification des contraintes de différence (notée FC_{Diff}) peut être faite à chaque fois qu'un nœud motif x est apparié à un nœud cible y (i.e. construction d'un appariement de départ $[[x, y]]$) en supprimant les correspondances contenant y des domaines de tous les nœuds non appariés. L'Algorithme 6 illustre le processus de filtrage par propagation FC_{Diff} . La complexité de cet algorithme est de l'ordre de $O(n_M)$ au pire des cas tel que $n_M = |N_M|$. Plus formellement :

$$\forall [[x, y]] \in N_M \times N_S, \forall [x', y'] \in N_M \times N_S, \text{ si } x \neq x' \Rightarrow y \neq y'$$

Algorithme 6. Algorithme Filtrage FC_{Diff}

Algorithme Filtrage $FC_{Diff}(A, S)$

Entrées

Un ensemble A d'appariements
Un ensemble S de correspondances à filtrer

Sorties

S

Début

```

1.  Pour chaque  $[[x, y]] \in A$  faire
2.    Pour chaque  $[x', y'] \in S$  faire
3.      Si  $y = y'$  alors
4.        Supprimer une correspondance  $[x', y']$  de S
5.      Fin Si
6.      Si  $D_x = \emptyset$  alors
7.         $S \leftarrow \emptyset$ 
8.      Fin Si
9.    Fin Pour
10. Fin Pour
Fin Algorithme

```

7.5.3.3. Propagation des contraintes d'arrêt FC_{Edges}

Une propagation par vérification des contraintes d'arrêt (notée FC_{Edges}) peut être effectuée à chaque fois qu'un nœud motif x est apparié à un nœud cible y en enlevant du domaine de chaque nœud adjacent à x tout nœud cible n'étant pas adjacent à y . L'Algorithme 7 illustre le processus de filtrage par propagation FC_{Edges} . Plus formellement :

$$\forall (x, x') \in R_M, [[x, y]] \in N_S \times N_M, [x', y'] \in N_S \times N_M, (y, y') \in R_S.$$

Théoriquement, la complexité de l'algorithme est de $O(d_M \cdot n_S)$ tel que d_M est le degré maximal du graphe G_M et n_S est le nombre de nœuds dans G_S . Cependant, le nombre de correspondances dans chaque domaine de valeurs n'atteint jamais n_S puisque : (i) nous avons appliqué une heuristique RSS pour réduire l'espace de recherche de n_S à n_{h-S} ; (ii) nous avons effectué des restrictions en considérant une comparaison entre attributs spécifiés dans G_M ; (iii) nous avons aussi appliqué FC_{Diff} pour notre cadre d'application, le nombre d'appariements de départ est au minimum 1. Donc $|D_x| \ll n_S$. La complexité de cette étape de l'algorithme d'isomorphisme de sous-graphe partiel est, alors, de $O(d_M \cdot n_{h-S})$.

Algorithme 7. Algorithme Filtrage FC_{Edges}

Algorithme Filtrage $FC_{Edges}(G_M, G_S, A, S)$

Entrées

Le graphe motif $G_M = \langle N_M, R_M \rangle$
 Le graphe social $G_S = \langle N_S, R_S \rangle$
 Un ensemble A d'appariements de G_M
 Un ensemble S de correspondances à filtrer

Sorties

S

Début

```

1.  Pour chaque  $[[x,y]] \in A$  faire
2.      Pour chaque  $x' = \text{adj}(x)$  faire
3.          Si  $\exists y'$  tel que  $[x',y'] \in S$  alors
4.               $\exists (y, y') \in R_S$ 
5.          Sinon
6.              Supprimer la correspondance  $[x',y']$  de  $S$ 
7.          Fin Si
8.          Si  $D_{x'} = \emptyset$  alors
9.               $S \leftarrow \emptyset$ 
10.         Fin Si
11.     Fin Pour
12. Fin Pour
Fin Algorithme
    
```

7.5.3.4. Propagation des contraintes d'arrêt AC_{Edges}

Une propagation par vérification des contraintes d'arrêt (notée AC_{Edges}) permet de vérifier pour chaque arrêt motif (x, x') et chaque correspondance $[x,y] \in D_x$, qu'il existe au moins un $[x',y'] \in D_{x'}$ tel que y' est adjacent à y (i.e. $(y,y') \in R_S$). Le nœud cible y' est dit support de la correspondance $[x,y]$ pour l'arrêt (x,x') . Si une correspondance $[x,y]$ n'a pas de support pour un arc, alors la correspondance $[x,y]$ doit être enlevée de D_x . Une telle propagation des contraintes d'arrêt est itérée jusqu'à ce que plus aucune valeur ne puisse être enlevée, assurant ainsi la cohérence d'arc des contraintes d'arrêt. Plus formellement :

$$\forall (x, x') \in R_M, \forall [x,y] \in D_x, \text{ si } \exists [x',y'] \in D_{x'} \text{ alors } (y,y') \in R_S.$$

L'Algorithme 8 illustre le processus de filtrage par propagation AC_{Edges} . De la ligne 13 à 21 : à chaque fois où un domaine de valeurs se réduit en une seule correspondance. La correspondance se transforme alors en appariement. L'ensemble d'appariements A prend, alors, les nouveaux appariements.

La complexité théorique de AC_{Edges} est de $O(r_M \cdot n_S^2)$ avec $r_M = |R_M|$ et $n_S = |N_S|$. Cependant, le nombre de correspondance dans chaque domaine de valeurs n'atteint jamais n_S puisque :

(i) nous avons appliqué une heuristique RSS pour réduire l'espace de recherche de n_S à n_{h-S} tel que n_{h-S} est proportionnelle au nombre de nœuds dans le graphe motif G_M ; (ii) Nous avons aussi appliqué FC_{Diff} et (iii) ainsi que FC_{Edges} . Pour notre cadre d'application, le nombre d'appariement de départ est au minimum 1. Donc $|D_x| \lll n_S$. La complexité de cette étape de l'algorithme est de $O(r_M \cdot n_{h-S}^2)$

Algorithme Filtrage $AC_{Edges}(G_M, G_S, S, A, Appariment_{GM})$

Entrées

Le graphe motif $G_M = \langle N_M, R_M \rangle$
 Le graphe social $G_S = \langle N_S, R_S \rangle$
 Un ensemble S de correspondances à filtrer
 Un ensemble d'appariements A
 L'ensemble d'appariements initial $Appariment_{GM}$

Sorties

$S, A, Appariment_{GM}$

Initialisation

$A' \leftarrow \emptyset$

Début

```

1.  Pour chaque  $(x, x') \in R_M$  faire
2.      Pour chaque  $[x, y] \in D_x$  faire
3.          Si  $\exists [x', y'] \in D_{x'}$  alors
4.               $\exists (y, y') \in R_S$ 
5.          Sinon
6.              Supprimer la correspondance  $[x, y]$  de  $S$ 
7.          Fin Si
8.          Si  $D_{x'} = \emptyset$  alors
9.               $S \leftarrow \emptyset$ 
10.         Fin Si
11.     Fin Pour
12. Fin Pour
13. Pour chaque  $D_x \in S$  faire
14.     Si  $|D_x| = 1$  alors
15.         Ajouter  $D_x$  à  $A'$ 
16.         Supprimer  $D_x$  de  $S$ 
17.     Fin Si
18. Fin Pour
19. Ajouter  $A$  à  $Appariment_{GM}$ 
20. Si  $A' \neq \emptyset$  alors
21.      $A \leftarrow A'$ 
22. Sinon
23.      $A \leftarrow \emptyset$ 
24. Fin Si
Fin Algorithme
    
```

Algorithme 8. Algorithme Filtrage AC_{Edges}

7.5.3.5. Considération de la similarité entre nœuds pour filtrer les domaines de valeurs

La dernière étape consiste à appliquer des restrictions (i.e., définies dans la section 7.5.3.5) tout en prenant en considération la comparaison entre les attributs des deux nœuds de part et d'autres d'une correspondance. Les fonctions de comparaison entre les attributs sont spécifiées par l'utilisateur.

Par exemple, pour le nœud f l'utilisateur a spécifié pour $a_1=(\text{Tag}, \text{'home'})$, $f_{c_{\text{thesaurus}}}(a_1, a') \geq \text{seuil}$. Ceci veut dire que si $\exists y$ tel que $[f, y] \in D_f$ alors y doit avoir un attribut Tag tel que $\langle \pi_{\text{Tag}}(f), \pi_{\text{Tag}}(y) \rangle \geq \text{seuil}$.

La complexité de cette étape est de l'ordre $O(|at| \times n_M \times \psi)$ tel que ψ est le coût d'une opération de comparaison, $|at|$ est le nombre maximal d'attribut par nœud spécifié dans un graphe motif G_M et n_M est le nombre de nœuds dans G_M .

7.5.3.6. Analyse de la complexité de notre algorithme *h-Pruning*

Considérons un nœud générique $n_M \in N_M$, d_M est le degré maximal du graphe motif G_M , n_{h-S} le nombre de nœuds considérés par l'algorithme *h-Pruning*, ψ est le coût d'une opération de comparaison, $|at|$ est le nombre maximal d'attributs par nœud spécifié dans un graphe motif G_M , r_M est le nombre de relations dans le graphe motif G_M ($|R_M|$).

Nous analysons la complexité de l'algorithme *h-Pruning* permettant d'instancier le graphe motif dans le graphe social dans les deux cas différents : le meilleur cas c'est-à-dire lorsque le nombre d'appariements initial est non nul et le pire cas c'est-à-dire lorsque le nombre d'appariements initial est nul.

- **Le meilleur cas**

Dans le meilleur cas, la complexité de l'instanciation du graphe motif sera une complexité polynomiale. Elle est égale à $O(r_M \cdot n_{h-S}^2 + d_M \cdot n_{h-S} + (|at| \times \psi + n_{h-S} + 1 + d^{em}) \cdot n_M)$. L'algorithme *h-Pruning* s'exécute en six étapes à savoir : (i) construction d'un espace de recherche réduit ; (ii) construction des domaines de valeurs D_x de chaque nœud de G_M ; (iii) filtrage des domaines de valeurs en propageant les contraintes de différence FC_{Diff} ; (iv) filtrage des domaines de valeur en propageant les contraintes d'arrête FC_{Edges} ; (v) filtrage des domaines de valeur en propageant les contraintes d'arrête AC_{Edges} ; et enfin (vi) filtrage des domaines de valeurs en comparant les attributs de chaque nœud. La complexité de chaque étape est égale respectivement : (i) $O(n_M \times d^{em})$; (ii) $O(n_{h-S} \cdot n_M)$; (iii) $O(n_M)$; (iv)

$O(d_M \cdot n_{h-S})$; (v) $O(r_M \cdot n_{h-S}^2)$; et (vi) $O(|at| \times n_M \times \psi)$. La première étape sera très coûteuse si le graphe est très dense car la complexité dépend du degré maximal d'un nœud dans le graphe G_S .

- **Le pire cas**

Dans le pire cas, nous ne pouvons pas construire un espace de recherche réduit en raison d'absence d'appariements initiaux. Dans ce cas l'algorithme *h-Pruning* se ramène à l'algorithme initial de C. Solnon. Sa complexité est exponentielle. En pratique, le pire cas n'est pas envisageable car nous appliquons cet algorithme en recherche d'information.

7.5.4. Exemple complet

Nous décrivons, dans cette section, un exemple complet pour expliquer le fonctionnement de l'algorithme d'isomorphisme de sous-graphe partiel. L'exemple illustré dans la Figure 36 est repris pour expliquer les étapes d'exécution.

On rappelle les caractéristiques des deux graphes de départ G_M et G_S :

$G_M = \langle N_M, R_M \rangle$

Tel que

$N_M = \{ab, c, d, e, f, i, h\}$

et

$R_M = \{(i, f), (e, f), (a, i), (b, i), (c, a), (c, d), (d, h), (h, b)\}$

et $G_S = \langle N_S, R_S \rangle$

$N_S = \{1, 2, 3, 4, 5, 6, 11, 12, 13, 14, 15, 16, 19, 21, 22, 23, 24, 31, 32, 33, 34, 35, 41, 42, 43, 44, 45\}$

et

$R_S = \{(1, 13), (1, 33), (1, 14), (1, 15), (1, 16), (1, 12), (2, 19), (2, 32), (2, 11), (3, 22), (3, 21), (4, 34), (4, 35), (6, 31), (11, 5), (12, 2), (13, 6), (14, 4), (15, 3), (16, 3), (31, 44), (31, 43), (32, 41), (32, 42), (32, 43), (33, 41), (33, 42), (33, 43), (34, 45), (35, 41), (41, 21), (41, 24), (42, 24), (43, 22), (45, 23)\}$

À l'étape 0 : Construction d'un espace de recherche réduit

L'étape 0 consiste à réduire l'espace de recherche et construire un ensemble appelé *EspaceDeRechercheReduit*, ceci, à partir des voisins de la destination des appariements qui

sont considérés à une distance de ce nœud égale à l'excentricité de sa source (source du même appariement) dans G_M .

Tout d'abord, l'algorithme *MarquerNœuds* procède par un remplissage de l'ensemble *NœudsMarqués* partant des voisins directs de la destination d'un appariement.

Dans l'exemple, nous considérons deux appariements :

- Construction des appariements

$[[e,10]]$ tel que $\text{exent}(e, G_M)=5$

$[[c,1]]$ tel que $\text{exent}(c, G_M)=4$

Pour le premier (i.e. e),

$NœudsMarqués_e = \{3,16,22,21,15,41,1,43,32,33,24,35,13,14,12,31,42,4,6,2,44,34,19,11\}$

Pour le deuxième (i.e. c),

$NœudsMarqués_c = \{1,13,33,14,15,16,12,6,43,41,42,4,3,2,31,32,22,21,24,35,34,19,11,44,45,5\}$

L'Algorithme 4 appelé RSS permet d'appeler la procédure *MarquerNœuds* au nombre de fois que des appariements sont détectés. Un vecteur $V[i]$ de *NœudsMarqués* est construit par appariements. L'algorithme RSS retourne un ensemble de nœuds, appelé *EspaceDeRechercheReducit*, correspondant à l'intersection de tous les vecteurs construits $V[i]$.

$V[0] = \{3,16,22,21,15,41,1,43,32,33,24,35,13,14,12,31,42,4,6,2,44,34,19,11\}$

$V[1] = \{1,13,33,14,15,16,12,6,43,41,42,4,3,2,31,32,22,21,24,35,34,19,11,44,45,5\}$

EspaceDeRechercheReducit prend l'intersection des deux vecteurs $V[0]$ et $V[1]$.

$\Rightarrow \text{EspaceDeRechercheReducit} = \{1, 2, 3, 4, 6, 11, 12, 13, 14, 15, 16, 19, 21, 22, 24, 31, 32, 33, 34, 35, 41, 42, 43, 44\}$.

La recherche sera effectuée seulement sur cet ensemble de nœuds et non pas sur tous les nœuds du graphe G_S .

À l'étape 1 : Construction de domaines de valeurs

L'étape 1 consiste à construire les domaines de valeurs en comparant le degré des nœuds. Cela réduit l'ensemble des nœuds cibles dont le degré est supérieur ou égal au degré de $x \in N_M$ car un nœud x ne peut être apparié à un nœud y que si $|adj(x)| \leq |adj(y)|$.

- Initialisation des domaines de valeurs

$$\deg(d)=2$$

$$\Rightarrow \mathbf{D}_d = \{[d, 1], [d, 2], [d, 3], [d, 4], [d, 6], [d, 11], [d, 12], [d, 13], [d, 14], [d, 15], [d, 16], [d, 21], [d, 22], [d, 24], [d, 31], [d, 32], [d, 33], [d, 34], [d, 35], [d, 41], [d, 42], [d, 43]\}$$

$$\deg(i)=3$$

$$\Rightarrow \mathbf{D}_i = \{[i, 1], [i, 2], [i, 3], [i, 4], [i, 31], [i, 32], [i, 33], [i, 41], [i, 42], [i, 43]\}$$

$$\deg(f)=2$$

$$\Rightarrow \mathbf{D}_f = \{[f, 1], [f, 2], [f, 3], [f, 4], [f, 6], [f, 11], [f, 12], [f, 13], [f, 14], [f, 15], [f, 16], [f, 21], [f, 22], [f, 24], [f, 31], [f, 32], [f, 33], [f, 34], [f, 35], [f, 41], [f, 42], [f, 43]\}$$

$$\deg(b)=2$$

$$\Rightarrow \mathbf{D}_b = \{[b, 1], [b, 2], [b, 3], [b, 4], [b, 6], [b, 11], [b, 12], [b, 13], [b, 14], [b, 15], [b, 16], [b, 21], [b, 22], [b, 24], [b, 31], [b, 32], [b, 33], [b, 34], [b, 35], [b, 41], [b, 42], [b, 43]\}$$

$$\deg(h)=2$$

$$\Rightarrow \mathbf{D}_h = \{[h, 1], [h, 2], [h, 3], [h, 4], [h, 6], [h, 11], [h, 12], [h, 13], [h, 14], [h, 15], [h, 16], [h, 21], [h, 22], [h, 24], [h, 31], [h, 32], [h, 33], [h, 34], [h, 35], [h, 41], [h, 42], [h, 43]\}$$

$$\deg(a)=2$$

$$\Rightarrow \mathbf{D}_a = \{[a, 1], [a, 2], [a, 3], [a, 4], [a, 6], [a, 11], [a, 12], [a, 13], [a, 14], [a, 15], [a, 16], [a, 21], [a, 22], [a, 24], [a, 31], [a, 32], [a, 33], [a, 34], [a, 35], [a, 41], [a, 42], [a, 43]\}$$

Ainsi, les ensembles S de correspondances et A d'appariements vont être composés comme suit:

$$\begin{aligned} \mathbf{S} = \{ & [d, 1], [d, 2], [d, 3], [d, 4], [d, 6], [d, 11], [d, 12], [d, 13], [d, 14], [d, 15], [d, 16], [d, 21], \\ & [d, 22], [d, 24], [d, 31], [d, 32], [d, 33], [d, 34], [d, 35], [d, 41], [d, 42], [d, 43], [i, 1], [i, 2], \\ & [i, 3], [i, 4], [i, 31], [i, 32], [i, 33], [i, 41], [i, 42], [i, 43], [f, 1], [f, 2], [f, 3], [f, 4], [f, 6], [f, 11], \\ & [f, 12], [f, 13], [f, 14], [f, 15], [f, 16], [f, 21], [f, 22], [f, 24], [f, 31], [f, 32], [f, 33], [f, 34], \\ & [f, 35], [f, 41], [f, 42], [f, 43], [b, 1], [b, 2], [b, 3], [b, 4], [b, 6], [b, 11], [b, 12], [b, 13], [b, 14], \\ & [b, 15], [b, 16], [b, 21], [b, 22], [b, 24], [b, 31], [b, 32], [b, 33], [b, 34], [b, 35], [b, 41], [b, 42], \\ & [b, 43], [h, 1], [h, 2], [h, 3], [h, 4], [h, 6], [h, 11], [h, 12], [h, 13], [h, 14], [h, 15], [h, 16], \\ & [h, 21], [h, 22], [h, 24], [h, 31], [h, 32], [h, 33], [h, 34], [h, 35], [h, 41], [h, 42], [h, 43], [a, 1], \\ & [a, 2], [a, 3], [a, 4], [a, 6], [a, 11], [a, 12], [a, 13], [a, 14], [a, 15], [a, 16], [a, 21], [a, 22], [a, 24], \\ & [a, 31], [a, 32], [a, 33], [a, 34], [a, 35], [a, 41], [a, 42], [a, 43] \} \end{aligned}$$

$$\mathbf{A} = \{[[e, 10]], [[c, 1]]\}$$

À l'étape 2 : Propagation de contrainte de différence FC(Diff)

On applique le filtrage FC(Diff), présenté dans l'Algorithme 6, pour filtrer les domaines en prenant en considération les appariements dans A . A cette étape $A=\{[[e,3]], [[c,1]]\}$. Ainsi, les domaines de D_a , D_h , D_b, D_d , D_i , et D_f , ne doivent pas contenir les nœuds 10 et 1. Ces nœuds (10 et 1) doivent être supprimés des domaines D_a , D_h , D_b, D_d , D_i et D_f

$D_d=\{ [d,2], [d,4], [d, 6], [d,11], [d,12], [d,13], [d,14], [d,15], [d,16], [d,21], [d,22], [d,24], [d,31], [d,32], [d,33], [d,34], [d,35], [d,41], [d,42], [d,43]\}$

$D_i=\{ [i,2], [i,4], [i,31], [i,32], [i,33], [i,41], [i,42], [i,43]\}$

$D_f=\{[f,2], [f,4], [f, 6], [f,11], [f,12], [f,13], [f,14], [f,15], [f,16], [f,21], [f,22], [f,24], [f,31], [f,32], [f,33], [f,34], [f,35], [f,41], [f,42], [f,43]\}$

$D_b=\{[b,2], [b,4], [b, 6], [b,11], [b,12], [b,13], [b,14], [b,15], [b,16], [b,21], [b,22], [b,24], [b,31], [b,32], [b,33], [b,34], [b,35], [b,41], [b,42], [b,43]\}$

$D_h=\{[h,2], [h,4], [h, 6], [h,11], [h,12], [h,13], [h,14], [h,15], [h,16], [h,21], [h,22], [h,24], [h,31], [h,32], [h,33], [h,34], [h,35], [h,41], [h,42], [h,43]\}$

$D_a=\{[a,2], [a,4], [a, 6], [a,11], [a,12], [a,13], [a,14], [a,15], [a,16], [a,21], [a,22], [a,24], [a,31], [a,32], [a,33], [a,34], [a,35], [a,41], [a,42], [a,43]\}$

$S=\{[d,2], [d,4], [d, 6], [d,11], [d,12], [d,13], [d,14], [d,15], [d,16], [d,21], [d,22], [d,24], [d,31], [d,32], [d,33], [d,34], [d,35], [d,41], [d,42], [d,43], [i,2], [i,4], [i,31], [i,32], [i,33], [i,41], [i,42], [i,43], [f,2], [f,4], [f, 6], [f,11], [f,12], [f,13], [f,14], [f,15], [f,16], [f,21], [f,22], [f,24], [f,31], [f,32], [f,33], [f,34], [f,35], [f,41], [f,42], [f,43], [b,2], [b,4], [b, 6], [b,11], [b,12], [b,13], [b,14], [b,15], [b,16], [b,21], [b,22], [b,24], [b,31], [b,32], [b,33], [b,34], [b,35], [b,41], [b,42], [b,43], [h,2], [h,4], [h, 6], [h,11], [h,12], [h,13], [h,14], [h,15], [h,16], [h,21], [h,22], [h,24], [h,31], [h,32], [h,33], [h,34], [h,35], [h,41], [h,42], [h,43], [a,2], [a,4], [a, 6], [a,11], [a,12], [a,13], [a,14], [a,15], [a,16], [a,21], [a,22], [a,24], [a,31], [a,32], [a,33], [a,34], [a,35], [a,41], [a,42], [a,43]\}$

$A=\{[[e,10]], [[c,1]]\}$

À l'étape 3 : Propagation de contraintes d'arcs FC_{Edges}

On applique le filtrage FC_{Edges} , présenté dans l'Algorithme 7 pour filtrer les domaines en considérant les arcs avec les nœuds qui construisent les appariements. Ce filtrage permet de supprimer 2, 4, 6, 11, 21, 22, 24, 31, 32, 34, 35, 41, 42 et 43 de D_d car ils ne sont pas adjacents au correspondant de c autrement dit, 1, supprime 2, 4, 6, 11, 21, 22, 24, 31, 32, 34, 35, 41, 42 et 43 de D_a car ils ne sont pas adjacents au correspondant de c c'est-à-dire 1.

Avec la contrainte FC_{Edges} on supprime, aussi, 2, 4, 6, 11, 12, 13, 14, 24, 31, 32, 33, 34, 35, 41, 42 et 43 de D_f car ils ne sont pas adjacents à 3.

$$D_d = \{ [d,12], [d,13], [d,14], [d,15], [d,16], [d,33] \}$$

$$D_i = \{ [i,2], [i,4], [i,31], [i,32], [i,33], [i,41], [i,42], [i,43] \}$$

$$D_f = \{ [f,15], [f,16], [f,21], [f,22] \}$$

$$D_b = \{ [b,2], [b,4], [b,6], [b,11], [b,12], [b,13], [b,14], [b,15], [b,16], [b,21], [b,22], [b,24], [b,31], [b,32], [b,33], [b,34], [b,35], [b,41], [b,42], [b,43] \}$$

$$D_h = \{ [h,2], [h,4], [h,6], [h,11], [h,12], [h,13], [h,14], [h,15], [h,16], [h,21], [h,22], [h,24], [h,31], [h,32], [h,33], [h,34], [h,35], [h,41], [h,42], [h,43] \}$$

$$D_a = \{ [a,12], [a,13], [a,14], [a,15], [a,16], [a,33] \}$$

$$S = \{ [d,12], [d,13], [d,14], [d,15], [d,16], [d,33], [i,2], [i,4], [i,31], [i,32], [i,33], [i,41], [i,42], [i,43], [f,15], [f,16], [f,21], [f,22], [b,2], [b,4], [b,6], [b,11], [b,12], [b,13], [b,14], [b,15], [b,16], [b,21], [b,22], [b,24], [b,31], [b,32], [b,33], [b,34], [b,35], [b,41], [b,42], [b,43], [h,2], [h,4], [h,6], [h,11], [h,12], [h,13], [h,14], [h,15], [h,16], [h,21], [h,22], [h,24], [h,31], [h,32], [h,33], [h,34], [h,35], [h,41], [h,42], [h,43], [a,12], [a,13], [a,14], [a,15], [a,16], [a,33] \}$$

$$A = \{ [[e,10]], [[c,1]] \}$$

À l'étape 4 : Propagation de contrainte d'arcs AC(Edges),

On applique le filtrage AC(Edges), illustré par l'Algorithme 7. Ce filtrage concerne les arcs entre les nœuds qui possèdent des domaines de valeurs. On supprime, alors, en plus, 2, 4, 31, 32 et 33 de D_i car les correspondances $[i,2]$, $[i,4]$, $[i,31]$, $[i,32]$ et $[i,33]$ n'ont pas de support pour l'arc (i,f) (aucun des nœuds adjacents à 2, 4, 31, 32, ou 33 n'appartient à D_i).

On supprime, en plus, 12,13, 14, 15 et 16 de D_a car les correspondances $[a, 12]$, $[a,13]$, $[a,14]$, $[a,15]$ et $[a,16]$ n'ont pas de support pour l'arc (a,i) . De même on supprime 2, 4, 6, 11, 12, 13, 14, 15, 16, 21, 22, 24, 34, 41, 42, et 43 de D_b et on supprime 11 et 5 de D_i . On supprime 15, 16 et 33 de D_d . On supprime également 11, 12, 13, 14, 15, 16, 21, 22, 24, 31, 32, 33, 34, 35, 41, 42, et 43 de D_h .

$$D_d = \{[d,12], [d,13], [d,14]\}$$

$$D_i = \{[i,41], [i,42], [i,43]\}$$

$$D_f = \{[f,15], [f,16], [f,21], [f,22]\}$$

$$D_b = \{[b,31], [b,32], [b,33], [b,35]\}$$

$$D_h = \{[h,2], [h,4], [h, 6]\}$$

$$D_a = \{[a,33]\}$$

$$S = \{[d,12], [d,13], [d,14], [i,41], [i,42], [i,43], [f,15], [f,16], [f,21], [f,22], [b,31], [b,32], [b,33], [b,35], [h,2], [h,4], [h, 6]\}$$

Création d'un nouveau appariement ayant comme source a:

$$A = \{[[e,10]], [[c,1]], [[a,33]]\}$$

On ré-applique $FC(Diff)$, pour supprimer 33 du domaine de D_b

$$D_d = \{[d,12], [d,13], [d,14]\}$$

$$D_i = \{[i,41], [i,42], [i,43]\}$$

$$D_f = \{[f,15], [f,16], [f,21], [f,22]\}$$

$$D_b = \{[b,31], [b,32], [b,35]\}$$

$$D_h = \{[h,2], [h,4], [h, 6]\}$$

$$S = \{[d,12], [d,13], [d,14], [i,41], [i,42], [i,43], [f,15], [f,16], [f,21], [f,22], [b,31], [b,32], [b,35], [h,2], [h,4], [h, 6]\}$$

$$A = \{[[e,10]], [[c,1]], [[a,33]]\}$$

À l'étape 5 : Comparaison des attributs

L'étape 5 consiste à filtrer les domaines de valeurs en considérant la comparaison des attributs pour les domaines de valeurs. On supprime, alors, 35 de D_b et ainsi que, 15 et 16 de D_f

$$\sigma_{Type='Role' \ Spec='actor'}(D_b) = \{[b,31], [b,32]\}$$

$$\sigma_{\text{Type}='socialRelationship'}(D_d) = \{[d,12], [d,13], [d,14]\}$$

$$\sigma_{\text{Type}='PersonalPhoto'}(D_i) = \{[i,41], [i,42], [i,43]\}$$

$$\sigma_{f_{\text{thesaurus}}}(D_f) = \langle \text{Tag}, 'home', \text{Tag}, ('home' \text{ Simthesaurus } y) \rangle \geq \text{seuil} = \{[f,21], [f,22]\}$$

$$\sigma_{\text{Type}='Person'}(D_h) = \{[h,2], [h,4], [h,6]\}$$

$$A = \{[e,10], [c,1], [a,33], \}$$

L'ensemble S et A finaux sont alors :

$$S = \{[d,12], [d,13], [d,14], [i,41], [i,42], [i,43], [f,21], [f,22], [b,31], [b,32], [h,2], [h,4], [h,6]\}$$

$$A = \{[e,10], [c,1], [a,33]\}$$

7.6. Résumé et discussion

Dans ce chapitre, nous avons présenté un algorithme permettant d'exploiter le graphe social en appariant une requête aux nœuds de ce graphe social. L'appariement entre requête et graphe social est ramené à un problème d'isomorphisme de sous-graphe partiel. Nous avons proposé un algorithme qui optimise un algorithme existant. L'optimisation proposé prend en compte les nœuds du graphe requête : appelé graphe motif dans ce travail de thèse.

La solution que nous avons proposée est de complexité théorique raisonnable. En effet, la recherche des solutions selon notre algorithme ne dépend pas du nombre des nœuds dans le graphe social mais d'un espace réduit de recherche qui dépend des nœuds appariés dans le graphe motif.

Nous planifions dans l'avenir d'étendre l'algorithme proposé en présentant à l'utilisateur des solutions qui concernent une partie du graphe motif (i.e. une projection selon un sous-graphe motif).

Chapitre 8. Implémentation et Evaluation

8.1.	INTRODUCTION.....	160
8.2.	PRESENTATION DU PROTOTYPE	161
8.2.1.	<i>Application mobile : PASMi</i>	163
8.2.2.	<i>Composante web pour l'organisation et la publication de photos : Phoox</i>	166
8.2.2.1.	Fonctionnalités.....	166
8.2.2.2.	Les bases de données	172
8.2.2.3.	Représentation et accès au profil social	175
8.3.	EXPERIMENTATION ET EVALUATION DES PERFORMANCES.....	176
8.3.1.	<i>Construction de collection de test</i>	177
8.3.2.	<i>Évaluation de l'approche de l'exploitation du graphe social</i>	178
8.3.3.	<i>Evaluation de l'approche de suggestion des dimensions sociales</i>	181
8.3.4.	<i>Evaluation de l'approche de suggestion d'événements</i>	183
8.4.	DISCUSSION.....	188
8.5.	RESUME.....	189

8.1. Introduction

Nous avons proposé des modèles et des plates-formes qui permettent l'annotation sensible au contexte social et la recherche des documents multimédias socio-personnels. Ce chapitre a pour objectif de les évaluer et de tester leur performance, grâce à un prototype d'annotation et de recherche sensible au contexte social développé dans le cadre de ces travaux de thèse. Ce prototype concrétisé comporte un système mobile appelé PASMi (Photo Annotation and information Sharing Middleware) ainsi qu'un système Web pour la gestion de photos, baptisé *Phoox*. Ainsi, nous souhaitons illustrer à l'aide de *Phoox* et *PASMi* les points suivants : (1) les avantages de l'exploitation de métadonnées contextuelles et des réseaux sociaux pour l'annotation, (2) la recherche, (3) le partage de documents multimédias socio-personnels, sans négliger, (4) l'impact des efforts de configuration ou d'annotation manuelle effectués par des utilisateurs finaux.

Dans ce chapitre, nous présentons plus particulièrement les trois expérimentations réalisées en utilisant notre prototype. La première série d'expérimentations décrit les tests de mesure de l'efficacité de notre approche de recherche appliqués sur des documents multimédias socio-personnels. Une deuxième série d'expérimentations se base sur le test de l'approche de suggestion des "tags événements", basé sur l'extraction des règles d'association. Enfin, la troisième série d'expérimentations porte sur l'approche d'enrichissement des "tags identité" des personnes par des dimensions sociales.

Notre approche de recherche exploitant le graphe social et basée sur l'isomorphisme de sous-graphe partiel a été testé en utilisant un graphe social construit à partir de Flickr. Ce graphe social comporte des photos géo-localisées et partiellement annotées ainsi que des profils sociaux. Nous avons enrichi ces photos avec notre plateforme d'enrichissement automatique. La nouvelle approche, appelée *h-Pruning*, a été comparée avec celle de C.Solnon [127] et une évaluation en terme de temps d'exécution a été effectuée.

Tout au long de la deuxième série d'expérimentations (test de l'approche de suggestion des tags événement), nous avons utilisé les mêmes données de la première série d'expérimentations. L'objectif de cette deuxième expérimentation est de montrer : (i) l'existence d'une cooccurrence entre le cadre spatio-temporel et l'événement ; (ii) l'importance de la notion de la proximité sociale, l'idée de base de l'algorithme *RSP* (proposée dans la section 6.3 Chapitre 6), dans les sites de médias sociaux.

Dans la dernière série d'expérimentations, le temps d'exécution de l'algorithme est calculé grâce à l'approche d'enrichissement des "tags identité" des personnes, baptisée *SQO* (détaillé dans la section 6.4 Chapitre 6). Ainsi, l'algorithme *SQO* est comparé à l'algorithme *BFS* (i.e. la recherche en largeur), ceci, en termes de temps de réponse sur une collection de test construit à partir de vrais profils d'utilisateurs spécifiés en "*foaf*" et "*SocialSphere*".

Les résultats du prototype largement détaillés dans ce chapitre ont été partiellement publiés dans [128], [129], et [111].

Ce chapitre est organisé comme suit : la section 8.2 a pour objectif de présenter le prototype développé avec les deux composantes mobile et web. Les résultats des expérimentations que nous avons menés pour évaluer les performances de nos propositions sont présentés en section 8.3. Nous terminons ce chapitre avec une discussion et un résumé présentés respectivement dans la section 8.4 et la section 8.5.

8.2. Présentation du prototype

Le prototype que nous avons développé comporte deux composants : Le composant mobile (*PASMi*) et le composant web (*Phoox*). Ces deux composants ont pour objectif de mettre en valeur les différents enjeux de nos travaux listés dans ce qui suit :

- **Automatiser au maximum le processus d'annotation** : Le système doit proposer un maximum d'annotations qui doivent être sémantiquement riches pour l'utilisateur.
- **Se servir du contexte physique de prise de vue comme point de départ pour construire des annotations riches** : Le système doit profiter de l'environnement ambiant de l'utilisateur pour apporter des informations contextuelles utiles servant à enrichir les annotations (e.g. temps, localisation, dispositifs autour). Le nombre important d'utilisateurs des réseaux sociaux et de services de médias sociaux sont des utilisateurs mobiles ce qui donne à l'approche un aspect réaliste.
- **Apporter des enrichissements d'annotations de plus haut niveau sémantique grâce à l'utilisation du profil social** : Le système doit permettre à l'utilisateur d'annoter et d'exploiter ses collections de photos selon des axes sémantiques et personnalisés. Des axes personnalisés comme 'ma maison', 'ma cousine', 'mon marie' doivent être possibles aussi bien en annotation qu'en recherche.

- **Exploitation efficace du graphe social :** Toutes informations dans le graphe social comme les documents multimédia, le réseau social et les annotations doivent être exploitables de manière efficace en termes de temps d'exécution.
- **Sensibilité sociale des suggestions d'annotations:** L'assistance apportée par le système à l'annotation doit être sensible au contexte social et s'adapter aux personnes lors de l'annotation ou de consultation des photos.
- **Exploiter la cooccurrence entre le contexte spatio-temporel-social et les événements :** L'assistance à l'annotation des tags doit profiter de la cooccurrence entre le contexte spatio-temporel et les événements pour construire des règles d'associations. Les règles construites pour chaque utilisateur construisent des potentiels de recommandation en cas d'annotation d'une nouvelle photo. La proximité sociale, quant à elle, peut contribuer à l'enrichissement des règles d'associations pour un utilisateur dans le cas où le contexte de l'annotation en cours ne correspond pas à aucune (ou peu) des règles existants pour chaque utilisateur.
- **Donner le choix à l'utilisateur d'annoter manuellement ses documents multimédias à différents moments du cycle de leurs gestion :** Kustanowitz et Shneiderman [4] discutent les différents moments du cycle de gestion de photos où les utilisateurs préfèrent annoter qui sont : au moment de la prise de vue, au moment du chargement, au moment du partage ou au moment de création de leurs collections.

Compte tenue de ces différents besoins, nous avons développé un prototype comportant deux composants : un mobile et un autre web. Le composant mobile de notre prototype est développé pour permettre à l'utilisateur de créer ses documents multimédias,, de capturer des éléments de contexte et de partager les documents multimédias capturés via le dispositif mobile. Le composant Web est développé dans l'objectif de mettre en œuvre notre approche d'enrichissement semi-automatique des annotations basé sur l'utilisation des ressources génériques comme GeoNames⁴⁵ ou spécifique à l'utilisateur comme son profil social (voir section 5.4 du Chapitre 5) ; enfin, pour mettre en œuvre aussi notre approche d'exploitation du graphe social proposé dans le Chapitre 7. La Figure 38 présente un schéma général illustrant les importantes fonctionnalités du prototype.

⁴⁵ www.geonames.org/

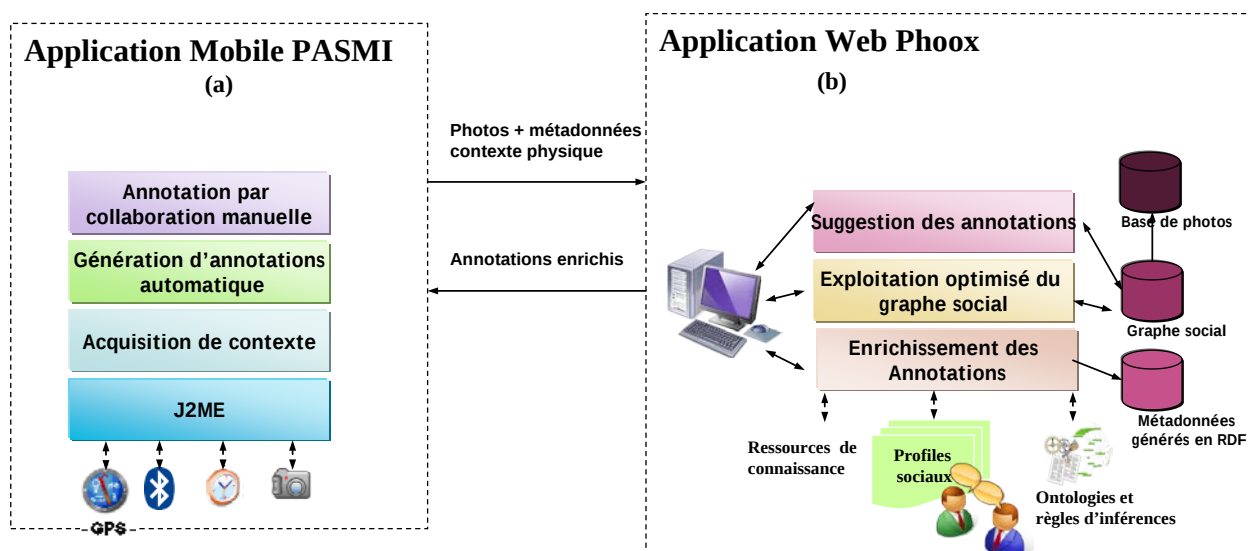


Figure 38. Architecture du prototype à deux composants: (a) *PASMi* (b) *Phoox*

8.2.1. Application mobile : *PASMi*

L'application mobile a été développée dans l'objectif de permettre aux utilisateurs de capturer leurs photos avec des métadonnées initiales (brutes). Les utilisateurs ont la possibilité d'organiser les photos capturées en des collections et de les annoter avec des tags ou des événements.

L'application mobile *PASMi* est déployée sur le dispositif mobile des utilisateurs. Cette application permet d'utiliser les senseurs de contexte pour apporter la première brique d'annotation riche de documents multimédias capturé par le dispositif mobile de l'utilisateur. L'application mobile *PASMi* permet, entre autres, l'annotation manuelle lors de la prise de vue ainsi que le partage et l'organisation des documents multimédias en collections (albums). Pour l'implémentation de l'application *PASMi*, nous avons utilisé J2ME⁴⁶ comme langage de programmation. *PASMi* doit être déployée sur les dispositifs des utilisateurs afin qu'ils puissent partager et capturer des photos munies des informations contextuelles.

Les principales fonctionnalités de l'application *PASMi* sont :

-*Acquisition de contexte de prise de vue* : Ce module permet aux utilisateurs de créer les documents multimédias et communiquer, lors de la prise de vue, avec les différents capteurs connectés au dispositif mobile. Les capteurs considérés par *PASMi* sont le GPS⁴⁷ pour

⁴⁶ Java 2 Micro Edition

⁴⁷ Global Positioning System

capturer la localisation, Bluetooth pour capturer les adresses Bluetooth proches lors de la prise de vue et l'horloge du système pour capturer le moment de la prise de vue.

-Génération d'annotations automatiques: Ce module communique avec le module acquisition de contexte et génère pour chaque photo des métadonnées. Nous avons utilisé un format texte simple de représentation des métadonnées car les parseurs *XML* ne sont pas intégrés dans la version actuelle du *MIDP*⁴⁸ (l'environnement d'exécution Java pour les petits dispositifs).

Les métadonnées décrites dans le fichier de représentation de métadonnées au format texte incluent l'identifiant de la photo, l'identifiant du créateur, les adresses Bluetooth des dispositifs autour lors de la prise de vue et les traces de partage des photos. Les traces de partage et d'annotation sont communiquées depuis le module '*Annotation manuelle collaborative*'.

La Figure 39 représente un exemple de représentation des métadonnées généré automatiquement par l'application *PASMi*, ceci, suite à une prise de vue et un partage de photo. La photo décrite dans la Figure 39 est identifiée par l'identifiant DSC32, créée par l'utilisateur d'identifiant 14165702aN05, prise à proximité de trois adresses Bluetooth: 001F6B579FAC, F00E5B679000 et 1F1F0B579F0A, enfin partagée avec les utilisateurs d'identifiants : 175248aN09 et 48325aN08.

<i>photo</i>
DSC032 14165702aN05 001F6B579FAC-F00E5B679000-1F1F0B579F0A 175248aN09-48325aN08

Figure 39. Exemple de fichier de représentation de métadonnées

Les autres éléments de contexte comme les coordonnées géographiques et la date sont stockées dans les métadonnées *EXIF* dans le même fichier (jpeg) de la photo.

-Annotation manuelle collaborative : Ce module permet aux utilisateurs de partager des photos et de collaborer lors de leur annotation. Dans le cadre de ce module, l'approche proposée par notre collègue Addisalem Negash [130] est intégrée dans l'application Mobile. Le travail d'Addisalem Negash est focalisé sur le partage d'informations dans un environnement ad-hoc et via des dispositifs mobiles. Dans le cadre de ses travaux de

⁴⁸ Mobile Information Device Profile

recherche, elle a proposé des algorithmes qui permettent le filtrage de la liste de diffusion des documents multimédias vue par chaque dispositif mobile appartenant à un réseau ad-hoc. Ce filtrage a pour objectif d'éviter la surcharge de l'environnement et de s'adapter avec l'interface minuscule des dispositifs mobiles. La politique de filtrage de diffusion est sensible à l'intérêt de l'utilisateur, à son contexte et son temps de connexion. Les intérêts sont représentés par des mots-clés de requêtes et/ou d'annotations manuelles. Les intérêts des utilisateurs sont déterminés à partir de la fouille de règles d'associations. Ces règles associent les intérêts des utilisateurs à leur contexte.

Grâce à l'intégration de cette politique dans notre application mobile, il serait plus facile à l'utilisateur de gérer ses documents multimédias et celles de son réseau de connaissance dans son dispositif mobile.

Pour la mise en œuvre, le prototype *PASMi* a été déployé sur : (i) un environnement émulateur (le Sun Wireless Toolkit) ; (ii) des téléphones de marque Sony Ericson w880i ; (iii) Sony Ericson w910i ; (iv) Nokia 5530 XpressMusic ; et (v) un PC portable doté d'un processeur de 2 GHz.

La Figure 40 illustre quelques captures d'écran des interfaces de l'application *PASMi*.

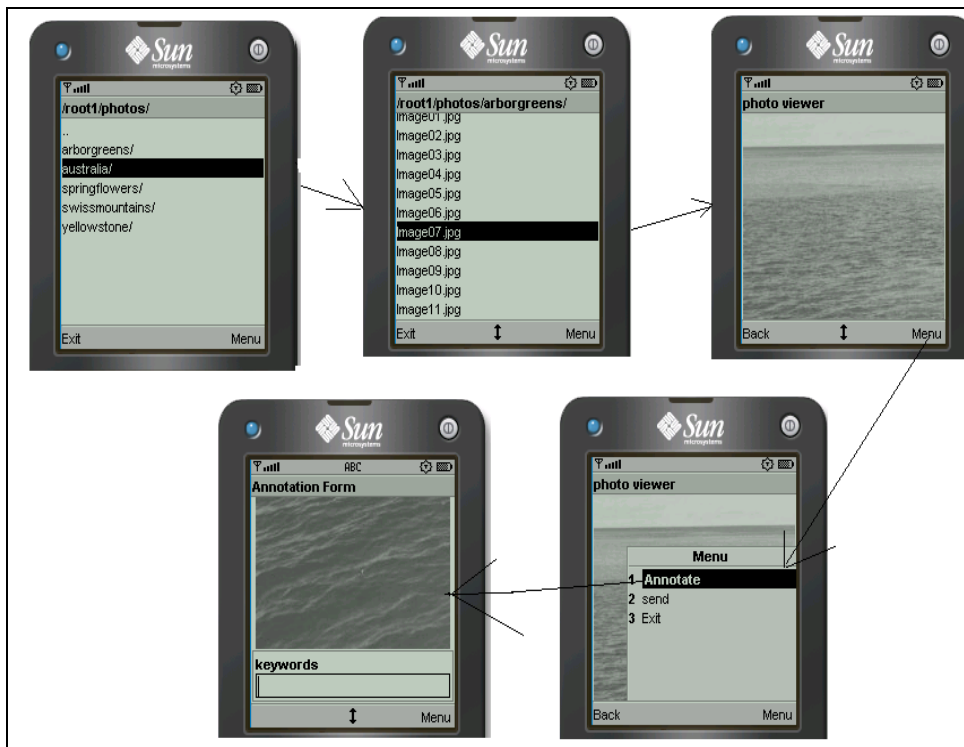


Figure 40. Quelques captures d'écran de l'application *PASMi*

PASMi synchronise les photos et leurs annotations à l'aide de l'application *Phoox* développée pour le web (Le composant web du prototype). La section suivante présente en détail les fonctionnalités fondamentales du composant web (*Phoox*).

8.2.2. Composante web pour l'organisation et la publication de photos : *Phoox*

L'objectif de la composante *Phoox* est de permettre à l'utilisateur de gérer ses photos, de les annoter plus aisément et ce en se servant de notre approche d'enrichissement automatique ou des suggestions proposées détaillé dans la section 5.4.

8.2.2.1. Fonctionnalités

L'application *Phoox*, dont l'interface de chargement de photos est présentée dans Figure 41, propose les fonctionnalités suivantes :

-*Enrichissement des annotations* : Ce module implémente la plateforme d'enrichissement des annotations présenté en détail dans la section 5.4 Chapitre 5. Nous avons implémenté le CCE (contexte de création enrichi) qui accède à des ressources sémantiques génériques pour enrichir les informations brutes de prise de vue par des données susceptibles d'être exploitées par l'utilisateur. La Figure 42 et la Figure 45 présentent respectivement les onglets Time et Environnement dans lesquels le système présente des informations plus riches comme le jour de la semaine, le mois, la saison et les conditions météorologiques.



Figure 41. Interface d'upload de l'application Phoox

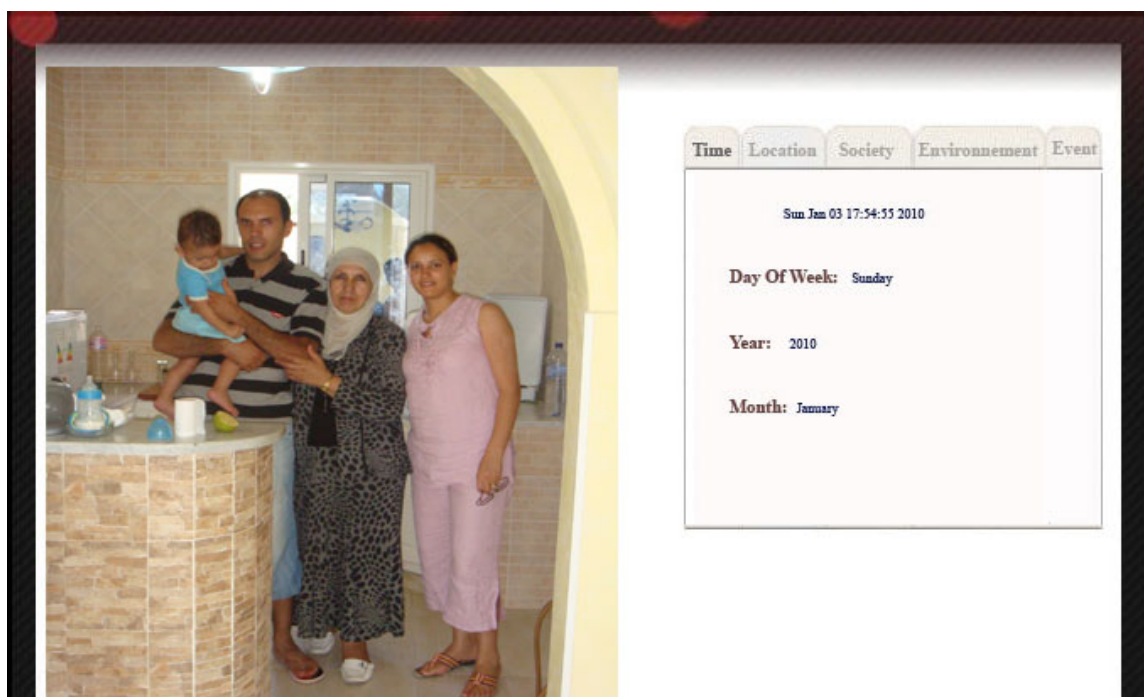


Figure 42. Interface d'annotation automatique de contexte temporel de prise de vue

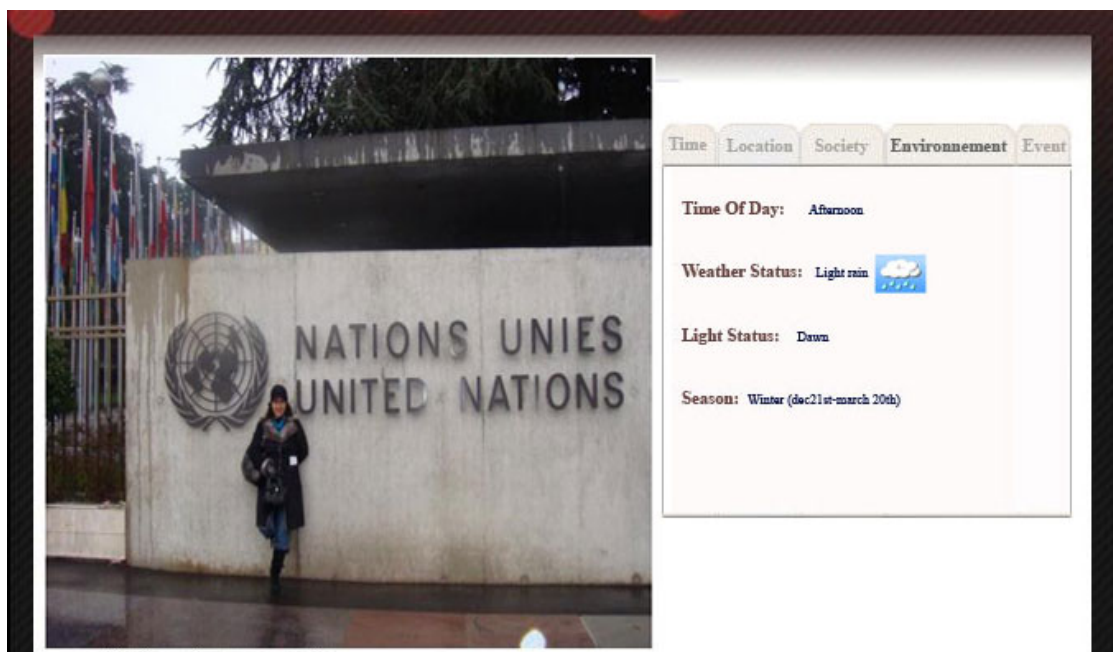


Figure 43. Interface d'annotation de contexte environnemental de prise de vue

L'onglet localisation (Location) est généré grâce à l'accès à deux modules : le *CCE* et le *CCP* (Contexte de Création Personnalisé) présentés respectivement dans la section 5.4.3 et section 5.4.4 Chapitre 5. La Figure 44 présente l'interface associée à l'enrichissement du contexte spatial (l'onglet location). Les premières informations sur la Figure 44 sont récupérées grâce à l'accès aux services de GeoNames et de Wikipédia. Dans l'exemple présenté dans la Figure 44, les objets proches selon le profil d'Alice⁴⁹ (i.e. *Nearby Objects according to your profil* (Alice) sur la Figure 44) sont obtenus par extraction des objets des utilisateurs les plus proches (à un seuil paramétré) des coordonnées géographiques du document multimédia en cours de consultation. Ces objets proches doivent appartenir au consultateur, aux personnes autour de la prise de vue ou à leurs réseaux sociaux. Dans notre exemple, un objet *Home* a été détecté dans le profil de Bob un participant à la prise de vue prise à un kilomètre près de la "place des Nations" à Genève. Alice (le consultateur dans l'exemple) a le choix d'accepter ou de décliner la proposition. Dans le cas d'acceptation, nous ajoutons la suggestion proposée comme une dimension sociale *SocialDimension* à cet utilisateur.

⁴⁹ Alice ici est la personne avec le profil actif

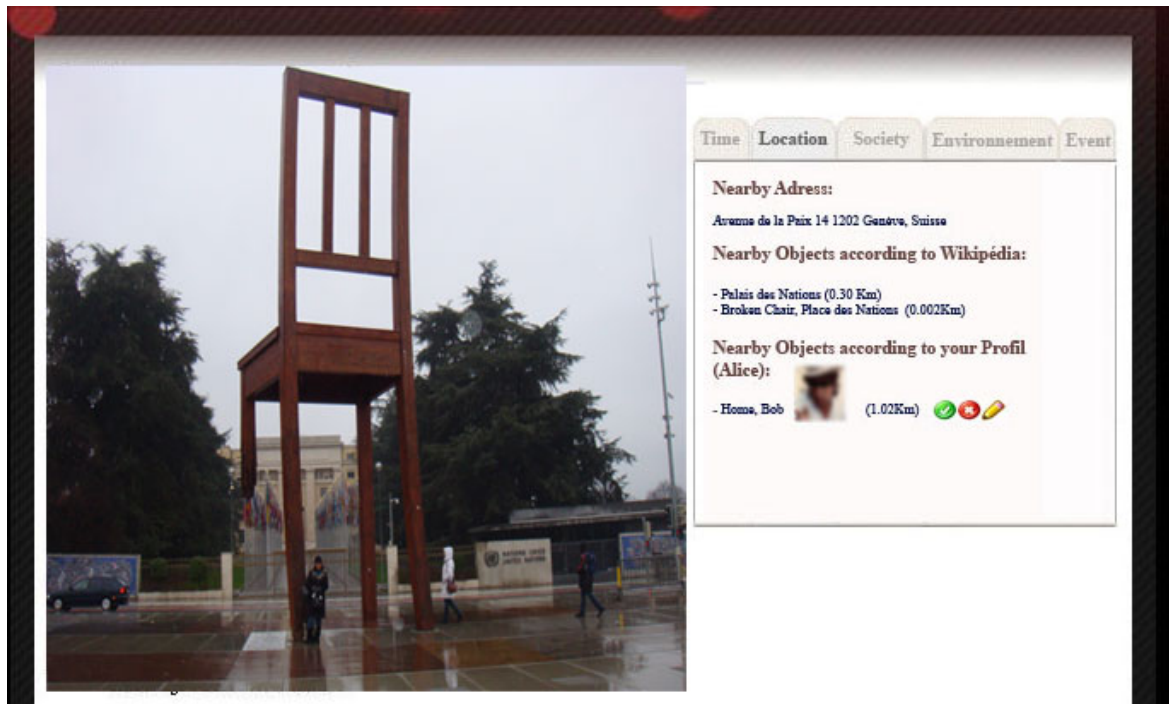


Figure 44. Interface d'annotation et suggestion de contexte spatial de prise de vue

- *Suggestion des annotations* : Dans ce module nous avons implémenté nos approches de suggestion de tags événements et tags identités des personnes dans les photos. L'utilisateur pourra choisir à travers l'interface homme-machine les propositions du système qui lui conviennent. Le système, quant à lui, stocke après validation effectuée par l'utilisateur, les annotations comme des tags valables pour le consultateur ou un groupe de consultateurs.

La Figure 45 représente l'interface correspondante au contexte social de la prise de vue de la photo. Dans cette photo, les dispositifs mobiles ayant des adresses Bluetooth : FO-OE-5B-67-90-00, 00-1F-6B-57-9F-AC, 1F-1F-0B-57-9F-0A, 00-1F-6B-57-9F-AC ont été capturés à proximité.

L'accès au profil social d'Alice permet la découverte des nœuds personnes correspondants aux adresses Bluetooth récupérées et mentionnées ci-dessus. Comme il est visible sur la Figure 45, Alice n'est pas directement liée à Bob mais grâce à l'accès au profil de Carol, la cousine à Alice, le système pourra découvrir l'ensemble des nœuds correspondants aux adresses Bluetooth des dispositifs mobiles capturés.



Figure 45. Interface d'annotation de contexte social de la prise de vue

Nous avons également élaboré une interface de suggestion d'événements correspondant à la prise de vue de la photo (illustré par le contenu de l'onglet Event dans la Figure 46). La suggestion des événements est effectuée grâce à notre algorithme de suggestion d'événements basé sur la découverte de règles d'association présentées en détails dans la section 6.3 du Chapitre 6. Dans l'exemple présenté par la Figure 46, le système propose les tags : fêter la nouvelle année "*celebrating new year*", voyage "*Trip*", visite d'amis "*Visiting Friend*" et anniversaire de Carol "*Carol's Birthday*". L'observateur accepte une ou plusieurs suggestions par un simple clic sur un bouton. Il peut aussi éditer un nouvel événement.

La suggestion des événements fait appel à un autre module appelé *Fouilles de règles d'association* qui permet d'extraire les informations importantes pour chaque utilisateur depuis la base de métadonnées et de préparer la base de données. Ensuite, le module *Fouilles de règles d'association* applique l'algorithme *APRIORI* proposé par Agrawal et al. dans [114] et [115] pour extraire des *itemSets* fréquents et de construire la base des règles d'association.

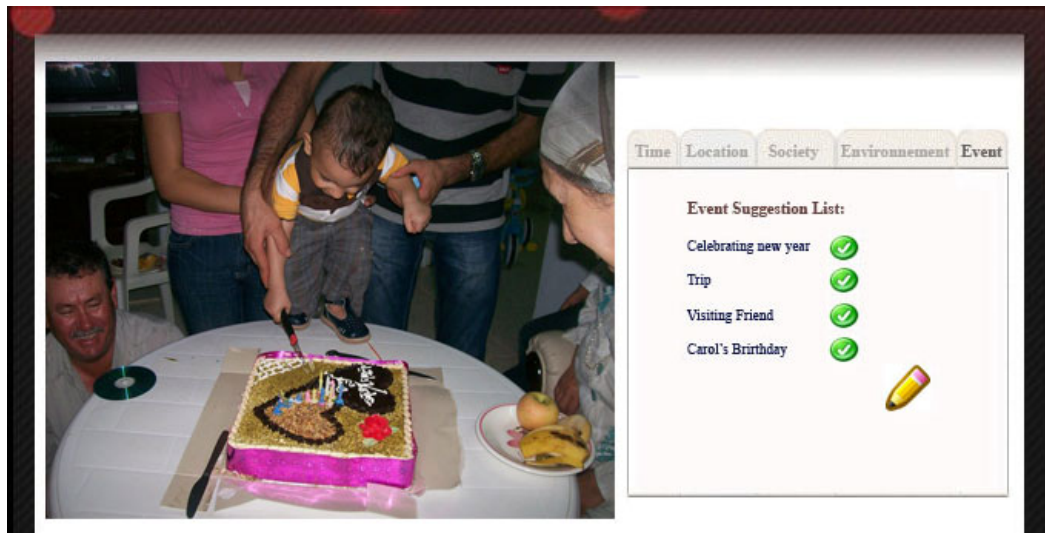


Figure 46. Interface de suggestion des événements

En identifiant des personnes sur les photos, le système propose selon l'observateur des dimensions sociales qui correspondent au chemin reliant l'observateur à la personne identifiée dans la photo.

Dans l'exemple de la Figure 47 la personne "Salwa" est identifiée. Ainsi, le système propose à Alice la dimension sociale ma tante "my aunt" grâce à l'enrichissement apporté par les règles d'inférence (voir section 5.3.3.2 du Chapitre 5 et annexe Annexe III).

Le mécanisme d'enrichissement des identités des personnes sur les photos est assuré en implémentant notre algorithme *SQO* de recherche optimisée dans le profil social, ce mécanisme est présenté dans la section 6.4 Chapitre 6.

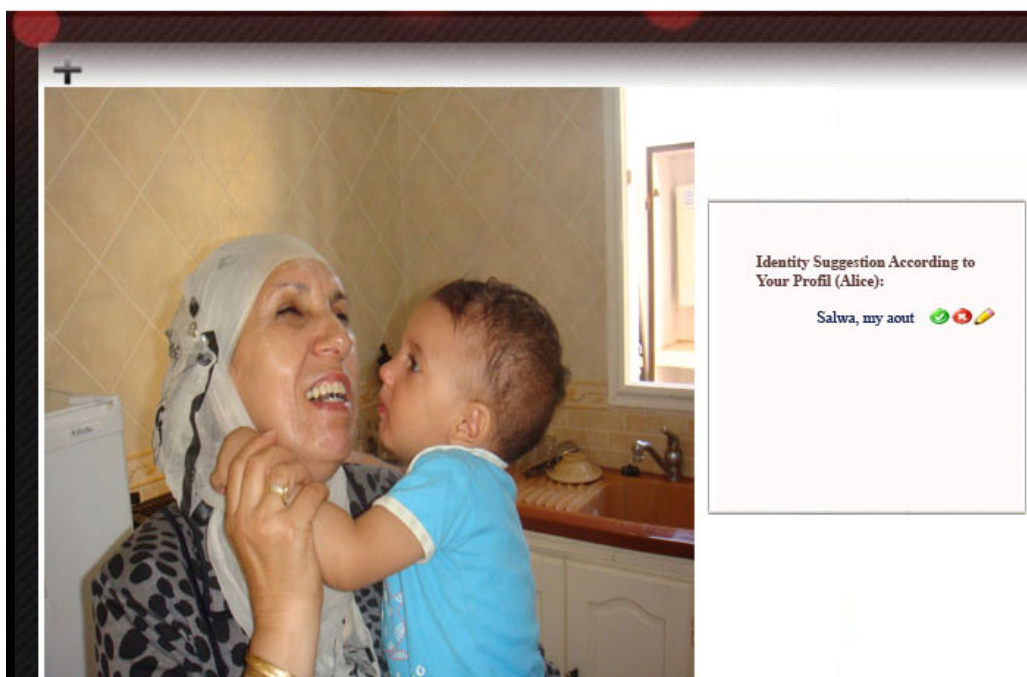


Figure 47. Interface de proposition des dimensions sociales

8.2.2.2. Les bases de données

L'architecture proposée repose sur un ensemble de base de données présenté dans ce qui suit :

-*Base de graphe social* : une représentation en RDF du graphe social. Cette base est pilotée par une API graphe, présentée dans la Figure 48 qui permet sa gestion.

-*Base de proximité sociale* : L'algorithme RSP de suggestion d'événements détaillé dans la section 6.3 du Chapitre 6 se sert de la proximité sociale afin de surmonter le problème de démarrage à froid du système de recommandation causé par l'absence ou la pauvreté de la base de documents multimédias propre à l'utilisateur. Comme solution à ce problème, nous avons mis en œuvre une base relationnelle représentant les valeurs de proximité sociale par couple d'utilisateurs. La base de proximité sociale est mise à jour périodiquement (i.e., lorsque le système est peu actif ou inactif⁵⁰).

-*Base de règles d'associations* : Cette base est nécessaire pour la production des recommandations d'événements (i.e. algorithme RSP).

D'un point de vu purement technique, l'application *Phoox* est développée en PHP, la communication avec la base RDF est assurée à l'aide de l'API RAP⁵¹. Les méthodes et les classes sont non transparentes à l'ensemble de l'application grâce à leur encapsulation par une API graphe que nous avons élaborée (voir Figure 48). Quant à l'édition des ontologies, nous avons utilisé le logiciel Semantic Work de la suite *Altova*⁵². Ce dernier propose une interface graphique et relativement simple pour l'édition et la vérification des ontologies en RDF/RDFS/OWL.

L'architecture de l'application *Phoox* se base sur le modèle d'architecture MVC (Modèle-Vue-Contrôleur) [131] afin de permettre la séparation entre les données, les traitements et la présentation.

L'organisation de la conception de l'application web *Phoox* en Modèle/Vue/Contrôleur (MVC) impose la séparation entre les données, les traitements et la présentation, ce qui engendre trois parties fondamentales dans l'application finale : le modèle, la vue et le contrôleur.

⁵⁰ Le mot inactif signifie que les utilisateurs ne sont pas entrain d'utiliser le système.

⁵¹ <http://www4.wiwiwiss.fu-berlin.de/bizer/rdfapi/>

⁵² <http://www.altova.com/semanticworks.html>

Dans le Modèle nous avons représenté le comportement de l'application selon deux grands traits: traitement des données et interaction avec la base de données. La base de données est gérée par des classes clés comme : *SocialProfil* , *SocialMultimediaDocuement*, *Event*, *Album*, *Time*, *Location*, etc.

Le diagramme de classe de l'ensemble de la couche Modèle de l'application est représenté dans la Figure 49. Ces classes représentent des types de nœuds du graphe social global. L'accès est assuré, par quatre classes principales qui sont : la classe *SocialGraph*, *Node*, *Attribut* et *Relation*. L'accès à la base de données est assuré uniquement par la réutilisation des méthodes de l'*API* graphe dont le diagramme de classe est représenté dans la Figure 48.

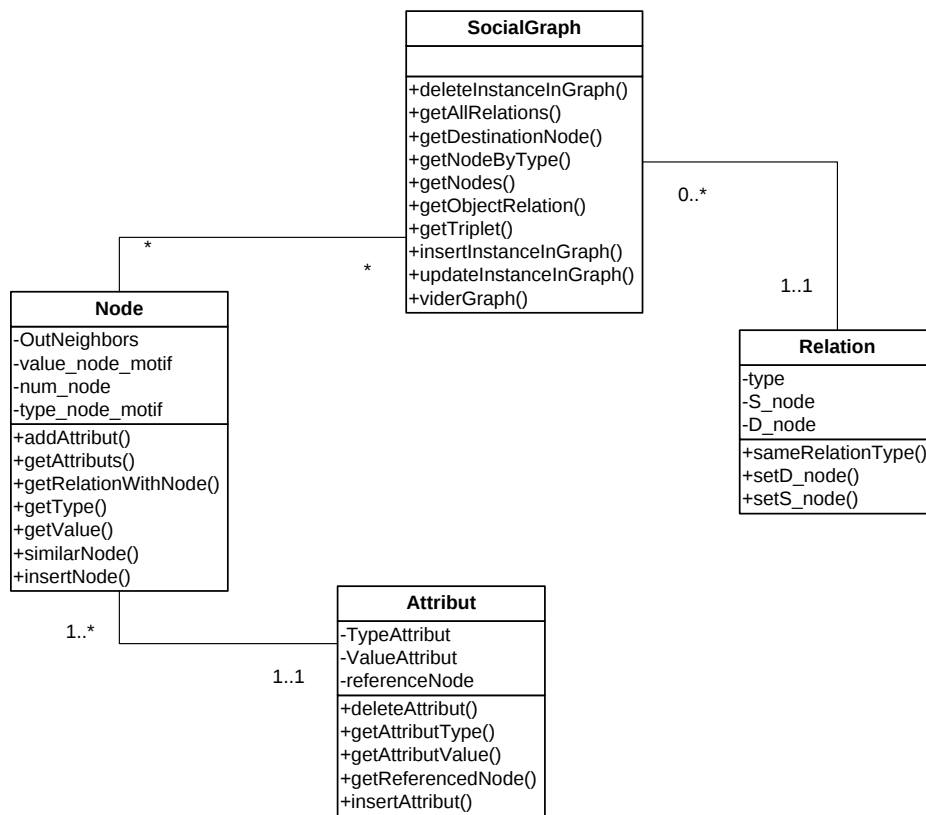


Figure 48. Diagramme de classe de l'API graphe

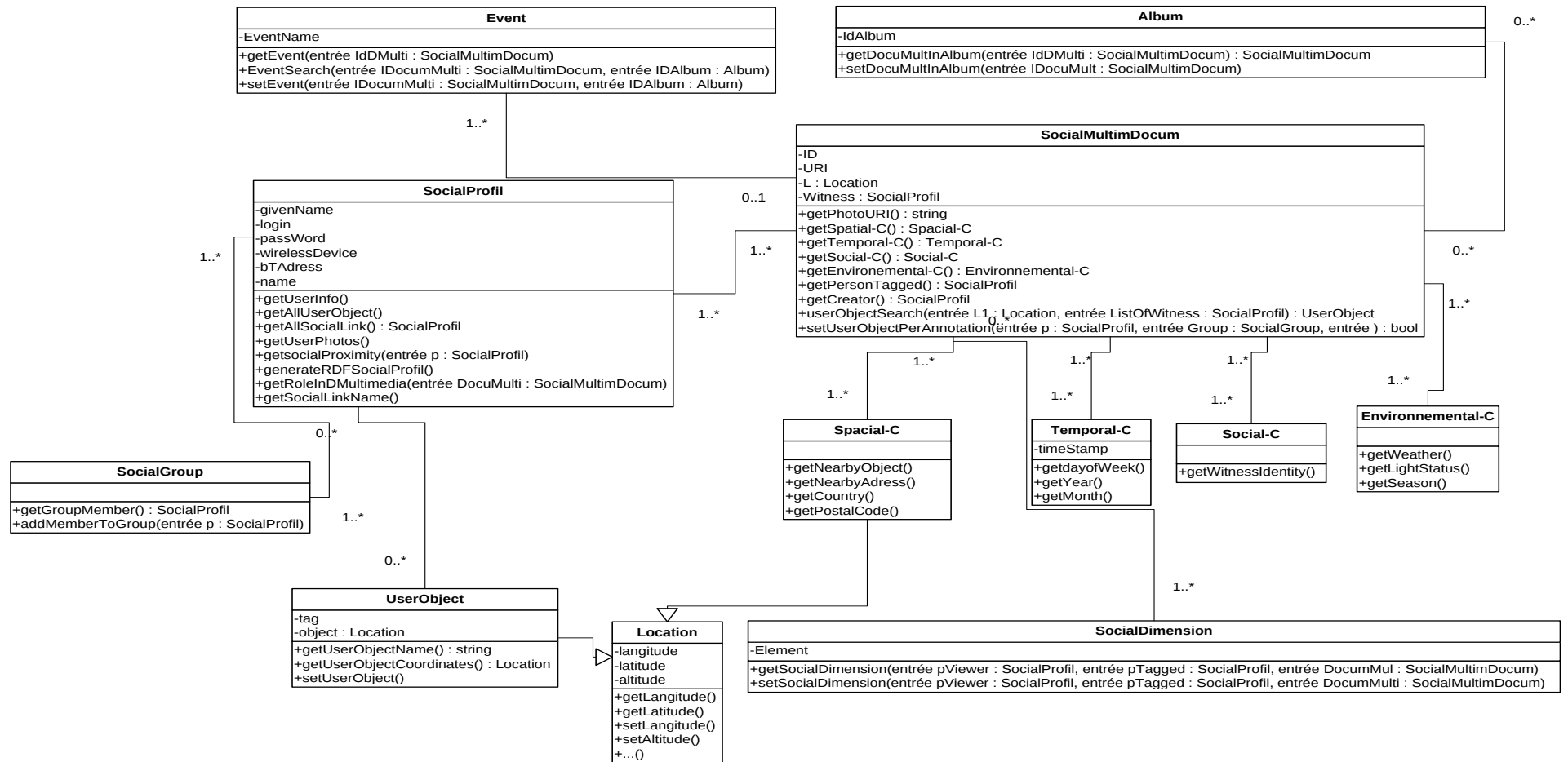


Figure 49. Diagramme de classes de l'application *Phoox*

8.2.2.3. Représentation et accès au profil social

Le profil social correspond à l'ensemble d'informations nécessaires au processus de suggestion (objet utilisateurs, événements et identité de personnes). Ces informations sont, comme nous les avons détaillées dans le Chapitre 5 de trois types : (i) informations personnelles, (ii) informations qui concernent les liens sociaux de l'utilisateur et ses sphères sociales et (iii) informations qui concernent ses objets géo-localisés. La Figure 50 représente une capture d'écran de l'interface de présentation et d'édition de la partie lien social du profil social de l'utilisateur Carol.

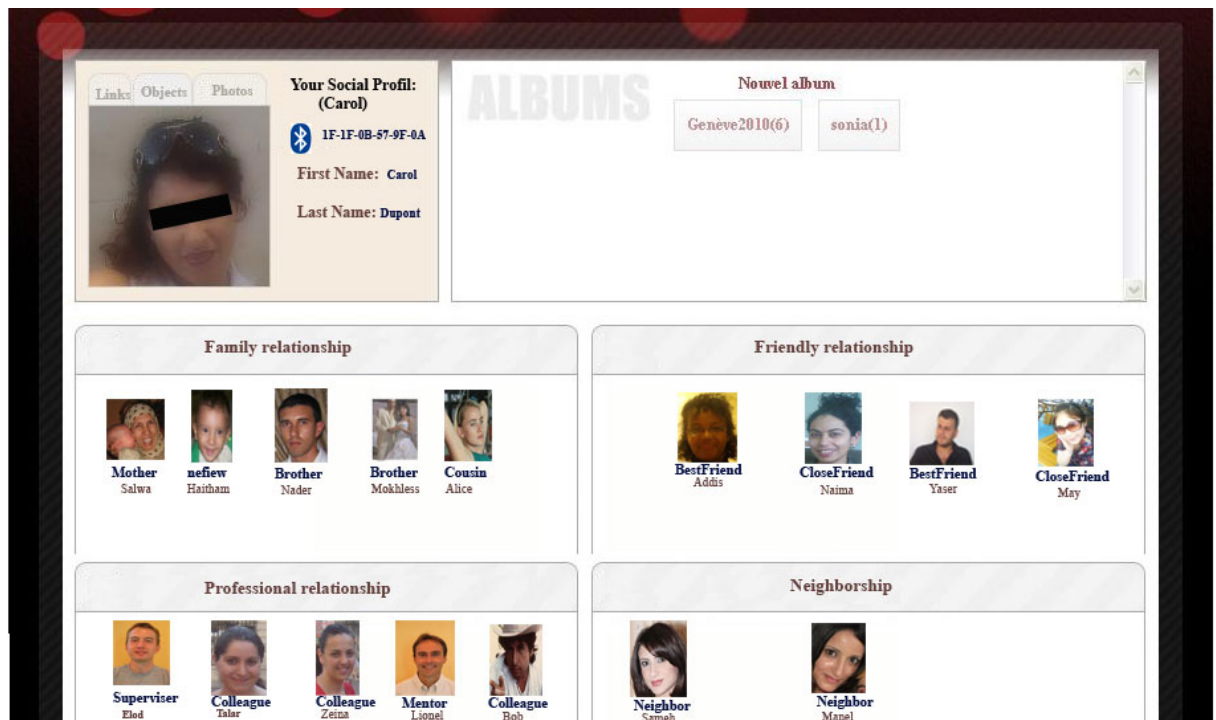


Figure 50. Interface de présentation et édition de profil sociale

La présentation de ces informations est faite grâce à un profil stocké en format RDF/XML dans la base de données et dont l'accès est effectué à l'aide de l'API graphe présenté ultérieurement. La Figure 51 représente une représentation en RDF/XML du profil social associé à Carol.

```

<foaf:Person rdf:ID="Dupont_Carol">
  <!-- personal informations -->
  <foaf:familyName>Dupont</foaf:familyName>
  <foaf:firstName>Carol</foaf:firstName>
  <acronym:carries rdf:resource="Device01"/>
  <!-- social network -->
  <foaf:know>
    <socialSphere:brother rdf:ID="nader_Dupont"/>
    <socialSphere:brother rdf:ID="mokless_Dupont"/>
    <socialSphere:BestFriend rdf:resource="Yaser_"/>
  </foaf:know>
  <!--social network category -->
  <socialSphere:FamilyRelationship>
    <foaf:member>
      <foaf:Person rdf:resource="nader_Dupont"/>
      <foaf:Person rdf:resource="mokless_Dupont"/>
    </foaf:member>
    <socialSphere:belongsTo rdf:resource="Dupont_Carol"/>
  </socialSphere:FamilyRelationship>
  <socialSphere:FriendlyRelationship>
    <foaf:member rdf:resource="Yaser_"/>
    <socialSphere:belongsTo rdf:resource="Dupont_Carol"/>
  </socialSphere:FriendlyRelationship>
  <!-- user objects-->
  <socialSphere:has rdf:resource="#UserObject01"/>
  <socialSphere:has rdf:resource="#UserObject02"/>
</foaf:Person>
<!-- other informations about resources described in the present social profil-->
<socialSphere:UserObject rdf:about="#UserObject01">
  <geonames:lat>34.7414</geonames:lat>
  <geonames:long>10.7567</geonames:long>
  <geonames:alt>0</geonames:alt>
  <socialSphere:tag>home</socialSphere:tag>
</socialSphere:UserObject>
<socialSphere:UserObject rdf:about="#UserObject02">
  <geonames:lat>34.7814</geonames:lat>
  <geonames:long>11.7607</geonames:long>
  <geonames:alt>0</geonames:alt>
  <socialSphere:tag>work</socialSphere:tag>
</socialSphere:UserObject>
<acronym:WirelessDevice rdf:about="Device01">
  <acronym:Bluetooth>00-1F-6B-57-9F-AC</acronym:Bluetooth>
</acronym:WirelessDevice>
<foaf:Person rdf:resource="nader_Dupont">
  <foaf:name>nader Dupont</foaf:name>
  <rdfs:seeAlso rdf:resource="http://www.exemple.com/Profilenaderlajmi.rdf"/>
</foaf:Person>

```

Figure 51. Représentation du profil social sous format RDF/XML

8.3. Expérimentation et évaluation des performances

Pour évaluer les performances de nos approches de correspondance entre requête et graphe social de recommandation, nous avons construit, en premier lieu, une collection de test provenant de vraies données utilisateurs. Nous avons mené, par la suite, trois expérimentations qui concernent: (i) l'approche de recherche dans le graphe social appelée *h-Pruning*; (ii) l'approche de suggestion de tags événements appelée *RSP* (iii) l'approche

de suggestion de tags identité des personnes appelée *SQO*. Les trois expérimentations sont décrites respectivement dans les sections qui suivent.

8.3.1. Construction de collection de test

Nous avons construit une collection de test à partir du service de médias sociaux Flickr⁵³. Avec plus de cent millions de photos et plus de huit millions d'utilisateurs, Flickr est un service de médias sociaux qui s'est développé rapidement durant ces dernières décennies ce qui lui donne l'avantage de devenir une des premières sources de référence et de test pour les chercheurs travaillant sur les réseaux sociaux et la gestion des photos. La raison la plus importante de cette attention portée sur Flickr est peut être due au fait qu'il se base sur le principe du tagging et qu'il met à la disposition des développeurs des interfaces d'accès automatique (*API*) faciles à utiliser et disponibles en plusieurs langages de programmation.

Nous avons utilisé l'API PHP de Flickr⁵⁴ disponible pour une utilisation non commerciale pour extraire des photos géo-localisées, des tags, des utilisateurs et des liens sociaux. Ainsi, celle-ci nous a permis de construire (i) des profils sociaux et (ii) le graphe social. Les profils sociaux construits comportent des liens sociaux et des informations sur les utilisateurs comme leurs avatars, leurs noms, leurs adresses. Quant au graphe social, nous l'avons construit à travers les photos géo-localisées, les tags ainsi qu'à travers les relations d'annotations et de création entre les utilisateurs et les photos

Le Tableau 8 récapitule des résultats statistiques quantitatifs sur ces données. En fait, nous avons extrait exactement 2851 photos géo-localisées. Les photos sont annotées par un nombre total de tags égal à 12 345. Nous avons extrait, à partir de Flickr, des identifiants d'utilisateurs et des liens sociaux exprimés explicitement par les utilisateurs eux mêmes. Au total, nous avons obtenu un réseau social de 101 899 relations qui relie 40 592 utilisateurs.

⁵³ www.flickr.com

⁵⁴ L'API de Flickr : <http://www.flickr.com/services/api/>

Tableau 8. . Statistiques sur la collection de données construite à partir de Flickr

Nombre de photos géo-localisés	Nombre de total de tags	Nombre d'utilisateurs	Nombre de liens sociaux explicites entre utilisateurs
2851	12 345	40 592	101 899

Les caractéristiques des tags dans Flickr sont présentées en détail dans [76]. Ces études montrent que, même si certaines photos sont annotées par plus de 50 tags, les statistiques montrent que 60% de l'ensemble de toutes les photos sont annotées par au minimum 1 à au maximum 4 tags. Ce qui correspond à la collection de données que nous avons réalisée. En effet, les photos extraites sont annotées en moyenne par 4 tags. Les types de ces tags⁵⁵ correspondent aux catégories WordNet de localisation, objets, personnes, actions (ou événement) et temps.

Les sections suivantes présentent les expérimentations menées en utilisant notre collection de test sur des profils de vrais utilisateurs extraits à partir de Flickr. À l'issue de ces expérimentations, nous avons prouvé l'efficacité de nos approches en termes de temps d'exécution.

8.3.2. Évaluation de l'approche de l'exploitation du graphe social

L'évaluation des performances de l'approche d'exploitation du graphe social s'est basée sur une comparaison entre notre algorithme *h-Pruning* (proposé dans la section 7.5.3 Chapitre 7) et l'algorithme initial proposé par C. Solnon [127], la comparaison s'est basée, particulièrement, sur le temps d'exécution. Il est nécessaire de mentionner que *h-Pruning* est une heuristique qui considère les appariements initiaux pour construire un espace de recherche réduit.

Pour mener à bien une expérimentation, il est important d'étudier les facteurs de variations de *h-Pruning*. En effet, comme nous avons discuté dans une étude théorique de la complexité de notre algorithme dans la section 7.5.3.6 du Chapitre 7, *h-Pruning* a une complexité polynomiale qui du nombre de nœuds dans l'espace de recherche réduit. Ce temps dépend du temps de construction de cet espace. Ce temps est fonction de

⁵⁵Les types des tags en comparant avec des mesures sémantiques les tags aux mots-clés location, event, time, object et person.

l'excentricité de chaque nœud connu et du nombre de ces nœuds connus. Le temps d'exécution qui suit la construction de cet espace de recherche réduit est fonction du nombre de nœuds dans le graphe motif G_M .

Selon les facteurs de variation de l'exécution de *h-Pruning*, nous avons développé des expérimentations.

Dans un premier temps, nous avons mesuré le temps de réponse en variant le nombre de nœuds dans le graphe social G_S et en fixant le nombre de nœuds dans le graphe requête G_M à 5, le nombre d'arcs à 4 et le nombre d'appariements initiaux à 2.

Le Tableau 9 détaille les résultats de l'expérimentation : par exemple la différence entre les comportements des deux courbes devient très large à partir de 2843 nœuds dans le graphe social G_S . En effet, le temps d'exécution de l'algorithme initial est de 3360 ms tandis que *h-Pruning* ne dépasse pas un temps d'exécution de 850 ms. Au passage à l'échelle (i.e., nous avons mesuré dans l'expérimentation un maximum de 5000 nœuds), la différence entre les deux algorithmes en termes de temps d'exécution devient remarquable. En effet, le résultat n'est pas surprenant car *h-Pruning* a une complexité polynomiale en fonction du nombre de nœuds et d'arcs dans le graphe motif G_M il dépend aussi de la densité du graphe social G_S . En contrepartie, notre heuristique est uniquement applicable lorsqu'il existe un ou plusieurs appariements initiaux (voir section 7.5.3 Chapitre 7) ce qui est le cas adopté dans l'expérimentation.

Tableau 9. Temps d'exécution des deux algorithmes en fonction du nombre de nœuds dans G_S

GS	temps d'exécution de l'algorithme de C.Solnon	temps d'exécution de <i>h-Pruning</i>
50	0	120
100	0	200
250	20	280
500	70	450
859	160	458
1000	300	454
1241	200	459
1558	950	650
2843	3360	850
3880	5270	1461
4228	5270	1872
5000	7350	2440

La lecture de la courbe représentée par la Figure 52 illustrant l'expérimentation montre clairement le comportement différent des deux algorithmes. La courbe ayant des nœuds

triangles étant celle qui illustre le comportement de *h-Pruning*. La courbe ayant des nœuds rectangles est celle qui illustre le comportement de l'algorithme initiale de C. Solnon.

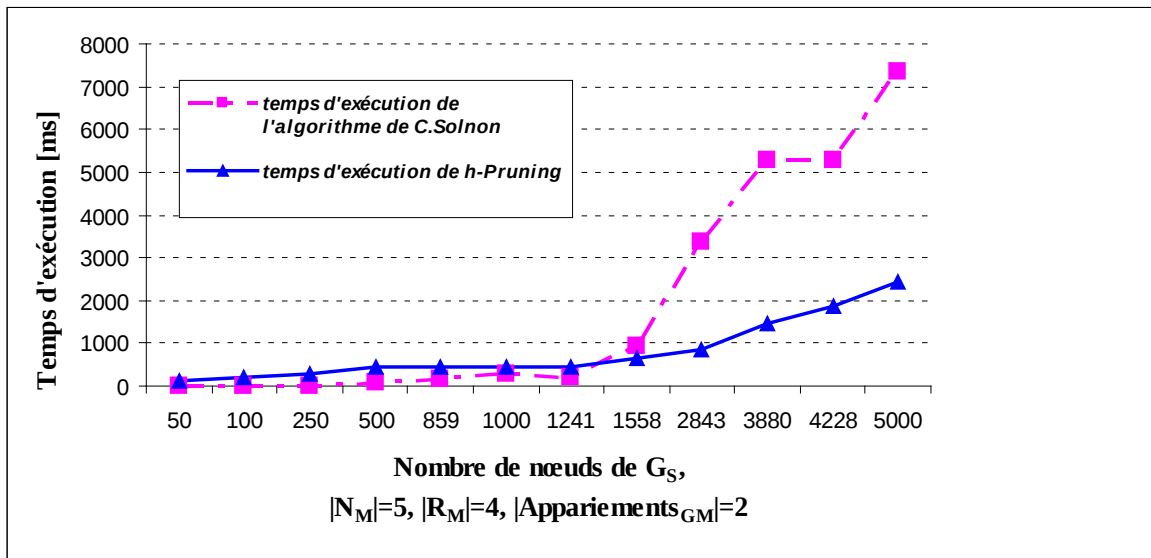


Figure 52. Temps d'exécution des deux algorithmes en fonction du nombre de nœuds dans G_S

Une deuxième série d'expérimentations a été menée dans l'objectif de vérifier le comportement des deux algorithmes vis-à-vis de la variation du nombre de nœuds dans le graphe motif G_M . Nous avons alors varié le nombre de nœuds dans le graphe motif de 2 à 20 et nous avons fixé le nombre de nœuds dans le graphe social à 1000 et un nombre d'arcs à 5000.

Le Tableau 10 représente les résultats de l'expérimentation. La différence entre les deux algorithmes devient importante à partir de 14 nœuds dans le graphe motif. Le temps d'exécution de *h-Pruning* reste raisonnable pour une requête de 20 nœuds (i.e., 2,578s) pour 8,542 (s) pour l'algorithme de C. Solnon.

Tableau 10. Temps d'exécution des algorithmes en variant le nombre de nœuds dans G_M

GM	temps d'exécution de l'algorithme de C.Solnon (ms)	temps d'exécution de notre algorithme <i>h-Pruning</i> (ms)
2	0	0
5	190	192
8	490	441
10	740	696
12	960	952
14	1530	1288
18	7408	2088
20	8542	2578

La courbe illustrée dans la Figure 53 montre la différence entre le comportement des deux algorithmes lors de la deuxième série d'expérimentations.

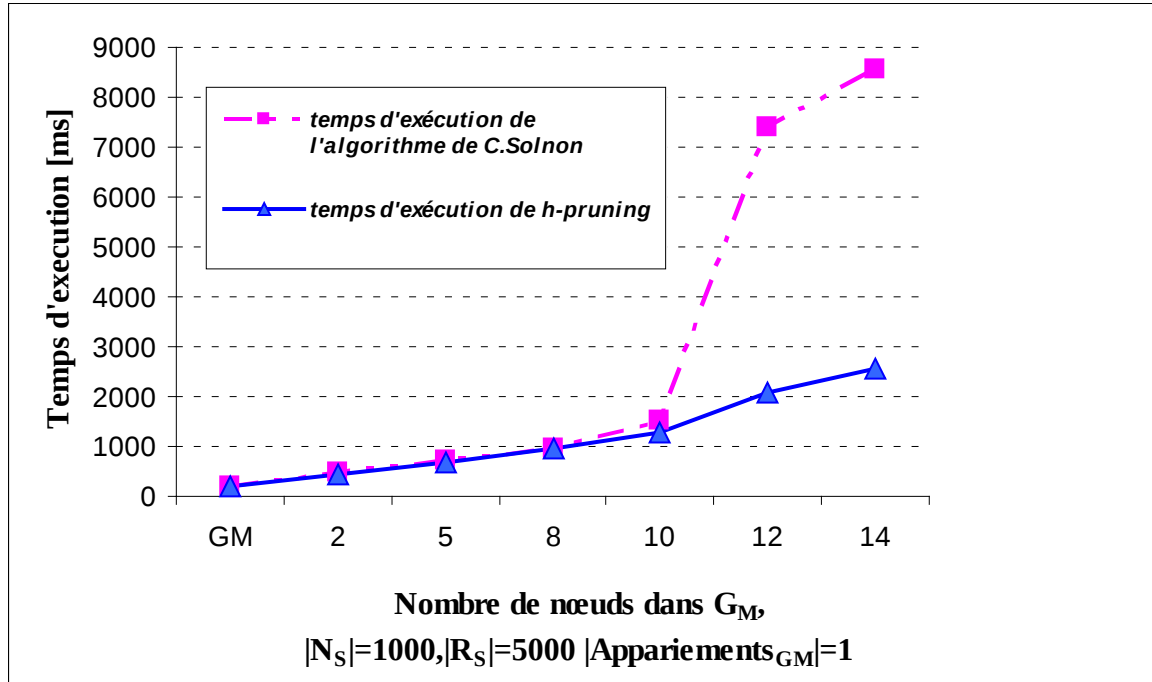


Figure 53. Temps d'exécution des deux algorithmes en fonction du nombre de nœuds dans G_M

En conclusion, comme illustré par des graphes et démontré par l'étude de complexité :

- (i) plus que le nombre d'appariements initiaux augmente plus la différence entre les deux algorithmes devient importante. Par contre, si le nombre d'appariements initiaux est nul h -Pruning se ramène à l'algorithme initial de C. Solnon. Le dernier cas n'est pas envisageable dans le cadre de la construction des requêtes. En effet, nous avons quasiment toujours au moins un appariement initial, ce qui apporte une valeur ajoutée à l'heuristique proposée.

- (ii) le temps d'exécution de h -Pruning est proportionnel à la taille du graphe motif car nous construisons un espace de recherche réduit proportionnel à cette taille ; ce qui résulte que seul l'espace de recherche réduit est exploité au lieu de la totalité du graphe social. En effet, la série d'expérimentations a prouvé qu'en variant le nombre de nœuds dans le graphe motif, le temps d'exécution augmente. En pratique, la taille de graphe motif n'atteint pas une grande taille puisqu'elle représente une requête.

8.3.3. Evaluation de l'approche de suggestion des dimensions sociales

L'évaluation des performances de l'approche de suggestion des dimensions sociales s'est basée sur une comparaison de notre algorithme *SQO* proposé dans la section 6.4 du

Chapitre 6 avec l'algorithme *BFS* (Breadth First Search) [119]. La comparaison que nous souhaitons étudier met en évidence la performance de notre approche en terme de temps d'exécution vis-à-vis à l'algorithme *BFS*. Nous avons implanté les deux algorithmes. Rappelons que notre algorithme se base sur un ensemble d'heuristiques. La première heuristique est une recherche en largeur. La deuxième est la considération de deux niveaux et pas toute la profondeur du graphe. La troisième consiste à limiter la recherche sur quelques sphères sociales suivant des règles d'associations.

Pour mener à bien notre expérimentation, il est important de préciser les facteurs de variation de notre algorithme. Selon l'analyse théorique de complexité de notre algorithme *SQO* présenté dans section 6.4.3 du Chapitre 6, le temps d'exécution de notre algorithme varie selon le nombre de nœuds voisins à chaque nœud.

La Figure 54 illustre les résultats d'exécution des deux algorithmes: *BFS* et *SQO*. Nous avons exécuté des calculs de chemin dans des réseaux sociaux représentés par des triplets RDF à l'aide des deux algorithmes. Nous avons mesuré le temps d'exécution des requêtes en fonction de la variation du nombre de triplets RDF. L'expérimentation menée s'est déroulée sur une machine avec une unité centrale de 2,10 GHz et de 2Go de RAM. Le jeu de données choisi correspond à un ensemble de profils distribués sur le web partant du profil de "*Libby Miller*"⁵⁶. Nous avons varié la taille du graphe *FOAF* de 153 triplets à 1446 triplets. Le temps d'exécution pour l'algorithme *BFS* a varié de 1,594s pour 153 triplets à 30,391s pour 1446 triplets.

⁵⁶ <http://swordfish.rdfweb.org/people/libby/rdfweb/webwho.xrdf>

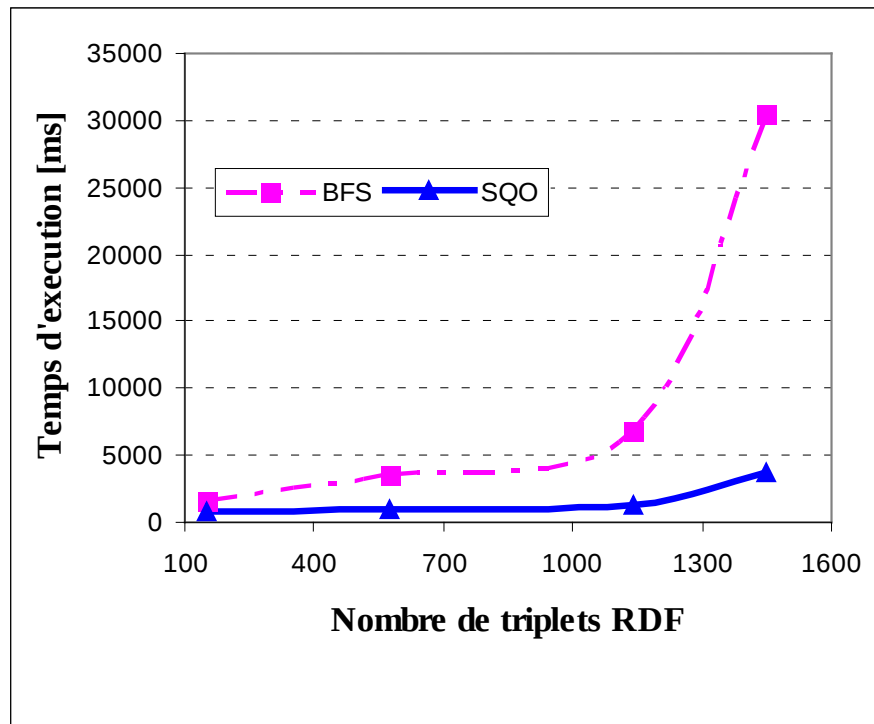


Figure 54. Temps d'exécution des deux algorithmes BFS et SQO en fonction de la taille du graphe

La lecture de la courbe représentée par la Figure 54 donne l'allure d'une fonction exponentielle.. Une étude de complexité bien détaillée de l'algorithme SQO est présentée dans la section 6.4.3 du Chapitre 6. La lecture de la courbe sur la Figure 54 montre que l'exécution de notre algorithme est beaucoup plus rapide que celui de *BFS* lors de graphes dépassent la taille d'une centaine de triplets.

8.3.4. Evaluation de l'approche de suggestion d'événements

Nous avons proposé dans la section 6.3 du Chapitre 6 l'algorithme RSP permettant la suggestion des recommandations d'événements dans le but d'assister les utilisateurs lors de l'annotation des photos. Rappelons que notre algorithme RSP prend en considération les règles d'association préparées et existantes pour chaque utilisateur ainsi que la proximité sociale, ceci, pour considérer de nouvelles règles d'association en vue de proposer des événements à un annotateur. Avant de proposer un protocole de test de l'algorithme, il est important d'étudier les facteurs de variation des résultats d'exécution de cet algorithme. En effet, le temps d'exécution d'un tel algorithme dépend de la variation de nombre de règles effectuées par personne, par le nombre d'utilisateurs ayant une proximité sociale élevée

avec l'utilisateur étudié, par le seuil de recommandation⁵⁷ préfixé et le seuil de proximité sociale qui en dessous desquels la proximité sociale ne sera pas considérée.

En ce qui concerne le temps d'exécution de l'algorithme de recommandation des événements de RSP, nous n'avons pas jugé nécessaire sa mesure. L'algorithme en tant que tel s'exécute en $O(|P_{seuil}| \cdot Avg(R(u,c)))$ avec $|P_{seuil}|$ le nombre d'utilisateurs ayant une proximité sociale $\geq$ à un seuil avec u et $Avg(R(u,c))$ est la moyenne de temps d'exécution nécessaire pour retrouver les règles d'association correspondantes à un contexte c . Cette moyenne est négligeable ; et elle correspond à la comparaison entre un contexte donné en paramètre et le contexte de chaque règle. De plus, la production des règles est faite a priori de manière cyclique et non lors de l'exécution de l'algorithme en temps réel.

Nos expérimentations se sont basées, alors, sur l'argumentation du choix de la considération des paramètres, notamment (i) la cooccurrence entre le contexte spatio-temporel et l'événement ainsi que (ii) la notion de proximité sociale et son importance dans les services de médias sociaux.

Nous avons mené une première série d'expérimentations pour prouver l'existence d'une cooccurrence entre le cadre spatio-temporel de prise de vue et les événements de manière répétitive. Nous avons, pour cela, extrait à partir de Flickr et pour un même utilisateur l'ensemble des photos géo-localisées dont il est le propriétaire.

Pour identifier l'événement depuis l'ensemble des tags décrivant chaque photo, nous avons comparé chaque tag et titre décrivant chaque photo au mot-clé '*event*' via des mesures de similarités sémantiques. Seul le mot-clé le plus similaire à '*event*' est considéré⁵⁸. Pour le calcul de similarité, nous avons dans un premier temps expérimenté plusieurs mesures de similarité sémantiques, à savoir : (i) celle de Wu & Palmer [52] ; (ii) celle de Jing & Conrath [55]; (iii) celle de Lin [53]; (iv) celle de Resnik [132]. Nous avons ensuite choisi d'utiliser Wu & Palmer pour sa meilleure performance vis-à-vis aux autres mesures de similarité sémantiques. Par exemple pour le mot-clé '*birthday*' selon Wu & Palmer la similarité entre '*event*' et '*birthday*' est égale à 0.533. Selon Resnick elle est égale à 0.516. Selon Jing & Conrath elle est égale à 0.089 et selon Lin elle est égale à 0.085. Pour cette raison, nous avons utilisé une bibliothèque java existante implémentant la plupart des mesures sémantiques existantes et manipulant l'accès à WordNet 2.1. Cette

⁵⁷ Le seuil de recommandation détermine le nombre de propositions recommandés à l'utilisateur

⁵⁸ Celui qui donne la valeur de similarité maximale parmi tous les autres

bibliothèque au nom Java WordNet::Similarity (beta.11.01) est proposée par Sussex University⁵⁹.

Le tag le plus similaire est considéré comme le plus représentatif de l'événement. Pour décrire la localisation et le temps, nous avons utilisé, en premier temps, notre plateforme d'enrichissement décrite dans la section 5.4.3 du Chapitre 5 pour enrichir les localisations et la notion de temps. Pour le temps, nous n'avons considéré que le moment de la journée en renseignant une des valeurs de l'ensemble {"*Afternoon*", "*Early morning*", "*Evening*", "*Late night*", "*Morning*" et "*Night*"}. Quant à la localisation, en raison de nombre limité de requêtes posées à la base GeoNames⁶⁰, nous n'avons malheureusement pas pu obtenir des informations concernant la localisation pour l'intégralité de la base de 960 photos. Pour remédier ce problème, nous avons octroyé le même principe que les tags qui décrivent l'événement en comparant des tags avec le mot-clé '*location*'.

Nous avons considéré une base de 589 utilisateurs et de 960 lignes de données. En appliquant l'algorithme de fouille de données APRIORI avec un MinSupp=0.2 et un MinConf=1.0 nous avons obtenu 343 règles au total spécifiques par utilisateur.

Le graphique représenté par Figure 55 représente la répartition des règles par nombre d'utilisateurs. Nous avons représenté le nombre de règles extraites pour chaque utilisateur et le nombre de personnes ayant x règles. La lecture du graphique montre que 74 utilisateurs parmi 589 utilisateurs possèdent au moins 2 règles spécifiques soit un pourcentage de $\sim 12,56\%$. Parmi ces utilisateurs un nombre considérable soit 49 utilisateurs ont 4 règles.

⁵⁹ <http://www.cogs.susx.ac.uk/users/drh21/>

⁶⁰ On ne pouvait pas poser plus que 100 requêtes pour un compte non professionnel

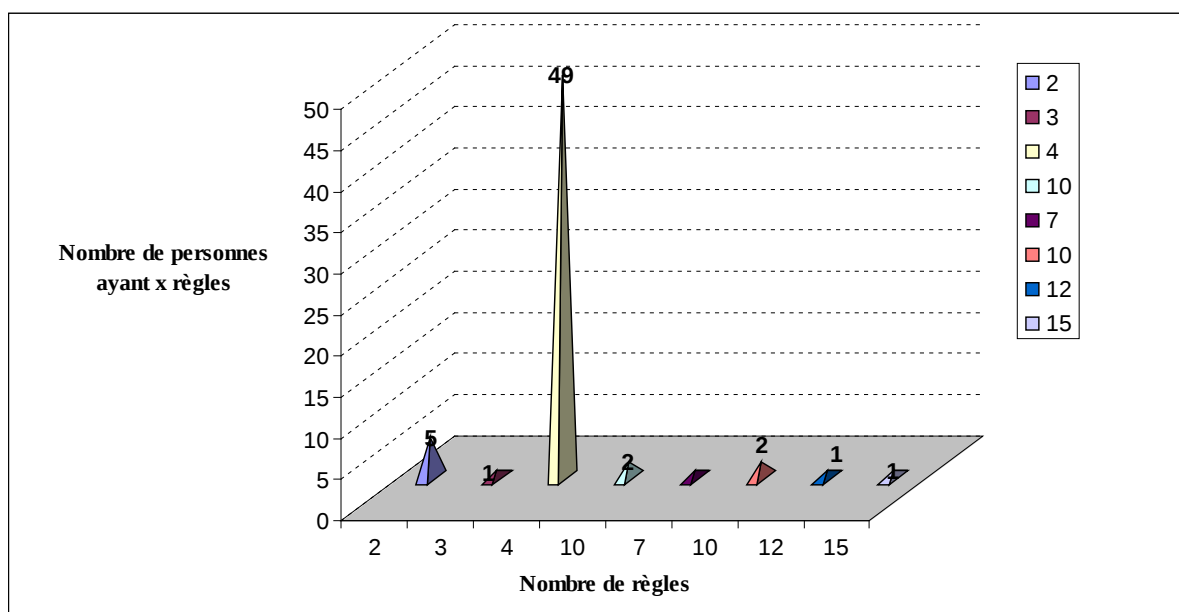


Figure 55. Nombres d'utilisateurs ayant x règles produites

L'expérimentation a prouvé l'efficacité de notre hypothèse de départ. Nous avons obtenu un nombre considérable de règles à savoir, un total de 343 règles sur une base de 960 photos avec un MinSupp de 0.2. En se basant sur les règles spécifiques pour chaque utilisateur, nous remarquons que la répartition de ces règles n'est pas équilibrée. Pour cette raison la deuxième partie de l'algorithme (e.g. se baser sur les règles des utilisateurs ayant une proximité sociale avec un utilisateur n'ayant pas ou peu de règles) est importante afin de proposer des recommandations plus riches.

Une deuxième série d'expérimentations a été menée, alors, pour prouver l'importance de la notion de proximité sociale et décider la limite de la valeur du seuil de proximité sociale considéré. L'expérimentation a été appliquée sur la même collection de données issues de Flickr. Nous avons mesuré la proximité sociale entre utilisateurs en prenant en considération les éléments existants dans la collection de test : la collaboration dans l'annotation et les relations d'amitié explicités. En considérant ces deux paramètres, nous avons dressé une matrice carrée ayant pour dimension le nombre d'utilisateurs soit 40 592 utilisateurs dont une grande partie ne possédant (seulement 589 utilisateurs) de photos géolocalisées sur Flickr. Un exemple de cette matrice est représenté dans le Tableau 6 de la section 6.3.4 du Chapitre 6.

Rappelons que nous avons considéré le concept de proximité sociale comme étant une valeur tenant compte des éléments suivants :

- Le nombre de cooccurrence sur les photos (le nombre de fois où les deux utilisateurs sont tagués ensemble sur une même photo) ;
- Le nombre de leurs collaborations dans l'annotation manuelle des photos (i.e. le nombre de fois où l'utilisateur u_1 a commenté ou annoté les photos de u_2) ; et
- Le nombre de fois où l'utilisateur u_1 a partagé avec u_2 des photos.
- Les liens explicites entre les nœuds du réseau social.

Donc, pour calculer la valeur de proximité sociale, nous avons proposé la fonction multicritère ci-dessous :

$$proxSocial(u_1, u_2) = \alpha \cdot freqPart(u_1, u_2) + \beta \cdot freqAnn(u_1, u_2) + \delta \cdot freqCoocPh(u_1, u_2) + \lambda \cdot N(u_1, u_2)$$

Avec $\alpha + \beta + \delta + \lambda = 1$

Dans cette partie d'expérimentations et en fonction des données mises à notre disposition nous avons considéré que :

$$proxSocial(u_1, u_2) = \beta \cdot freqAnn(u_1, u_2) + \lambda \cdot N(u_1, u_2)$$

Avec $\beta + \lambda = 1$

Nous avons considéré $\beta=0.5$ et $\lambda=0.5$ pour donner une considération égale aux deux critères de la fonction.

La courbe illustrée par la Figure 56 décrit le nombre d'utilisateurs ayant une proximité x avec d'autre(s) utilisateurs(s). La lecture de la courbe montre que la plupart des utilisateurs ont une proximité sociale qui varie entre 0.56 et 0.47. Donc le seuil de proximité sociale pris en compte doit être compris entre ces deux valeurs.

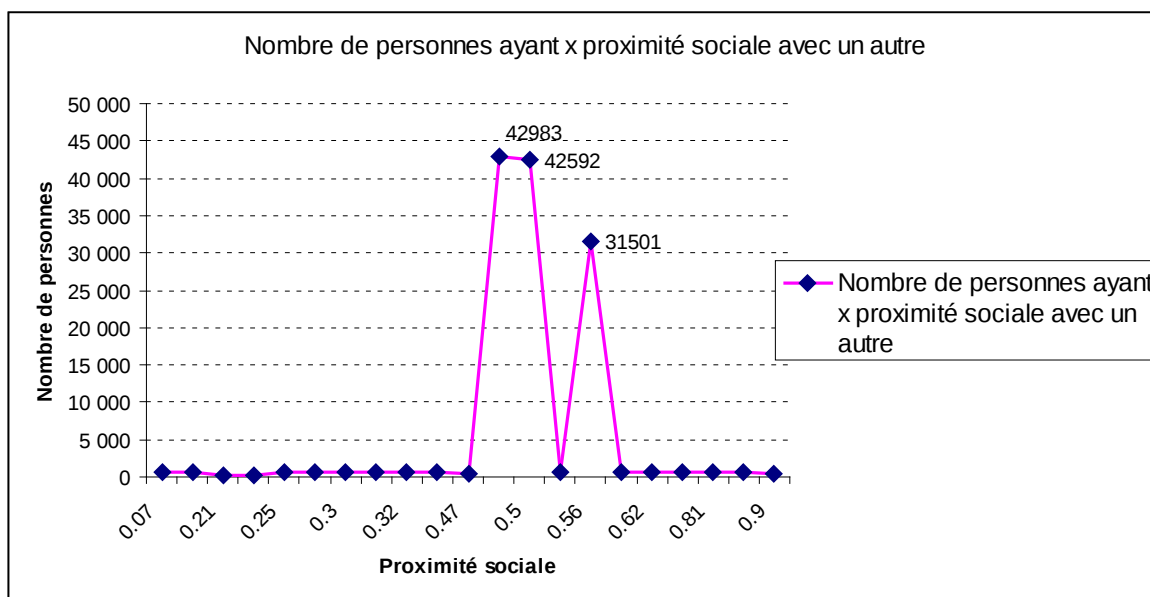


Figure 56. Nombre d'utilisateurs ayant une proximité sociale donnée avec d'autres utilisateurs

8.4. Discussion

Afin de valider notre travail nous avons développé en premier lieu une application mobile dans un environnement réel. J2ME a été utilisé comme langage de programmation. La version actuelle du MIDP⁶¹ (l'environnement d'exécution Java pour les petits dispositifs) ne supporte pas l'accès aux métadonnées sous format XML ou RDF/XML comme DublinCore, MPEG-7 ou notre modèle *SeMAT*. Il ne fournit pas un parseur XML. Pour cette raison nous avons stocké les traces de partage de photos et les commentaires des utilisateurs dans le champ commentaire des métadonnées EXIF. Les données EXIF souffrent malheureusement de l'inconvénient du non extensibilité. L'application Web lit les métadonnées EXIF et génère par la suite les métadonnées enrichies.

En ce qui concerne l'utilisation des services web nous avons utilisé les services GeoNames et Wikipédia pour enrichir les coordonnées géographiques. Il est à noter que la base de GeoNames n'est pas très complète par rapport à la base de Wikipédia qui comporte plus d'articles et donc s'avère plus fiable.

⁶¹ Mobile Information Device Profile

8.5. Résumé

Dans ce chapitre, nous avons présenté le prototype développé pour valider notre approche d'enrichissement d'annotations basée sur des ressources génériques et des ressources spécifiques à l'observateur de la photo que nous l'avons appelé 'profil social'. La mise en œuvre des approches proposées a été suivie par une série d'expérimentations. La première série a prouvé l'efficacité de notre approche *h-Pruning* par rapport à l'algorithme existant de C. Solnon en termes de temps d'exécution. La deuxième a prouvé le rationalisme de nos hypothèses qui ont mené à la construction de notre approche de suggestion des "tags événements" à savoir la notion de proximité sociale et la cooccurrence entre contexte spatio-temporel et événement. La dernière a prouvé l'efficacité de notre approche SQO par rapport à l'approche existante de BFS en termes de temps d'exécution.

Chapitre 9. Conclusion et perspectives

9.1. CONCLUSION 193

9.2. BILAN DES CONTRIBUTIONS..... 194

9.3. PERSPECTIVES 197

9.1. Conclusion

Les travaux présentés dans ce travail de thèse, abordent essentiellement, la sensibilité sociale des annotations et l'exploitation du graphe social dans le cadre de la gestion du document multimédia socio-personnel.

Nous nous sommes intéressés à la problématique de l'expressivité des annotations qui décrivent les documents multimédias socio-personnels vis-à-vis de chaque utilisateur. Nous nous sommes intéressés aussi à la problématique de l'exploitation sémantique du graphe social.

Divers modèles existent pour la représentation des annotations des documents multimédias socio-personnels. Certains sont basés sur des ontologies qui semblent mieux répondre aux problèmes d'expressivité et d'interopérabilité comparées aux approches d'annotations classiques (telles que l'usage des mots-clés). De plus, les standards et les modèles existant dans la littérature sont principalement destinés à la caractérisation du contenu et ils sont exploités pour une annotation extensive et manuelle des documents. Par ailleurs, mis à part les travaux de Viana et al. [44] qui proposent une ontologie (*ContextMultimedia*), les travaux de modélisation des annotations n'offrent quasiment pas de description dédiée au contexte de création des documents multimédias dont la prise en compte constitue un point de départ très prometteur pour les annotations. Enfin, aucun des modèles et standards proposés ne considère l'utilisateur comme élément central du modèle d'annotation. Ceci est peut être dû à la nature non spécifique des standards existants. En effet, ils ont été créés pour annoter des documents multimédias en général et pas spécifiquement des documents multimédias socio-personnels.

En ce qui concerne les systèmes d'annotation proposés qui prennent en considération la capture du contexte de prise de vue, tels que MMM Image Gallery [85], ZoneTag [84] et PhotoMap [97], ils enrichissent les annotations quasiment de la même manière : Ils utilisent des ressources sémantiques non spécifiques à l'utilisateur pour construire des annotations consistantes et dont l'utilité a été prouvée par les études empiriques [5], [6], [1]. De plus, en raison de l'intégration des capteurs de contexte dans la plupart des dispositifs mobiles, ce type d'approche demeure plus réaliste.

Néanmoins, cet enrichissement n'atteint pas le degré d'expressivité souhaité par les utilisateurs. En effet, des études empiriques [9] ont prouvé que l'utilisateur préfère organiser, annoter et exploiter (rechercher) les photos sociales en utilisant des axes sémantiques personnalisés. Pour atteindre un tel degré d'expressivité, il est enrichissant

d'utiliser d'autres sources de données propres à l'utilisateur. Dans ce contexte, les travaux de Sarin et al. [98] sont les plus remarquables ; ils proposent des techniques évoluées pour enrichir les annotations partant d'un client Google desktops implanté dans la machine de l'annotateur. Cependant, cette source de données est insatisfaisante. En effet, Sarin et al. négligent la nature sociale des documents multimédias socio-personnels. En effet, les documents de chaque utilisateur correspondent à des évènements qui ocurrent dans son monde réel avec des personnes appartenant à son réseau social. ; ceci explique qu'un même évènement vécu dans un document multimédia socio-personnel peut être partagé par plusieurs personnes rassemblées dans un même réseau social.

Donc, les sources d'enrichissement de contexte ne doivent pas se restreindre à des services web impersonnels (de types non spécifiques à l'utilisateur) pour enrichir les annotations par des informations comme le climat, le jour du mois. Ces sources doivent être spécifiques à l'annotateur et/ou à son réseau social. En effet, l'utilisation d'agenda ou de profil d'utilisateur sur le web, importé à partir de systèmes de réseaux sociaux ou de services de médias sociaux s'avère très bénéfique.

Compte tenu de ces différentes limites, la mise en place d'un nouveau modèle de représentation d'annotations s'avère cruciale. Ce modèle doit prendre en considération le contexte de prise de vue comme point de départ pour construire des annotations sensibles au contexte social. Ce modèle doit considérer l'utilisateur comme élément central. De plus, la mise en place de nouvelles techniques de recommandation d'annotations est essentielle. Ces nouvelles techniques doivent être sensibles au contexte social et doivent réutiliser des sources de données propres à l'utilisateur et aux membres de son réseau de connaissances.

En ce qui concerne l'exploitation, la nature du modèle de données construit (graphe) nous a été une motivation pour résoudre le problème de l'exploitation des métadonnées avec des techniques qui ont été initialement proposées dans la théorie des graphes.

9.2. Bilan des contributions

Nous avons, donc, tenté de résoudre ces problèmes en apportant les solutions suivantes:

I. Le modèle *SeMAT* de représentation d'informations contextuelles

Nous avons proposé un modèle d'annotation avec une nouvelle vision du contexte de prise de vue. La vision proposée couvre un spectre plus important d'éléments de contexte. Ce sont ceux qui peuvent être utilisées pour faciliter la gestion du document multimédia

socio-personnel. De plus, le modèle *SeMAT* proposé considère trois différentes couches du contexte. Une première est la couche physique, qui, permet de créer les autres couches partant de la capture des éléments de contexte. Une deuxième est la couche enrichie, elle contient des éléments consistants pour tous les utilisateurs. La dernière couche est la couche personnalisée, qui contient des éléments personnalisés selon l'utilisateur.

II. Le modèle *SocialSphere* de représentation de la notation du profil social

Pour instrumentaliser des moyens de personnalisations des annotations qui ne sont autres que les éléments de la couche personnalisée du contexte de prise de vue, nous avons modélisé la notion de *profil social* avec des outils du web sémantique en se focalisant, plus particulièrement, sur la notion de liens sociaux et au mécanisme de raisonnement permettant d'inférer de nouveaux liens sociaux non explicités.

Le fruit de cette modélisation est l'ontologie *SocialSphere*.

III. La plateforme d'enrichissement des métadonnées

Comme les approches existantes, nous avons utilisé des ressources externes (GeoNames⁶² et Wikipédia⁶³) pour enrichir les annotations partant des éléments de contexte capturés lors de la prise de vue. L'utilisation de ces services web génériques, nous a permis de peupler, de manière automatique, les éléments de la couche enrichie du contexte de prise de vue et construire, donc, des annotations consistantes pour tous les utilisateurs.

Le peuplement des éléments de la couche personnalisée est réalisé de manière semi-automatique.

IV. Des algorithmes de recommandation des tags sensibles au contexte social

Le module CCP de la plateforme d'enrichissement est mis en place afin d'implémenter nos algorithmes de recommandation de tags sensibles au contexte social. Nous avons proposé d'utiliser le profil social, modélisé au sein de l'ontologie *SocialSphere*, dans le but de suggérer à l'utilisateur une liste d'annotations sémantiques. Nous avons considéré des objets utilisateurs (maison, travail, etc.), des dimensions sociales (ma femme, le cousin de mon mari, etc.) ou des événements (mariage, anniversaire, etc.) pour enrichir respectivement les "*tags localisations*", "*tags identité*" et "*tags événements*".

⁶² www.geonames.org/

⁶³ fr.wikipedia.org/

Nous avons pour cet effet, proposé *SQO*, un algorithme permettant une recherche optimisée des dimensions sociales. Nous avons prouvé l'efficacité d'un tel algorithme en termes de temps d'exécution. Ensuite, nous avons réutilisé l'expérience de l'utilisateur et des membres de son réseau de connaissance en appliquant un algorithme de fouille de données sur l'ensemble de sa base de documents multimédias annotés. La réutilisation de l'expérience de l'utilisateur était dans l'objectif de lui proposer de nouveaux "*tags évènements*". L'algorithme *RSP* a été proposé pour cet effet.

V. Des mécanismes d'exploitation du graphe social basé sur l'isomorphisme de sous-graphe partiel.

Nous avons abordé le problème de correspondance (ou d'appariement) entre requête et graphe social formé par l'ensemble de photos, d'annotations, d'utilisateurs et de liens sociaux. Nous avons proposé de ramener le problème de recherche de correspondance à un problème d'isomorphisme de sous-graphe partiel. Nous avons, alors, développé un algorithme, appelé *h-Pruning*, d'isomorphisme de sous-graphe partiel permettant de faire une correspondance rapprochée entre les nœuds des deux graphes : motif (représentant la requête) et social. La comparaison entre les nœuds des deux graphes prend en considération la similarité de leurs attributs. Cette similarité peut être calculée en spécifiant un ou plusieurs mesures de similarité sémantique.

9.3. Perspectives

Plusieurs perspectives scientifiques sont envisagées comme suites de notre travail. Elles peuvent être classées en perspectives réalisables à court terme et perspectives réalisables à plus long terme. Nous soulignons celles d'entre elles qui nous paraissent pertinentes pour l'évaluation des propositions réalisées.

À court terme :

- *Approfondir les mécanismes de suggestion d'objets des utilisateurs*: nous avons proposé comme moyen de personnalisation des annotations selon les consultants trois directives : une première concerne les "*tags identité*", qui consiste à ajouter des dimensions sociales. Une deuxième concerne les "*tags événements*" qui consiste à se servir de l'expérience passée de l'utilisateur et de son réseau de connaissance afin de lui proposer des "*tags événements*" pertinents. Une dernière concerne les "*tags localisation*" qui propose des objets d'utilisateurs pour caractériser sémantiquement le lieu de prise de vue du document multimédia socio-personnel selon le consultant. Pour cette dernière approche nous avons considéré les objets d'utilisateur géo-localisés présents dans le profil social de l'observateur et des participants à la prise de vue. Nous avons proposé de calculer la distance euclidienne [108] entre les coordonnées géographiques de ces objets et celles du document multimédia socio-personnel en question. Par définition, la distance euclidienne est la distance à "*vol d'oiseau*" que l'on mesure sur la carte ; elle correspond rarement à la distance que l'on parcourt réellement sur le terrain. Dans un futur proche, nous envisageons d'approfondir cette approche investiguant une indexation à base de quadtree [133]. Cette dernière s'avère très rapide au niveau de la recherche des coordonnées géographiques les plus proches.
- *Evaluations des approches de suggestions de tags avec de vrais utilisateurs* : à l'état actuel, nous avons mené des expérimentations sur un corpus butiné depuis le service de médias sociaux Flickr. Le nombre de photos dans ce corpus (~2500 photos) ne représente pas un nombre considérable d'échantillons de comparaison. De plus, ce corpus encapsule des données d'utilisateurs ayant des profils hétérogènes choisis arbitrairement. Il serait, intéressant voire nécessaire d'étudier la cooccurrence et la récurrence du contexte spatio-temporel et des événements sur un échantillon représentatif de la société. Par exemple, considérer un échantillon

comportant des adultes, des adolescents, des hommes et des femmes. Une autre idée serait d'étudier la variation de ces motifs en considérant l'âge des individus. Nous envisageons, dans un futur proche d'utiliser le prototype développé (*Phoox*) pour valider nos approches pour ce genre d'étude sociale.

À long terme:

- *Ajouter plus d'expressivité au concept de graphe motif*: en ce qui concerne l'approche de correspondance entre la requête et le graphe social. A l'état actuel, seul la correspondance simple est permise. Nous souhaitons l'améliorer en ajoutant plus d'expressivité à la requête comme la jointure, l'agrégation, les opérateurs ensemblistes, etc.
- *Ajouter une interface conviviale à l'interrogation*: sur le plan de l'interface homme-machine, l'interrogation textuelle n'est pas la meilleure solution en termes de simplicité pour l'utilisateur. Nous souhaitons l'améliorer en proposant une approche de requête graphique à l'aide d'un visualiseur de graphes et de fonctionnalités de glisser-déplacer.
- *Vie privée de l'utilisateur*: les utilisateurs nous ont souvent reproché le fait que le prototype développé dans notre travail ne permet pas de conserver la vie privée des utilisateurs. Compte tenu de ces exigences primordiales, nous planifions d'apporter des solutions à ce genre de problèmes en intégrant par exemple l'approche proposée par notre collègue Omar Hasan dans son travail de thèse [134].

GLOSSAIRE DES ACRONYMES

GLOSSAIRE DES ACRONYMES

AAT	Art and Architecture Thesaurus
AKT	Advanced Knowledge Technologies
API	Application Program Interface
BFS	Breadth First Search
CCE	Contexte de Création Enrichit
CCP	Contexte de Création Personnalisé
Cell-ID	Identificateur Cellulaire
ConceptNet	Base de connaissance accessible librement
CONFOTO	Un système doté de services d’annotation et de navigation sémantique pour les photos de conférences
CSV	Comma Separated Values
DAML+OIL	Un langage ontologique pour le Web sémantique. Il est le successeur du langage DAML et OIL en combinant les caractéristiques des deux langages
DCMES	Dublin Core Metadata Element Set
DDL	Definition Descriptions Language
DS	Description Schemes
Dublin Core	Dublin Core Metadata Initiative est une organisation dédiée à la création et la diffusion de vocabulaires destinés à la description des métadonnées sur des ressources numériques y compris les photos
DyNetML	In format basé sur XML pour stocker un réseau social enrichi
EXIF	EXchangeable Image File
FOAF	Friend Of Friend
FOAF-X	Une extension de FOAF
GeoNames	base de données géographique couvrant plusieurs pays et contient plus de huit millions de lieux qui sont disponibles en téléchargement gratuit.
GIF	Graphics Interchange Format

GPS	Global Positioning System
GSM	Global System for Mobile Communications
HTML	Hypertext Markup Language
Icon-class	Icon-class est un système de classification conçu pour l'art et l'iconographie
IIM	Information Interchange Model
IPTC	IPTC Photo Metadata Standards
Jess	Rule Engine for the Java Platform
JPEG	Joint Photographic Experts Group
MPEG-7	Moving Picture Experts Group
OIL	Ontology Inference Layer
OWL	Web Ontology Language
OWL DL	Ontology Web Language Description Logics
OWL Full	Ontology Web Language Full
OWL Lite	Web Ontology Language Lite
PASMI	self-adaptive Photo Annotation and Sharing Middleware
Pellet	Raisonneur sur les ontologies écrites en OWL DL
PNG	Portable Network Graphics : un format d'images
PSD	Photoshop Document : un format natif du logiciel Adobe Photoshop
Racer	Raisonneur ontologiques pour des ontologies écrites en OWL DL
RDF	Resource Description Framework
RDFS	Schema for RDF
RSP	Recommendation via Social Proximity
SeMAT	Semantic Model for self Adaptive Tag
SMIL	Synchronized Multimedia Integration Language
SNC	Social Network Category
SocialSphere	Ontologie décrivant des niveaux de granularités des liens sociaux

SPARQL	SPARQL un langage de requête pour des bases en RDF
SQL	Structured Query Language. Language de requête pour des bases de données relationnelles
SQO	Social Query Optimisation
SQO	Social Query Optimisation
SWRL	Semantic Web Rule Language
TIFF	Tagged Image File Format
Time	Ontologie décrivant les concepts temporels exprimés en OWL
UML	Unified Modeling Language
URI	Uniform Resource Identifier
W3C	World Wide Web Consortium
WordNet	Une base de données lexicale développée par des linguistes du laboratoire des sciences cognitives de l'université de Princeton.
WWW	World Wide Web
XFN	Xhtml Friends Network
XHTML	eXtensible HyperText Markup Language
XML	eXtensible Markup Language
XML schema	XML schema
XMP	Extensible Metadata Platform

BIBLIOGRAPHIE

BIBLIOGRAPHIE

- [1] M. Naaman, "Leveraging geo-referenced digital photographs," Thèse de doctorat, Stanford University, Stanford, 2005.
- [2] J. Bretin, "Statistiques Facebook - Février 2010 [en ligne]," *Disponible sur* <<http://www.seomanager.fr/statistiques-facebook-fevrier-2010.html>> (consulté le 2010-06-17 12:17:00), 2010.
- [3] Y. Chai, X. Zhu, et J. Jia, "OntoAlbum : An Ontology Based Digital Photo Management System," *Image Analysis and Recognition*, 2008, p. 263-270.
- [4] J. Kustanowitz et B. Shneiderman, *Motivating Annotation for Personal Digital Photo Libraries: Lowering Barriers While Raising Incentives*, Maryland, USA: Univ. of Maryland, 10 pages, 2005.
- [5] A. Matellanes, A. Evans, et B. Erdal, "Creating an application for automatic annotations of images and video," *In Proc. of the 1er International Workshop on Semantic Web Annotations for Multimedia (SWAMM)*, Edinburgh, Scotland: 2006, p. 10 pages.
- [6] M. Naaman, Y.J. Song, A. Paepcke, et H. Garcia-Molina, "Automatic Organization for Digital Photographs with Geographic Coordinates," *In Proc. of the 4th ACM/IEEE-CS joint conference on Digital libraries (JCDL'04)*, Tucson, Arizona: ACM Press, 2004, p. 53-62.
- [7] S. Favre, "Social Network Analysis in Multimedia Indexing: Making Sense of People in Multiparty Recordings," *In Proc. of the Doctoral Consortium of the International Conference on Affective Computing & Intelligent Interaction (ACII)*, Amsterdam: 2009, p. 8 pages.
- [8] Z. Kunda, *Social Cognition: Making Sense of People*, The MIT Press, 1999.
- [9] K. Rodden et K.R. Wood, "How do people manage their digital photographs?," *In Proc. of the SIGCHI conference on Human factors in computing systems*, Ft. Lauderdale, Florida, USA: ACM, 2003, p. 409-416.
- [10] R. W3C, "Synchronized Multimedia Integration Language (SMIL 2.1)," *Disponible sur* <<http://www.w3.org/TR/SMIL2/>> (consulté le 2010-11-22 08:06:49), 2005.
- [11] Macromedia, "Macromedia Flash FLA Project File Format [en ligne]," *Disponible sur* <<http://www.digitalpreservation.gov/formats/fdd/fdd000132.shtml>> (consulté le 2010-11-29 12:37:06), 2007.

-
- [12] Wikipédia, “HTML5 — Wikipédia, l'encyclopédie libre [en ligne],” *Disponible sur* <http://fr.wikipedia.org/w/index.php?title=HTML5&oldid=59678318l> (consulté le 2010-11-29), 2010.
- [13] J. Lehtikainen, A. Aaltonen, P. Huuskonen, et I. Salminen, *Personal Content Experience: Managing Digital Life in the Mobile Age*, Wiley-Interscience, 382 pages, 2007.
- [14] W. Viana, “Mobilité et sensibilité au contexte pour la gestion de documents multimédias personnels : CoMMedia,” Thèse de doctorat, Université de Grenoble, 266 pages, 2010.
- [15] A. Mitschick et K. Meißner, “Generation and Maintenance of Semantic Metadata for Personal Multimedia Document Management,” *In Proc. of the International Conference on Advances in Multimedia*, Los Alamitos, CA, USA: IEEE Computer Society, 2009, p. 74-79.
- [16] M. Crampes, J. de Oliveira-Kumar, S. Ranwez, et J. Villerd, “Indexation de photos sociales par propagation sur une hiérarchie de concepts,” *In the Proc. of Conférence Francophone en Ingénierie des Connaissances, IC*, Tunisie: 2009, p. 12 pages.
- [17] E. Wiki, “Tagging — EduTech Wiki [en ligne],” *Disponible sur* <http://edutechwiki.unige.ch/fmediawiki/index.php?title=Tagging&oldid=16060> (consulté le 2010-11-29 12:37:06), 2010.
- [18] Wikipédia, “Folksonomie — Wikipédia, l'encyclopédie libre [en ligne],” *Disponible sur* <http://fr.wikipedia.org/w/index.php?title=Folksonomie&oldid=58549839> (consulté le 2010-11-29), 2010.
- [19] H. Zargayouna, “Indexation sémantique de documents XML,” Thèse de Doctorat en Informatique, Université Paris XI UFR scientifique d’Orsay, 227 pages, 2005.
- [20] E. Desmontils et C. Jacquin, “Annotation sur le web :notes de lectures,” *Disponible sur* <http://enssibal.enssib.fr/autres-sites/RTP/websemantique/octobre/octobre1/Desmontils.pdf> (consulté le 25-11-2010): 2002, p. 5 pages.
- [21] T. Berners-Lee et M. Fischetti, *Weaving the Web: The Original Design and Ultimate Destiny of the World Wide Web by Its Inventor*, Harper San Francisco, 226 pages, 1999.
- [22] T. Berners-Lee, J. Hendler, et O. Lassila, “The Semantic Web: Scientific American,” *Scientific American*, Mai. 2001.

-
- [23] R. W3C, "Extensible Markup Language (XML) 1.0 (Fifth Edition) [en ligne]," *Disponible sur* <<http://www.w3.org/TR/REC-xml/>> (consulté le 2010-06-12), 2008.
- [24] R. W3C, "XML Schema Part 1: Structures Second Edition [en ligne]," *Disponible sur* <<http://www.w3.org/TR/xmlschema-1/>> (consulté le 2010-06-12), 2004.
- [25] H. Yuh-Jong, "RuleML [en ligne]," *Disponible sur* <<http://ruleml.org/>> (consulté le 2010-11-01), 2010.
- [26] W3C Note, "SWRL: A Semantic Web Rule Language Combining OWL and RuleML [en ligne]," *Disponible sur* <<http://www.w3.org/Submission/SWRL/>> (consulté le 2010-11-01), 2004.
- [27] HP Labs, "Jena Semantic Web Framework [en ligne]," *Disponible sur* <<http://jena.sourceforge.net>> (consulté le 2010-11-01).
- [28] R. W3C, "RDF - Semantic Web Standards," *Disponible sur* <<http://www.w3.org/RDF/>> (consulté le 2010-06-12), 2004.
- [29] T.R. Gruber, "A translation approach to portable ontology specifications," *Knowledge Acquisition*, 1993, p. 199-220.
- [30] R. W3C, "RDF Vocabulary Description Language 1.0: RDF Schema [en ligne]," *Disponible sur* <<http://www.w3.org/TR/rdf-schema/>> (consulté le 2010-11-01), 2004.
- [31] R. W3C, "SPARQL Query Language for RDF [en ligne]," *Disponible sur* <<http://www.w3.org/TR/rdf-sparql-query/>> (consulté le 2010-11-02), 2008.
- [32] J. Pérez, M. Arenas, et C. Gutierrez, "Semantics and complexity of SPARQL," *In ACM Trans. Database Syst.*, vol. 34, 2009, p. 1-45.
- [33] DCMI, "Dublin Core Metadata Initiative (DCMI) [en ligne]," *Disponible sur* <<http://dublincore.org/>> (consulté le 2010-06-12), 2010.
- [34] IPTC, "international Press Telecommunications Council [en ligne]," *Disponible sur* <<http://www.iptc.org/cms/site/index.html?channel=CH0100>> (consulté le 2010-06-12), 2010.
- [35] JEITA, "EXIF 2.2 Specifications [en ligne]," *Disponible sur* <<http://www.exif.org/specifications.html>> (consulté le 2010-06-12), 2002.
- [36] H. Kosch, *Distributed Multimedia Database Technologies Supported Mpeg-7 and by Mpeg-21*, CRC Press Inc, 2003.
- [37] M. Lux, "Caliph & Emir: MPEG-7 photo annotation and retrieval," *In Proc. of the*

-
- seventeen *ACM international conference on Multimedia*, Beijing, China: ACM, 2009, p. 925-926.
- [38] M. Lux, J. Becker, et H. Krottmaier, "Caliph & Emir: Semantic Annotation and Retrieval," *IN Proc. of CAISE '03 forum PERSONAL DIGITAL PHOTO LIBRARIES at 15TH conference on advanced information systems engineering*, Klagenfurt/Velden, Austria: Springer-Verlag, 2003, p. 85-89.
- [39] G. Martens,, C. Poppe, et E. Mannens, "DIG35: Metadata Standard for Digital Images [en ligne]," *Disponible sur* <<http://multimedialab.elis.ugent.be/users/gmartens/Ontologies/DIG35/v0.2/>> (consulté le 2010-06-12), 2007.
- [40] W3C Note, "Describing and retrieving photos using RDF and HTTP [en ligne]," *Disponible sur* <<http://www.w3.org/TR/photo-rdf/>> (consulté le 2010-06-19), 2002.
- [41] Visual Resources Association, "VRA Core 4.0 Specification," 2007.
- [42] L. Hollink, A.T. Schreiber, J. Wielemaker, et B. Wierlinga, "Semantic Annotation of Image Collections," Florida: 2003, p. 8.
- [43] Adobe Systems, "Adobe XMP: Extensible Metadata Platform specification," 2001.
- [44] W. Viana, "Mobilité et sensibilité au contexte pour la gestion de documents multimédias personnels : CoMMedia," Thèse de doctorat, Université de Grenoble, 266 pages, 2010.
- [45] W.D. Nooy, A. Mrvar, et V. Batagelj, *Exploratory Social Network Analysis with Pajek*, Cambridge University Press, 2005.
- [46] M. Tsvetovat, J. Reminga, et K.M. Carley, "DyNetML: Interchange Format for Rich Social Network Data [en ligne]," *Disponible sur* <http://www.casos.cs.cmu.edu/publications/papers/tsvetovat_2003_dynetmlinterchange.pdf> (consulté le 2010-07-11), 2004.
- [47] S. Amer-yahia, L.V.S. Lakshmanan, et C. Yu, "SocialScope: Enabling Information Discovery on Social Content Sites," *In Proc. of the 4th biennial conference on innovative data systems CIDR 2009*, Berkeley, CA, USA: 2009, p. 11 pages.
- [48] D. Brickley et L. Miller, *The Friend Of A Friend (FOAF) Vocabulary Specification*, 2007.
- [49] P. Mika, "Social Networks and the Semantic Web," *Social Networks and the Semantic Web*, Springer, 2007, p. 285-291.

-
- [50] R. Rada, H. Mili, E. Bicknell, et M. Blettner, "Development and application of a metric on semantic nets," *In IEEE Transactions on Systems, Man, and Cybernetics*, vol. 19, 1989, p. 17-30.
- [51] G. Hirst et D. St-Onge, "Lexical Chains as Representations of Context for the Detection and Correction of Malapropisms," *In WordNet: An electronic lexical database, Cambridge, MA: The MIT Press*, 1998, p. 107-116.
- [52] Z. Wu et M. Palmer, "Verbs semantics and lexical selection," *In Proc. of the 32nd annual meeting on Association for Computational Linguistics*, Las Cruces, New Mexico: Association for Computational Linguistics, 1994, p. 133-138.
- [53] D. Lin, "An Information-Theoretic Definition of Similarity," *IN Proc. of the 15th international conference on machine learning*, Madison, Wisconsin USA: Morgan Kaufmann Publishers, 1998, p. 296-304.
- [54] P. Resnik, "Using Information Content to Evaluate Semantic Similarity in a Taxonomy," *IN Proc. of the 14th International Joint conference on artificial intelligence (IJCAI-95)*, Montreal, Quebec, Canada: 1995, p. 448-453.
- [55] J.J. Jiang et D.W. Conrath, "Semantic Similarity Based on Corpus Statistics and Lexical Taxonomy," *Proc. of In the International Conference Research on Computational Linguistics*, Taiwan: 1997, p. 19-33.
- [56] E. Agirre et G. Rigau, "Word sense disambiguation using Conceptual Density," *In Proc. of the 16th conference on Computational linguistics - Vol. 1*, Copenhagen, Denmark: Association for Computational Linguistics, 1996, p. 16-22.
- [57] A. Smeulders, M. Worring, S. Santini, A. Gupta, et R. Jain, "Content-based image retrieval at the end of the early years," *In Transactions on Pattern Analysis and Machine Intelligence*, vol. 22, 2000, p. 1349-1380.
- [58] C. Fellbaum, "WordNet: An Electronic Lexical Database (Language, Speech, and Communication)," The MIT Press, 1998, p. 423 pages.
- [59] T. Peterson, "Introduction to the Art and Architecture Thesaurus [en ligne]," *Disponible sur* <<http://www.getty.edu/research/tools/vocabulary/aat/>> (consulté le 21-10-2010), Oxford University Press: 1994.
- [60] B. Elliott, "Hierarchical and semantic data management and querying for patient records and personal photos," Thèse de doctorat, Case Western Reserve University, 384 pages, USA, 2009.
- [61] S. Deerwester, S.T. Dumais, G.W. Furnas, T.K. Landauer, et R. Harshman,

-
- “Indexing by Latent Semantic Analysis,” *In Journal of the american society for information science*, vol. 41, 1990, p. 391-407.
- [62] J. Lim, P. Mulhem, et Q. Tian, “Home Photo Content Modeling for Personalized Event-Based Retrieval,” *In IEEE Multimedia Special Issue on Image Retrieval*, vol. 10, 2003, p. 28-37.
- [63] The Getty Foundation, “ULAN: Union List of Artist Names [en ligne],” *Disponible sur* < <http://www.getty.edu/research/tools/vocabulary/ulan/>> (consulté le 2010-06-12), 2000.
- [64] H. van der Waal, *ICONCLASS: An iconographic classification system*, Royal Dutch Academy of Sciences (KNAW), 1985.
- [65] R. W3C, “OWL Web Ontology Language Overview [en ligne],” *Disponible sur* < <http://www.w3.org/TR/owl-features/>> (consulté le 2010-11-01), 2004.
- [66] B. Shevade, H. Sundaram, et N.M. Yen-Kan, “A Collaborative Annotation Framework,” *In Proc. of the IEEE International Conference on Multimedia and Expo*, Los Alamitos, CA, USA: IEEE Computer Society, 2005, p. 1346-1349.
- [67] H. Liu et P. Singh, “ConceptNet A Practical Commonsense Reasoning Tool-Kit,” *In BT Technology Journal*, vol. 22, 2004, p. 211-226.
- [68] L.V. Ahn et L. Dabbish, “Labeling images with a computer game,” *In Proc. of the SIGCHI conference on Human factors in computing systems*, Vienna, Austria: ACM, 2004, p. 319-326.
- [69] google, “Google Image Labeler [en ligne],” *Disponible sur* <<http://images.google.com/imagelabeler/>> (consulté le 2010-06-16), 2007.
- [70] M. Truran, J. Goulding, et H. Ashman, “Co-active intelligence for image retrieval,” *In Proc. of the 13th annual ACM international conference on Multimedia*, Hilton, Singapore: ACM, 2005, p. 547-550.
- [71] B. Shevade, H. Sundaram, et L. Xie, “Modeling personal and social network context for event annotation in images,” *In Proc. of the 7th ACM/IEEE-CS joint conference on Digital libraries*, Vancouver, BC, Canada: ACM, 2007, p. 127-134.
- [72] A. Zunjarwad, H. Sundaram, et L. Xie, “Contextual wisdom: social relations and correlations for multimedia event annotation,” *In Proc. of the 15th international conference on Multimedia*, Augsburg, Germany: ACM, 2007, p. 615-624.
- [73] Z. Stone, T. Zickler, et T. Darrell, “Autotagging facebook: Social network context

-
- improves photo annotation,” *In Proc. of the 1st IEEE Workshop on Internet Vision (CVPR 2008)*, 2008, p. 8 pages.
- [74] J.R. Anderson, “A spreading activation theory of memory,” *In Journal of Verbal Learning and Verbal Behavior*, vol. 22, 1983, p. 261-295.
- [75] L. Wenyin, S. Dumais, Y. Sun, H. Zhang, M. Czerwinski, et B. Field, “Semi-Automatic Image Annotation,” *In Proc. of INTERACT 2001: 8th TC13 IFIP International Conference on Human-Computer Interaction*, Tokyo, Japan: 2001, p. 326-333.
- [76] B. Sigurbjörnsson et R.V. Zwol, “Flickr tag recommendation based on collective knowledge,” *In Proc. of the 17th international conference on World Wide Web*, Beijing, China: ACM, 2008, p. 327-336.
- [77] O. Kucuktunc, S. Sevil, A. Tosun, H. Zitouni, P. Duygulu, et F. Can, “Tag Suggestr: Automatic Photo Tag Expansion Using Visual Information for Photo Sharing Websites,” *In Semantic Multimedia*, 2008, p. 61-73.
- [78] I. Ivanov, P. Vajda, L. Goldmann, J. Lee, et T. Ebrahimi, “Object-based tag propagation for semi-automatic annotation of images,” *In Proc. of the international conference on Multimedia information retrieval*, Philadelphia, Pennsylvania, USA: ACM, 2010, p. 497-506.
- [79] Wikipédia, “Bluetooth — Wikipédia, l'encyclopédie libre [en ligne],” *Disponible sur* <http://fr.wikipedia.org/w/index.php?title=Bluetooth&oldid=57191968> (consulté le 2010-11-29), 2010.
- [80] Wikipédia, “Global Positioning System — Wikipédia, l'encyclopédie libre [en ligne],” *Disponible sur* http://fr.wikipedia.org/w/index.php?title=Global_Positioning_System&oldid=58345154 (consulté le 2010-11-29), 2010.
- [81] Wikipedia, “Cell Global Identity — Wikipedia, The Free Encyclopedia [en ligne],” *Disponible sur* http://en.wikipedia.org/w/index.php?title=Cell_Global_Identity&oldid=354241758 (consulté le 2010-11-29), 2010.
- [82] B. Nowack, “CONFOTO: A Semantic Browsing and Annotation Service for Conference Photos,” *In Proc. of the 4th international Semantic Web conference (ISWC 2005)*, LNCS, Galway, Ireland: SpringerLink, 2005, p. 1067-1070.
- [83] M. Naaman, R.B. Yeh, H. Garcia-Molina, et A. Paepcke, “Leveraging context to
-

-
- resolve identity in photo albums,” *In Proc. of the 5th ACM/IEEE-CS joint conference on Digital libraries*, Denver, CO, USA: ACM, 2005, p. 178-187.
- [84] S. Ahern, S. King, M. Naaman, R. Nair, et J. Yang, “ZoneTag: Rich, Community-supported Context-Aware Media Capture and Annotation,” *In Proc. of Mobile Spatial Interaction workshop (MSI) at the SIGCHI conference on Human Factors in computing systems (CHI 2007)*, San Jose, California, USA: ACM, 2007.
- [85] R. Sarvas, E. Herrarte, A. Wilhelm, et M. Davis, “Metadata creation system for mobile images,” *In Proc. of the 2nd international conference on Mobile systems, applications, and services*, Boston, MA, USA: ACM, 2004, p. 36-48.
- [86] M.M. Tuffield, S. Harris, C. Brewster, N. Gibbins, F. Ciravegna, D. Sleeman, N.R. Shadbolt, et Y. Wilks, “Image annotation with photocopain,” *IN Proc. of semantic web annotation of multimedia (SWAMM-06) Workshop at the World Wide Web Conference (WWW 06)*, Southampton, United Kingdom: 2006, p. 22-26.
- [87] M.M. Tuffield, D.P. Dupplaw, K. O'Hara, N.R. Shadbolt, A. Chakravarthy, C. Brewster, F. Ciravegna, et Y. Wilks, “Photocopain - Annotating Memories For Life,” *In Proc. of Memories for life colloquium*, London, United Kingdom: 2006, p. 2 pages.
- [88] A. Chakravarthy, “AKTiveMedia: Cross-media Document Annotation and Enrichment,” *CEUR-WS*, Disponible sur: < <http://eprints.aktors.org/537/01/poster-camera.pdf>> (consulté le 21-06-2010): 2006.
- [89] A. Chakravarthy, “Cross-Media document annotation and enrichment,” *In Proc. of the 1st semantic Web authoring and annotation workshop (SAAW2006) co-located with the The 5th International Semantic Web Conference (ISWC2006)*, Athens, GA, USA, Monday: 2006, p. 3 pages.
- [90] N. O'Hare, C. Gurrin, G.J.F. Jones, H. Lee, N.E. O'Connor, et A.F. Smeaton, “Using text search for personal photo collections with the MediAssist system,” *In Proc. of the 2007 ACM symposium on Applied computing*, Seoul, Korea: ACM, 2007, p. 880-881.
- [91] N. O'Hare, H. Lee, S. Cooray, C. Gurrin, G. Jones, J. Malobabic, N. O'Connor, A. Smeaton, et B. Uscilowski, “MediAssist: Using Content-Based Analysis and Context to Manage Personal Photo Collections,” *In Proc. of the Image and Video Retrieval, LNCS, Vol. 4071*, Springer Berlin / Heidelberg, 2006, p. 529-532.
- [92] F. Monaghan, “Context-aware photograph annotation on the Social Semantic Web,”

-
- Thèse de doctorat, Digital Enterprise Research Institute, National University of Ireland, Galway, 206 pages, 2008.
- [93] F. Monaghan et D. O'Sullivan, "Leveraging Ontologies, Context and Social Networks to Automate Photo Annotation," *In Semantic Multimedia*, 2007, p. 252-255.
 - [94] J. Rodgers et A. Nicewander, "Thirteen Ways to Look at the Correlation Coefficient," *In The American Statistician*, vol. 42, 1988, p. 59-66.
 - [95] Bryant et Yarnold, "Principal components analysis and exploratory and confirmatory factor analysis," *Reading and understanding multivariate analysis*, 1994, p. 373 pages.
 - [96] G. Bradski et A. Kaehler, "Learning OpenCV: Computer Vision with the OpenCV Library," O'Reilly Media, 555 pages, 2008.
 - [97] W. Viana, J.B. Filho, J. Gensel, M. Villanova-Oliver, et H. Martin, "PhotoMap: from location and time to context-aware photo annotations," *J. Locat. Based Serv.*, vol. 2, 2008, p. 211-235.
 - [98] S. Sarin, T. Nagahashi, T. Miyosawa, et W. Kameyama, "On the design and exploitation of user's personal and public information for semantic personal digital photograph annotation," *In Adv. MultiMedia Vol. 2*, vol. 2008, 2008, p. 1-16.
 - [99] M. Naaman, "Leveraging geo-referenced digital photographs," Thèse de doctorat, Stanford University, 202 pages, 2005.
 - [100] G. Deleuze, "Qu'est ce qu'un événement ?," *Disponible sur* <<http://cerphi.net/lec/even.htm>> (consulté le 2010-06-12), 1984.
 - [101] A.K. Dey, "Understanding and Using Context," *In Personal Ubiquitous Comput.*, vol. 5, 2001, p. 4-7.
 - [102] W. Viana, J.B. Filho, J. Gensel, M. Villanova-Oliver, et H. Martin, "PhotoMap: from location and time to context-aware photo annotations," *In Journal of Location Based Service*, vol. 2, 2008, p. 211-235.
 - [103] GeoNames, "GeoNames [en ligne]," *Disponible sur* <<http://www.geonames.org/>> (consulté le 2010-06-12), 2008.
 - [104] Q. Zhou et Q.R. Fikes, "A Reusable Time Ontology," *In AAAI Workshop on Ontologies for the Semantic Web*, 2002, p. 9 pages.
 - [105] M. Naaman, Y. Song, A. Paepcke, et H. Garcia-Molina, "Automatic Organization
-

-
- for Digital Photographs with Geographic Coordinates,” *In Proc. of the 4th ACM/IEEE-CS joint conference on Digital libraries (JCDL'04)*, Tucson, Arizona: ACM Press, 2004, p. 53-62.
- [106] R. Jerry et P. Feng, “An Ontology of Time for the Semantic Web,” *In ACM Transactions on Asian Language Processing (TALIP): Special issue on Temporal Information Processing, Vol 3 No. 1*, 2004, p. 66-85.
- [107] R. Jerry et P. Feng, “Ontology Engineering Patterns Task Force of the Semantic Web Best Practices and Deployment Working Group, World Wide Web Consortium (W3C) notes,” *Disponible sur* <http://www.w3.org/TR/owl-time/> (consulté le 2010-10-25), 2006.
- [108] Rosset, “Distance euclidienne [en ligne],” *Disponible sur* <http://www.limsi.fr/Individu/rosset/similarite2/node13.html> (consulté le 2010-11-02), 2008.
- [109] S. Lajmi, E. Egyed-Zsigmond, A.B. Hamadou, et J. Pinon, “Enrichissement de l'annotation des identités des personnes dans les photos personnelles,” *In Prise en Compte de l'Usager dans les Systèmes d'Information (PeCUSI 2009) Atelier conjoint à INFORSID 2009*, Toulouse, France: 2009.
- [110] S. Lajmi, J. Stan, H. Hacid, E. Egyed-Zsigmond, et P. Maret, “Extended Social Tags: Identity Tags Meet Social Networks,” *In Proc. of the IEEE International Conference on Computational Science and Engineering (SocialCom09)*, Vancouver, USA: IEEE Computer Society, 2009, p. 181-187.
- [111] A. Negash, S. Lajmi, V. Scuturici, et L. Brunie, “PASMi: self-adaptive Photo Annotation and Sharing Middleware of Mobile Ad-hoc Networks,” *In Proc. of the IEEE International Conference on Pervasive Computing and Communications Workshops (PerComW 2010)*, IEEE, Éd., Mannheim, Germany: 2010.
- [112] M. Malek, *Fouille de données : Approche de règles d'association*, Cergy, France: Ecole internationale des sciences du traitement de l'information (EISTI), 6 pages, 2008.
- [113] M. Naaman, S. Harada, Q. Wang, H. Garcia-Molina, et A. Paepcke, “Context data in geo-referenced digital photo collections,” *In Proc. of the 12th annual ACM international conference on Multimedia*, New York, NY, USA: ACM, 2004, p. 196-203.
- [114] R. Agrawal et R. Srikant, *Fast Algorithms for Mining Association Rules*, IBM
-

-
- Research Report RJ9839, IBM Almaden Research Center, San Jose, California, USA: 1994.
- [115] R. Agrawal, H. Mannila, R. Srikant, H. Toivonen, et A.I. Verkamo, “Fast discovery of association rules,” *In Advances in knowledge discovery and data mining*, American Association for Artificial Intelligence, 1996, p. 307-328.
- [116] S. Perugini, M.A. Gonçalves, et E.A. Fox, “Recommender Systems Research: A Connection-Centric Survey,” *In Journal of Intelligent Information Systems*, vol. 23, 2004, p. 107-143.
- [117] M. Chen, J. Han, et P.S. Yu, “Data Mining: An Overview from a Database Perspective,” *In IEEE transactions on knowledge and data engineering*, vol. 8, 1996, p. 866-883.
- [118] W.A. Wagenaar, “My memory: A study of autobiographical memory over six years,” *Cognitive Psychology*, vol. 18, 1986, p. 225-252.
- [119] D.E. Knuth, *The art of computer programming, volume 1 (3rd ed.): fundamental algorithms*, Redwood City, CA, USA: Addison Wesley Longman Publishing Co., Inc., 1997.
- [120] O. Hasan, L. Brunie, et J. Pierson, “Evaluation of the iterative multiplication strategy for trust propagation in pervasive environments,” *In Proc. of the 2009 international conference on Pervasive services*, London, United Kingdom: ACM, 2009, p. 49-54.
- [121] W. Viana, J. Gensel, M. Villanova-Oliver, et H. Martin, “Indexation sémantique d'images géoréférencées,” *In Revue internationale de géomatique*, vol. 19, 2009, p. 169-189.
- [122] W. Viana, S. Hammiche, B. Moisuc, M. Villanova-Oliver, J. Gensel, et H. Martin, “Semantic keyword-based retrieval of photos taken with mobile devices,” *Proc. of the 6th International Conference on Advances in Mobile Computing and Multimedia*, Linz, Austria: ACM, 2008, p. 192-199.
- [123] W. Viana, S. Hammiche, M. Villanova-Oliver, J. Gensel, et H. Martin, “Photo Context as a Bag of Words,” *In Proc. of the International Symposium on Multimedia*, Los Alamitos, CA, USA: IEEE Computer Society, 2008, p. 310-315.
- [124] Y. Prié, “Représentation de documents audiovisuels en Strates-Interconnectées par les Annotations pour l'exploitation contextuelle,” Thèse de Doctorat en Informatique, Institut National des Sciences Appliquées (INSA) de Lyon, 276 pages,

-
- 1999.
- [125] E. Egyed-Zsigmond, “Modélisation des connaissances dans une base de documents multimédias,” Thèse de Doctorat en Informatique, Institut National des Sciences Appliquées (INSA) de Lyon, 216 pages, 2003.
- [126] S. Fortin, *The Graph Isomorphism Problem*, rapport technique, n: TR 96-20, Edmonton, Alberta, Canada: Dept of Computing Science, Univ. Alberta, 1996.
- [127] C. Solnon, “AllDifferent-based Filtering for Subgraph Isomorphism,” *In Artificial Intelligence*, vol. 174, 2010, p. 850–864.
- [128] E. Egyed-Zsigmond, S. Lajmi, et Z. Iszlai, “Concurrent use in an image management system,” *In Proc. of the 2006 conference on Leading the Web in Concurrent Engineering: Next Generation Concurrent Engineering*, Antibes, France: IOS Press, 2006, p. 403-417.
- [129] E. Egyed-zsigmond, S. Lajmi, et Z. Iszlai, “PhotoMot, Collaborative Image Management, Interacting with Use Traces,” *In Proc. of the 2nd International Conference on Digital Information Management, ICDIM '07*, Lyon, France: IEEE, 2007, p. 104-109.
- [130] A. Negash Shiferaw, “Mobility and Interest Aware Information Sharing in MANETs,” Thèse de Doctorat en Informatique, L’Institut National des Sciences Appliquées de Lyon (INSA de Lyon), 255 pages, France, 2010.
- [131] Wikipédia, “Modèle-Vue-Contrôleur — Wikipédia, l'encyclopédie libre [en ligne],” Disponible sur <<http://fr.wikipedia.org/w/index.php?title=Mod%C3%A8le-Vue-Contr%C3%B4leur&oldid=58193212>> (consulté le 2010-11-29), 2010.
- [132] P. Resnik, “Semantic Similarity in a Taxonomy: An Information-Based Measure and its Application to Problems of Ambiguity in Natural Language,” *In Journal of artificial intelligence research*, vol. 11, 1999, p. 95-130.
- [133] H. Samet, *Foundations of Multidimensional and Metric Data Structures*, Morgan Kaufmann, 1024 pages, 2006.
- [134] O. Hasan, “Privacy Preserving Reputation Systems for Decentralized Environments,” Thèse de Doctorat en Informatique, INSA de Lyon, 174 pages, France, 2010.

ANNEXES

Annexe I. L'Ontologie *SocialSphere*

Code de l'ontologie *SocialSphere* sérialisé en XML/RDF

```
<rdf:RDF
  xmlns:foaf="http://xmlns.com/foaf/0.1/"
  xmlns:swap="http://www.w3.org/2000/10/swap/pim/contact#"
  xmlns:wordnet="http://xmlns.com/wordnet/1.6/"
  xmlns:rdf="http://www.w3.org/1999/02/22-rdf-syntax-ns#"
  xmlns:xsd="http://www.w3.org/2001/XMLSchema#"
  xmlns:rdfs="http://www.w3.org/2000/01/rdf-schema#"
  xmlns:owl="http://www.w3.org/2002/07/owl#"
  xmlns="http://liris.cnrs.fr/~slajmi/Ontology/SocialSphere.owl#"
  xmlns:geo="http://www.w3.org/2003/01/geo/wgs84_pos#"
  xml:base="http://liris.cnrs.fr/~slajmi/ Ontology/SocialSphere.owl">
  <rdf:Description rdf:about="">
    <rdfs:label rdf:datatype="http://www.w3.org/2001/XMLSchema#string"
    >SocialSphere</rdfs:label>
    <rdfs:isDefinedBy rdf:resource="Sonia LAJMI"/>
  </rdf:Description>
  <owl:Ontology rdf:about="http://xmlns.com/foaf/0.1/">
  <rdfs:Class rdf:ID="SocialNetworkCategory">
    <rdfs:subClassOf>
      <rdfs:Class rdf:about="http://xmlns.com/foaf/0.1/Group"/>
    </rdfs:subClassOf>
  </rdfs:Class>
  <rdfs:Class rdf:ID="IntimeRelationship">
    <rdfs:subClassOf rdf:resource="#SocialNetworkCategory"/>
  </rdfs:Class>
  <rdfs:Class rdf:ID="Neighborship">
    <rdfs:subClassOf rdf:resource="#SocialNetworkCategory"/>
  </rdfs:Class>
  <rdfs:Class rdf:about="http://xmlns.com/foaf/0.1/Person"/>
  <rdfs:Class rdf:ID="ProfessionalRelationship">
    <rdfs:subClassOf rdf:resource="#SocialNetworkCategory"/>
  </rdfs:Class>
  <rdfs:Class rdf:ID="FamilyRelationship">
    <rdfs:subClassOf rdf:resource="#SocialNetworkCategory"/>
  </rdfs:Class>
  <rdfs:Class rdf:ID="FriendlyRelationship">
    <rdfs:subClassOf rdf:resource="#SocialNetworkCategory"/>
  </rdfs:Class>
  <owl:ObjectProperty rdf:ID="niece">
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    <owl:inverseOf>
      <owl:ObjectProperty rdf:ID="aunt"/>
    </owl:inverseOf>
  </owl:ObjectProperty>
  </owl:Ontology>
</rdf:RDF>
```

```

    </owl:inverseOf>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:ID="belong">
    <rdfs:domain rdf:resource="#SocialNetworkCategory"/>
    <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:ID="husband">
    <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    <owl:inverseOf>
      <owl:ObjectProperty rdf:ID="spouse"/>
    </owl:inverseOf>
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:ID="parentOf">
    <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    <owl:inverseOf>
      <owl:ObjectProperty rdf:ID="child"/>
    </owl:inverseOf>
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:ID="supervisor">
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    <owl:inverseOf>
      <owl:ObjectProperty rdf:ID="employed"/>
    </owl:inverseOf>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:ID="divorced">
    <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:ID="ex-boyFriend">
    <owl:inverseOf>
      <owl:ObjectProperty rdf:ID="ex-girlFriend"/>
    </owl:inverseOf>
    <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:ID="mentor">
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    <owl:inverseOf>
      <owl:ObjectProperty rdf:ID="apprentice"/>
    </owl:inverseOf>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:ID="ancestor">
    <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    <owl:inverseOf>
      <owl:ObjectProperty rdf:ID="descendant"/>
    </owl:inverseOf>
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>

```

```

</owl:ObjectProperty>
<owl:ObjectProperty rdf:ID="grandChild">
  <owl:inverseOf>
    <owl:ObjectProperty rdf:ID="grandParent"/>
  </owl:inverseOf>
  <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
  </owl:ObjectProperty>
  <owl:ObjectProperty rdf:about="#aunt">
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <owl:inverseOf rdf:resource="#niece"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:ID="uncle">
      <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
        <owl:inverseOf>
          <owl:ObjectProperty rdf:ID="nephew"/>
        </owl:inverseOf>
      </owl:ObjectProperty>
      <owl:ObjectProperty rdf:ID="ex-spouse">
        <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
        <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
        <owl:inverseOf>
          <owl:ObjectProperty rdf:ID="ex-husband"/>
        </owl:inverseOf>
      </owl:ObjectProperty>
      <owl:ObjectProperty rdf:about="#ex-husband">
        <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
        <owl:inverseOf rdf:resource="#ex-spouse"/>
        <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      </owl:ObjectProperty>
      <owl:ObjectProperty rdf:about="#ex-girlFriend">
        <owl:inverseOf rdf:resource="#ex-boyFriend"/>
        <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      </owl:ObjectProperty>
      <owl:ObjectProperty rdf:ID="boyFriend">
        <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
          <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
          <owl:inverseOf>
            <owl:ObjectProperty rdf:ID="girlFriendOf"/>
          </owl:inverseOf>
        </owl:ObjectProperty>
        <owl:ObjectProperty rdf:about="#employed">
          <owl:inverseOf rdf:resource="#supervisor"/>
          <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
            <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
          </rdfs:subPropertyOf>
        </owl:ObjectProperty>
      </owl:ObjectProperty>
    </owl:ObjectProperty>
  </owl:ObjectProperty>

```

```

    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:about="#child">
      <owl:inverseOf rdf:resource="#parent"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:about="#spouse">
      <owl:inverseOf rdf:resource="#husband"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:ID="sister">
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:about="#girlFriend">
      <owl:inverseOf rdf:resource="#boyFriend"/>
      <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:about="#grandParent">
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <owl:inverseOf rdf:resource="#grandChild"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:about="#apprentice">
      <owl:inverseOf rdf:resource="#mentor"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:about="#descendant">
      <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <owl:inverseOf rdf:resource="#ancestor"/>
    </owl:ObjectProperty>
    <owl:ObjectProperty rdf:about="#nephew">
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <owl:inverseOf rdf:resource="#uncle"/>
      <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </owl:ObjectProperty>
    <owl:TransitiveProperty rdf:ID="brother">
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </owl:TransitiveProperty>

```

```

    <owl:TransitiveProperty rdf:ID="sibling">
      <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <owl:inverseOf rdf:resource="#sibling"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#SymmetricProperty"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    </owl:TransitiveProperty>
    <owl:SymmetricProperty rdf:ID="worksWith">
      <owl:inverseOf rdf:resource="#worksWith"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="ex-lifePartner">
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <owl:inverseOf rdf:resource="#ex-lifePartner"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="lostContactWith">
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <owl:inverseOf rdf:resource="#lostContactWith"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="lifePartnerOf">
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <rdfs:range rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <owl:inverseOf rdf:resource="#lifePartner"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="livesWith">
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <owl:inverseOf rdf:resource="#livesWith"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="colleague">
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <owl:inverseOf rdf:resource="#colleague"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    </owl:SymmetricProperty>

```

```

    <owl:SymmetricProperty rdf:ID="acquaintance">
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <owl:inverseOf rdf:resource="#acquaintance"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="cousin">
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <owl:inverseOf rdf:resource="#cousin"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="fiance">
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <owl:inverseOf rdf:resource="#fiance"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="hasMet">
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <owl:inverseOf rdf:resource="#hasMet"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="engaged">
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <owl:inverseOf rdf:resource="#engaged"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="closeFriend">
      <owl:inverseOf rdf:resource="#closeFriend"/>
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="ex-engaged">
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <owl:inverseOf rdf:resource="#ex-engaged"/>
      <rdfs:domain rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="collaborator">
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
      <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="neighbor">

```

```

    <owl:inverseOf rdf:resource="#neighbor"/>
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    </owl:SymmetricProperty>
    <owl:SymmetricProperty rdf:ID="friend">
    <owl:inverseOf rdf:resource="#friend"/>
    <rdfs:subPropertyOf
rdf:resource="http://xmlns.com/foaf/0.1/knows"/>
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    </owl:SymmetricProperty>
    <owl:FunctionalProperty rdf:ID="stable">
    <rdfs:domain rdf:resource="#SocialNetworkCategory"/>
    <rdfs:range>
    <owl:DataRange>
    <owl:oneOf rdf:parseType="Resource">
    <rdf:first
rdf:datatype="http://www.w3.org/2001/XMLSchema#string"
    >yes</rdf:first>
    <rdf:rest rdf:parseType="Resource">
    <rdf:first
rdf:datatype="http://www.w3.org/2001/XMLSchema#string"
    >no</rdf:first>
    <rdf:rest rdf:resource="http://www.w3.org/1999/02/22-rdf-
syntax-ns#nil"/>
    </rdf:rest>
    </owl:oneOf>
    </owl:DataRange>
    </rdfs:range>
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#DatatypeProperty"/>
    <rdfs:comment
rdf:datatype="http://www.w3.org/2001/XMLSchema#string"
    >property allowed to define if SocialNetworkCategory is stable or
unstable</rdfs:comment>
    </owl:FunctionalProperty>
</rdf:RDF>

```

Annexe II. L'ontologie *SeMAT*

Code de l'ontologie SEMAT sérialisé en XML/RDF

```
<?xml version="1.0"?>
  <rdf:RDF xmlns:foaf="http://xmlns.com/foaf/0.1/"
xmlns:owl="http://www.w3.org/2002/07/owl#"
xmlns:rdf="http://www.w3.org/1999/02/22-rdf-syntax-ns#"
xmlns:rdfs="http://www.w3.org/2000/01/rdf-schema#">
    <rdf:Description
rdf:about="http://liris.cnrs.fr/~slajmi/SematOntology">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Ontology"/>
    </rdf:Description>
    <rdf:Description rdf:about="#User">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
      <rdfs:subClassOf
rdf:resource="http://xmlns.com/foaf/0.1/Person"/>
    </rdf:Description>
    <rdf:Description rdf:about="#Personal_Photo">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
      <rdfs:subClassOf
rdf:resource="http://xmlns.com/foaf/0.1/Image"/>
    </rdf:Description>
    <rdf:Description rdf:about="#Activity">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
    </rdf:Description>
    <rdf:Description rdf:about="#Annotation">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
      <rdfs:subClassOf rdf:resource="#Activity"/>
    </rdf:Description>
    <rdf:Description rdf:about="#Event">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
    </rdf:Description>
    <rdf:Description rdf:about="#SocialInteraction">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
    </rdf:Description>
    <rdf:Description rdf:about="#Context">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
    </rdf:Description>
    <rdf:Description rdf:about="#ContextElement">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
    </rdf:Description>
    <rdf:Description rdf:about="#Group">
      <rdf:type rdf:resource="http://www.w3.org/2002/07/owl#Class"/>
    </rdf:Description>
    <rdf:Description rdf:about="#relatedWith">
      <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
```

```

        <rdfs:domain rdf:resource="#User"/>
        <rdfs:domain rdf:resource="#Group"/>
        <rdfs:range rdf:resource="#SocialInteraction"/>
    </rdf:Description>
    <rdf:Description rdf:about="#to">
        <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
        <rdfs:range rdf:resource="#User"/>
        <rdfs:domain rdf:resource="#SocialInteraction"/>
        <rdfs:range rdf:resource="#Group"/>
    </rdf:Description>
    <rdf:Description rdf:about="#belongs">
        <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
        <rdfs:domain rdf:resource="#User"/>
        <rdfs:range rdf:resource="#Group"/>
    </rdf:Description>
    <rdf:Description rdf:about="#actsBy">
        <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
        <rdfs:domain rdf:resource="#Activity"/>
        <rdfs:range rdf:resource="#User"/>
    </rdf:Description>
    <rdf:Description rdf:about="#doneOn">
        <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
        <rdfs:domain rdf:resource="#Activity"/>
        <rdfs:range rdf:resource="#Personal_Photo"/>
    </rdf:Description>
    <rdf:Description rdf:about="#displayedOn">
        <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
        <rdfs:domain rdf:resource="#SocialInteraction"/>
        <rdfs:range rdf:resource="#Personal_Photo"/>
    </rdf:Description>
    <rdf:Description rdf:about="#representedBy">
        <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
        <rdfs:domain rdf:resource="#Event"/>
        <rdfs:range rdf:resource="#Personal_Photo"/>
    </rdf:Description>
    <rdf:Description rdf:about="#contextActivity">
        <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
        <rdfs:domain rdf:resource="#Activity"/>
        <rdfs:range rdf:resource="#Context"/>
    </rdf:Description>
    <rdf:Description rdf:about="#hasContextElement">
        <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
        <rdfs:domain rdf:resource="#Context"/>
        <rdfs:range rdf:resource="#ContextElement"/>
    </rdf:Description>

```

```
<rdf:Description rdf:about="#tag">
  <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#DatatypeProperty"/>
    <rdfs:domain rdf:resource="#Annotation"/>
    <rdfs:range
rdf:resource="http://www.w3.org/2001/XMLSchema#string"/>
  </rdf:Description>
  <rdf:Description rdf:about="#acts">
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    <owl:inverseOf rdf:resource="#actsBy"/>
  </rdf:Description>
  <rdf:Description rdf:about="#displayedTo">
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    <owl:inverseOf rdf:resource="#role"/>
  </rdf:Description>
  <rdf:Description rdf:about="#hasRole">
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    <rdfs:domain rdf:resource="#User"/>
    <rdfs:range rdf:resource="#Personal_Photo"/>
  </rdf:Description>
  <rdf:Description rdf:about="#actor">
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    <rdfs:subPropertyOf rdf:resource="#hasRole"/>
  </rdf:Description>
  <rdf:Description rdf:about="#photograph">
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    <rdfs:subPropertyOf rdf:resource="#hasRole"/>
  </rdf:Description>
  <rdf:Description rdf:about="#witness">
    <rdf:type
rdf:resource="http://www.w3.org/2002/07/owl#ObjectProperty"/>
    <rdfs:subPropertyOf rdf:resource="#hasRole"/>
  </rdf:Description>
</rdf:RDF>
```


Annexe III. Règles d'inférences explicites de sens commun

Les règles ci-dessous permettent le raisonnement sur les liens sociaux

```
Jena Rules
@prefix SocialSphere:
<http://liris.cnrs.fr/~slajmi/ProfilFoaf/socialSphere.owl#>

#####
#social_Rules
#####

[FamilyRule1:
 ( ?P1 SocialSphere:sibling ?P2)
 (?P1 foaf:gender "women")
 --> (?P1 SocialSphere:sister ?P2)]
[FamilyRule2:
 ( ?P1 SocialSphere:sibling ?P2)
 (?P1 foaf:gender "men")
 --> (?P1 SocialSphere:brother ?P2)]
[FamilyRule3:
 (?P1 SocialSphere:parent ?P2)
 (?P1 foaf:gender "men")
 --> (?P1 SocialSphere:father ?P2)]
[FamilyRule4:
 (?P1 SocialSphere:parent ?P2)
 (?P1 foaf:gender "women")
 --> (?P1 SocialSphere:mother ?P2)]
[FamilyRule5:
 (?P1 SocialSphere:grandParent ?P2)
 (?P1 foaf:gender "women")
 --> (?P1 SocialSphere:grandMother ?P2)]
[FamilyRule6:
 (?P1 SocialSphere:grandParent ?P2)
 (?P1 foaf:gender "men")
 --> (?P1 SocialSphere:grandFather ?P3)]
[FamilyRule7:
 ( ?P1 SocialSphere:sister ?P2)
 ( ?P3 SocialSphere:child ?P2)
 --> (?P1 SocialSphere:aunt ?P3)]
[FamilyRule8:
 ( ?P1 SocialSphere:brother ?P2)
 ( ?P3 SocialSphere:child ?P2)
```

```

--> (?P1 SocialSphere:uncle ?P3)]
[FamilyRule9:
  (?P1 SocialSphere:sibling ?P2)
  (?P3 SocialSphere:child ?P2)
--> (?P1 SocialSphere:uncle ?P3)]
[FamilyRule1:
  ( ?P1 SocialSphere:sibling ?P2)
  (?P1 foaf:gender "women")
  ( ?P3 SocialSphere:child ?P1)
  ( ?P4 SocialSphere:child ?P2)
--> (?P1 SocialSphere:aunt ?P3)
[FamilyRule1:
  ( ?P1 SocialSphere:sibling ?P2)
  (?P1 foaf:gender "men")
  ( ?P3 SocialSphere:child ?P1)
  ( ?P4 SocialSphere:child ?P2)
--> (?P1 SocialSphere:aunt ?P3)
and (?P3 SocialSphere:cousin ?P4)]
[FamilyRule2:
  ( ?P1 SocialSphere:sibling ?P2)
  (?P1 foaf:gender "men")
  ( ?P3 SocialSphere:child ?P1)
  ( ?P4 SocialSphere:child ?P2)
--> (?P1 SocialSphere:uncle ?P3)
and (?P3 SocialSphere:cousin ?P4)]
[FamilyRule3:
  ( ?P1 SocialSphere:sibling ?P2)
  (?P1 foaf:gender "men")
  ( ?P3 SocialSphere:child ?P1)
--> (?P1 SocialSphere:uncle ?P3)
]
[FamilyRule4:
  ( ?P1 SocialSphere:sibling ?P2)
  (?P1 foaf:gender "women")
  ( ?P3 SocialSphere:child ?P1)
--> (?P1 SocialSphere:aunt ?P3)
]
[FamilyRule5:
  ( ?P1 SocialSphere:aunt ?P2)
  or ( ?P1 SocialSphere:uncle ?P2)
  (?P2 foaf:gender "women")
--> (?P2 SocialSphere:niece ?P3)
]
[FamilyRule6:
  ( ?P1 SocialSphere:aunt ?P2)
  or ( ?P1 SocialSphere:uncle ?P2)
  (?P2 foaf:gender "men")
--> (?P2 SocialSphere:nephew ?P3)
]
[FamilyRule7:
  ( ?P1 SocialSphere:spouse ?P2)
  (?P3 SocialSphere:sibling ?P2)

```

```

(?P3 foaf:gender "women")
--> (?P3 SocialSphere:sister-in-law ?P2)
]
[FamilyRule8:
  ( ?P1 SocialSphere:spouse ?P2)
  (?P3 SocialSphere:sibling ?P2)
  (?P3 foaf:gender "men")
--> (?P3 SocialSphere:brother-in-law ?P2)
]
[FamilyRule9:
  ( ?P1 SocialSphere:sibling ?P2)
  (?P1 foaf:gender "men")
--> (?P1 SocialSphere:brother ?P2)
]
[FamilyRule10:
  ( ?P1 SocialSphere:sibling ?P2)
  (?P1 foaf:gender "women")
--> (?P1 SocialSphere:sister ?P2)
]
[FamilyRule11:
  ( ?P1 SocialSphere:parent ?P2)
  (?P3 SocialSphere:child ?P2)
--> (?P1 SocialSphere:grandParent ?P2)
]
[FamilyRule12:
  ( ?P1 SocialSphere:parent ?P2)
  (?P1 foaf:gender "women")
--> (?P1 SocialSphere:mother ?P2)
]
[FamilyRule13:
  ( ?P1 SocialSphere:parent ?P2)
  (?P1 foaf:gender "men")
--> (?P1 SocialSphere:father ?P2)
]
[FamilyRule14:
  ( ?P1 SocialSphere:lifePartner ?P2)
  (?P1 foaf:gender "men")
  (?P2 foaf:gender "women")
--> (?P1 SocialSphere:boyFriend ?P2)
  or (?P1 SocialSphere:husband ?P2)
  or (?P1 SocialSphere:fiance ?P2)
]
[FamilyRule15:
  ( ?P1 SocialSphere:engagedTo ?P2)
--> (?P1 SocialSphere:fiance ?P2)
]
[FamilyRule16:
  ( ?P1 SocialSphere:divorcedWith ?P2)
  (?P1 foaf:gender "men")
--> (?P1 SocialSphere:ex-husband ?P2)
]
[FamilyRule17:

```

```

(?P1 SocialSphere:ex-lifePartner ?P2)
(?P1 foaf:gender "men")
  (?P2 foaf:gender "women")
--> (?P1 SocialSphere:ex-boyFriend ?P2)
or (?P1 SocialSphere:ex-husband ?P2)
  or (?P1 SocialSphere:ex-fiance ?P2)
]
[FriendRule18:
  (?P1 SocialSphere:hasMet ?P2)
  --> (?P1 SocialSphere:acquaintance ?P2)
]
#####
#SocialSphere_Rules
#####
[NeighborhoodCategory Rule:
  (?P1 foaf:member SocialSphere:Neighborhood)
  (SocialSphere:Neighborhood foaf:name ?C)
  (?C SocialSphere:belongsTo ?P2)
  --> (?P1 SocialSphere:livesWith ?P2)
  or (?P1 SocialSphere:neighbor ?P2)
]
[ProfessionalRelationshipCategory Rule:
  (?P1 foaf:member SocialSphere:Neighborhood)
  (SocialSphere:Neighborhood foaf:name ?C)
  (?C SocialSphere:belongsTo ?P2)
  --> (?P1 SocialSphere:worksWith ?P2)
  or (?P1 SocialSphere:colleague ?P2)
  or (?P1 SocialSphere:collaboratesWith ?P2)
  or (?P1 SocialSphere:employedBy ?P2)
  or (?P1 SocialSphere:mentor ?P2)
  or (?P1 SocialSphere:apprenticeTo ?P2)
  or (?P1 SocialSphere:supervisor ?P2)
]
[FamilyRelationshipCategory Rule:
  (?P1 foaf:member SocialSphere:FamilyRelationship)
  (SocialSphere:FamilyRelationship foaf:name ?C)
  (?C SocialSphere:belongsTo ?P2)
  --> (?P1 SocialSphere:parent ?P2)
  or (?P1 SocialSphere:child ?P2)
  or (?P1 SocialSphere:grandChild ?P2)
  or (?P1 SocialSphere:grandParent ?P2)
  or (?P1 SocialSphere:descendant ?P2)
  or (?P1 SocialSphere:sibling ?P2)
  or (?P1 SocialSphere:uncle ?P2)
  or (?P1 SocialSphere:cousin ?P2)
  or (?P1 SocialSphere:niece ?P2)
  or (?P1 SocialSphere:nephew ?P2)
]
[FriendlyRelationshipCategory Rule:
  (?P1 foaf:member SocialSphere:FriendlyRelationship)
  (SocialSphere:FriendlyRelationship foaf:name ?C)
  (?C SocialSphere:belongsTo ?P2)

```

```

--> (?P1 SocialSphere:friend ?P2)
or (?P1 SocialSphere:acquaintance ?P2)
or (?P1 SocialSphere:lostContactWith ?P2)
or (?P1 SocialSphere:closeFriend ?P2)
or (?P1 SocialSphere:hasMet ?P2)
]
[IntimeRelationshipCategory Rule:
(?P1 foaf:member SocialSphere:IntimeRelationship)
(SocialSphere:IntimeRelationship foaf:name ?C)
(?C SocialSphere:belongTo ?P2)
--> (?P1 SocialSphere:spouse ?P2)
or (?P1 SocialSphere:lifePartner ?P2)
or (?P1 SocialSphere:engagedTo ?P2)
or (?P1 SocialSphere:girlfriend ?P2)
or (?P1 SocialSphere:boyFriend ?P2)
]
#####
#Substitution_Rules
#####
[Uncle_Rule:
(?P1 SocialSphere:husband ?P2)
(?P2 SocialSphere:aunt ?P3)
--> (?P1 SocialSphere:uncle ?P3)
]

```


Annexe IV. Exemples d'utilisation du Framework Jena

Le Code 9 représente un exemple de création d'un graphe, ou d'un modèle au sens Jena. Nous commençons par quelques définitions des constantes, puis nous créons un modèle vide, en utilisant une méthode baptisée *createDefaultModel()* de la classe *ModelFactory*. La bibliothèque Jena fournit d'autres implémentations de l'interface modèle (*Model*). Nous citons, par exemple, celle qui utilise une base de données relationnelle. Dans l'exemple présenté dans le Code 9, la ressource de *John Smith* est créée et une propriété est ajoutée. La propriété est fournie par une classe *VCARD* qui contient des objets représentant toutes les définitions dans le schéma *VCARD*.

```
// quelques définitions
static String personURI = "http://somewhere/JohnSmith";
static String fullName = "John Smith";
// creation d'un modèle vide
Model model = ModelFactory.createDefaultModel();
// creation d'une ressource
Resource johnSmith = model.createResource(personURI);
// ajout d'une propriété
johnSmith.addProperty(VCARD.FN, fullName);
```

Code 11. Création d'une ressource avec le modèle de Jena

Ajoutons un peu plus des détails à la *VCARD* (voir Code 12), en explorant quelques fonctionnalités de *RDF* et de la bibliothèque Jena. Dans l'exemple décrit par le Code 9, la valeur de la propriété était un littéral. Les propriétés *RDF* peuvent également prendre d'autres valeurs que leurs valeurs littérales. Pour étendre cet exemple, nous pouvons ajouter de nouvelles propriétés, *vcard:N*, pour représenter le nom de *John Smith*. La propriété *vcard:N* prend comme valeur la ressource. Nous pouvons, aussi, représenter le nom composé qui n'a pas *URI* connu comme un nœud blanc. Le code java pour mettre en place l'exemple étendu est représenté par le Code 12.

```
// quelques définitions
String personURI    = "http://somewhere/JohnSmith";
String givenName    = "John";
String familyName   = "Smith";
String fullName     = givenName + " " + familyName;
// creation d'un modèle vide
Model model = ModelFactory.createDefaultModel();
// creation d'une ressource
//et y ajouter les propriétés associées
    = model.createResource(personURI)
        .addProperty(VCARD.FN, fullName)
        .addProperty(VCARD.N,
            model.createResource()
                .addProperty(VCARD.Given, givenName)
                .addProperty(VCARD.Family, familyName));
```

Code 12. Une déclaration RDF

Chaque arc dans un modèle RDF est appelé une déclaration. Chaque déclaration affirme un fait d'une ressource. Une déclaration possède trois parties : le sujet est la ressource à partir de laquelle l'arc part ; le prédicat est la propriété qui étiquète l'arc et l'objet est la ressource ou la littérale destination de l'arc.

Un modèle RDF est représenté comme un ensemble de déclarations. Dans le Code 12, chaque appel de *addProperty* permet d'ajouter une autre déclaration au modèle initiale. La liste des déclarations du modèle est obtenue en utilisant la méthode *listStatements()* qui retourne un *StmtIterator* : un sous-type de *Iterateur* de Java. La classe *StmtIterator* possède une méthode *nextStatement()* qui retourne la prochaine déclaration de l'itérateur. L'interface *Statement* fournit des méthodes d'accesseur sur le sujet, prédicat et objet d'une déclaration.

Le Code 13 illustre un exemple permettant de lister toute les déclarations d'un modèle et d'afficher pour chaque déclaration le sujet, le prédicat et l'objet correspondant.

```
// liste des declarations du modèle
StmtIterator iter = model.listStatements();
// afficher le prédicate, le sujet et l'objet
while (iter.hasNext()) {
    Statement stmt = iter.nextStatement();
    Resource subject = stmt.getSubject();
    Property predicate = stmt.getPredicate();
    RDFNode object = stmt.getObject();
    System.out.print(subject.toString());
    System.out.print(" " + predicate.toString() + " ");
    if (object instanceof Resource) {
        System.out.print(object.toString());
    } else {
        // l'objet est un littéral
        System.out.print(" \"" + object.toString() + "\"");
    }
}
```

Code 13. Affichage du prédicat, sujet et objet

Puisque l'objet d'une déclaration peut être soit une ressource ou un littéral, la méthode *getObject()* retourne un objet de type *RDFNode*, qui est une superclasse commune des deux types ressource (*Ressource*) et littéral (*Literal*).

La bibliothèque Jena fournit, également, des opérations pour manipuler des modèles dans leurs ensembles. L'ensemble des opérations communes comme union (*Model1.union (Model2)*), intersection (*Model1.intersection (Model2)*) et différence (*Model1.difference (Model2)*); peut être utilisé pour manipuler les modèles de données.

Pour permettre le raisonnement et l'application de règles d'inférence a un modèle, la bibliothèque Jena fournit une autre classe l'*OntModel* qui permet de créer un modèle ontologique. L'*OntModel* étend, donc, le modèle RDF (*Model*) en ajoutant un support pour les types d'objets: les classes (dans une hiérarchie de classes), les propriétés (dans une hiérarchie de propriétés) et d'autres.. Par exemple, un *OntClass* a une méthode pour lister ses super-classes, ce qui correspond aux valeurs de la propriété *subClassOf*. Lorsque la méthode *listSuperClasses()* de la classe *OntClass* est appelée, les informations sont extraites des déclarations RDF sous-jacente.

Les modèles *Ontologie* sont créées par le *ModelFactory*. La plus simple manière pour créer un modèle d'ontologie est la suivante :

```
OntModel m = ModelFactory.createOntologyModel();
```

Pour créer un modèle avec une spécification donnée, il est recommandé d'invoquer le *ModelFactory* comme suit:

```
OntModel m =ModelFactory.createOntologyModel  
(OntModelSpec.OWL_MEM);
```

.

Liste de publications

- Addisalem Negash Shiferaw, **Sonia Lajmi**, Vasile-Marian Scuturici, Lionel Brunie (2010) "PASMi: self-Adaptive Photo Annotation and Sharing Middleware" in proceeding of Annual IEEE International Conference on Pervasive Computing and Communications (PerCom 2010), march 29 - april 2, Mannheim, Germany.
- **Sonia Lajmi**, Johann Stan, Hakim Hacid, Elöd Egyed-Zsigmond, Pierre Maret (2009) "Extended Social Tags: Identity Tags Meet Social Networks" in 2009 IEEE International Conference on Social Computing (SocialCom-09), 29-31 august, Vancouver, Canada.
- **Sonia Lajmi**, Elöd Egyed-Zsigmond, Abdelmajid Ben Hamadou, Jean-Marie Pinon (2009) "Enrichissement de l'annotation des identités des personnes dans les photos personnelles". Dans Prise en Compte de l'Usager dans les Systèmes d'Information PeCUSI 2009 Atelier conjoint à INFORSID 2009, Toulouse, France.
- Youssef Matar, Elöd Egyed-Zsigmond, **Sonia Lajmi** (2008). "KWSim: Concept Similarity Measure". CONFérence en Recherche d'Information et Applications, (CORIA 2008), ARIA ed. Trégastel, France. pp. 475-482. 2008.
- Elöd Egyed-Zsigmond, **Sonia Lajmi** Zoltan Iszlai, (2007) «PhotoMot, collaborative image management Interacting with use traces». In IEEE The Second International Conference on Digital Information Management (ICDIM'07), Lyon, France. pp. 104-109. IEEE . ISBN 1-4244-1476-8. 2007.
- **Sonia Lajmi**, Abdelmajid Ben Hamadou, Elöd Egyed-Zsigmond et Jean-Marie Pinon (2007). « Système en ligne de gestion collaborative d'images médicales », Dans RJCIA'07, Plate-forme AFIA, Grenoble, France. pp. 255-260. ISBN 978-2-85428-792-9. 2007.
- **Sonia Lajmi** (2007). « Tracer l'utilisation d'une base d'images », Colloque Jeunes Chercheurs en Sciences Cognitives (CJC-SC'07) 30-31 mai, 1er juin 2007. Lyon, France.
- **Sonia Lajmi**, Elöd Egyed-Zsigmond, Jean-Marie Pinon et Abdelmajid Ben Hamadou (2007). « SyLGeCoM: Système en ligne de gestion collaborative d'images médicales », 3èmes Rencontres Inter-Associations, (RIA's2007), Toulouse, 12 et 13 Mars 2007. Accessible en ligne depuis < <http://www.irit.fr/RIA07/intervenants.html> >.
- Elöd Egyed-Zsigmond, **Sonia Lajmi**, Zoltan Iszlai (2006). « Concurrent use in an image management system », 13th international conference on Concurrent Engineering: research and applications, (CE2006), Antibes, September 2006.

- **Sonia Lajmi**, Zoltan Iszlai, Elöd Egyed-Zsigmond (2006). « User Centered Image Management System », Workshop on Knowledge Management and Organizational Memories, (ECAI'2006), Riva del Garda - Italy, 28 août 2006.

Curriculum Vitae

Information Personnelle

Nom : LAJMI
Prénom : Sonia
Année et lieu de naissance : 1982 à Sfax (Tunisie)
Nationalité : Tunisienne
Adresse : INSA de Lyon, LIRIS Bâtiment B. Pascal (501). Bureau (315)
20 Avenue Albert Einstein, 69621 Villeurbanne, France
Tel. (Tunisie) + 216 25 11 25 55
Tel. (France) +33 (0)4 72 43 72 29
Fax. (France) +33 (0)4 72 43 87 13

Formation

- Doctorat en Informatique - LIRIS, INSA – Lyon- France (Décembre 2006-présent)
- Master de recherche en informatique, INSA- Lyon. France (2006)
- Maîtrise en "Informatique Système et Multimédia" ; ISIMSF de Sfax Tunisie (2005)
- DUT en "Informatique Système et Multimédia"; ISIMSF de Sfax Tunisie (2003)
- Baccalauréat scientifique spécialité Mathématiques Lycée ‘20 Mars 1956’ Sfax-Tunisie (2001).

Expérience Professionnelle

- ATER à temps plein l’Université Jean Monnet Faculté des sciences et IUT de Saint-Etienne (2009-2010)
- ATER à temps partiel à l’Institut national des sciences appliquées (INSA) de Lyon (2008-2009)
- Enseignante chez la société ‘Après la classe’ (2007-2008)
- Enseignante Vacataire à l’Institut National des Sciences Appliquées (INSA) de Lyon Premier cycle (2007-2008).
- Enseignante Vacataire au Département Génie Electrique Informatique Industrielle - IUT B - UCB Lyon1 (2007-2008).
- Assistante contractuelle à l’Institut Supérieur en Informatique et Multimédia de SFAX (ISIMSF)- Tunisie (2006-2007)

Connaissances Linguistiques

- Français (courant)
- Anglais (courant)
- Arabe (langue maternelle)

FOLIO ADMINISTRATIF
THESE SOUTENUE DEVANT L'INSTITUT NATIONAL DES SCIENCES APPLIQUEES DE
LYON

NOM : Lajmi

DATE de SOUTENANCE : 11 mars 2011

Prénom : Sonia

TITRE : Annotation et recherche contextuelle des documents multimédias socio-personnels

NATURE : Doctorat

Numéro d'ordre : 2011-ISAL-0025

Ecole doctorale : Informatique et Mathématiques (INFOMATHS)

Spécialité : Informatique

Cote B.I.U. - Lyon : T 50/210/19 / et bis CLASSE :

RESUME :

L'objectif de cette thèse est d'instrumentaliser des moyens, centrés utilisateur, de représentation, d'acquisition, d'enrichissement et d'exploitation des métadonnées décrivant des documents multimédias socio-personnels.

Afin d'atteindre cet objectif, nous avons proposé un modèle d'annotation, appelé SeMAT avec une nouvelle vision du contexte de prise de vue. Nous avons proposé d'utiliser des ressources sémantiques externes telles que GeoNames , et Wikipédia pour enrichir automatiquement les annotations partant des éléments de contexte capturés. Afin d'accentuer l'aspect sémantique des annotations, nous avons modélisé la notion de profil social avec des outils du web sémantique en focalisant plus particulièrement sur la notion de liens sociaux et un mécanisme de raisonnement permettant d'inférer de nouveaux liens sociaux non explicites. Le modèle proposé, appelé SocialSphere, construit un moyen de personnalisation des annotations suivant la personne qui consulte les documents (le consultateur). Des exemples d'annotations personnalisées peuvent être des objets utilisateurs (e.g. maison, travail) ou des dimensions sociales (e.g. ma mère, le cousin de mon mari). Dans ce cadre, nous avons proposé un algorithme, appelé SQO, permettant de suggérer au consultateur des dimensions sociales selon son profil pour décrire les acteurs d'un document multimédia. Dans la perspective de suggérer à l'utilisateur des événements décrivant les documents multimédias, nous avons réutilisé son expérience et l'expérience de son réseau de connaissances en produisant des règles d'association. Dans une dernière partie, nous avons abordé le problème de correspondance (ou appariement) entre requête et graphe social. Nous avons proposé de ramener le problème de recherche de correspondance à un problème d'isomorphisme de sous-graphe partiel. Nous avons proposé un algorithme, appelé h-Pruning, permettant de faire une correspondance rapprochée entre les nœuds des deux graphes : motif (représentant la requête) et social.

Pour la mise en œuvre, nous avons réalisé un prototype à deux composantes : web et mobile. La composante mobile a pour objectif de capturer les éléments de contexte lors de la création des documents multimédias socio-personnels. Quant à la composante web, elle est dédiée à l'assistance de l'utilisateur lors de son annotation ou consultation des documents multimédias socio-personnels. L'évaluation a été effectuée en se servant d'une collection de test construite à partir du service de médias sociaux Flickr. Les tests ont prouvé : (i) l'efficacité de notre approche de recherche dans le graphe social en termes de temps d'exécution ; (ii) l'efficacité de notre approche de suggestion des événements (en effet, nous avons prouvé notre hypothèse en démontrant l'existence d'une cooccurrence entre le contexte spatio-temporel et les événements) ; (iii) l'efficacité de notre approche de suggestion des dimensions sociales en termes de temps d'exécution.

MOTS-CLES : Annotation et Recherche Sémantique, Web 2.0, Web Sémantique, Ontologies, Réseaux sociaux, Raisonnement Sémantique, Sensibilité au contexte social, Profil social, Graphe, Appariement de graphe

Laboratoire (s) de recherche : LIRIS - Laboratoire d'InfoRmatique en Image et Systèmes d'information

Directeur de thèse: Pr. Jean Marie Pinon

Co-directeur de thèse :Pr. Abdelmajid Ben Hamadou

Rapporteurs :

Mohamed Mohsen GAMMOUDI Maître de conférences HDR à l'ESSAI, Univ. de Carthage, Tunisie

Philippe MULHEM Chargé de recherche Première Classe au CNRS, LIG, Université de Grenoble

Président de jury : Corine CAUVET

Composition du jury :

Mohamed Mohsen GAMMOUDI Maître de conférences HDR à l'ESSAI, Univ. de Carthage, Tunisie (Rapporteur)

Philippe MULHEM Chargé de recherche Première Classe au CNRS, LIG, Université de Grenoble (Rapporteur)

Christine LARGERON épouse LETENO Professeur à l'univ. Jean Monnet - Saint-Etienne (UJM) (Examinatrice)

Corine CAUVET Professeur à l'univ. Paul Cezanne - Aix-Marseille 3 (Examinatrice)

Jean Marie Pinon, Professeur à l'INSA de Lyon, France (Directeur de thèse)

Abdelmajid Ben Hamadou, Professeur à l'ISIMSF de Sfax, Tunisie (Directeur de thèse)

Elöd Egyed Zsigmond, Maître de conférences à l'INSA de Lyon, France (Co-Directeur de thèse)