



Vers un système hybride pour l'annotation sémantique d'images IRM du cerveau

Ammar Mechouche, Christine Golbreich, Bernard Gibaud

► To cite this version:

Ammar Mechouche, Christine Golbreich, Bernard Gibaud. Vers un système hybride pour l'annotation sémantique d'images IRM du cerveau. 7èmes Journées Francophones en Extraction et Gestion des Connaissances, EGC 2007, Jan 2007, Namur, Belgium. pp.433-444. inserm-00151070

HAL Id: inserm-00151070

<https://inserm.hal.science/inserm-00151070>

Submitted on 4 Jun 2007

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.

Vers un système hybride pour l'annotation sémantique d'images IRM du cerveau

Ammar Mechouche*, Christine Golbreich**, Bernard Gibaud*

* INSERM, U746, Campus de Villejean, F-35043 Rennes, France

* INRIA, VisAGeS Unité/Projet, F-35042 Rennes, France

* Univ. Rennes I, CNRS, UMR 6074, IRISA, F-35042 Rennes, France

{Ammar.Mechouche, Bernard.Gibaud}@irisa.fr

<http://www.irisa.fr/visages/>

** LIM, Univ Rennes I, Av du Pr. Léon Bernard, 35043 Rennes, France

Christine.Golbreich@univ-rennes1.fr

<http://www.med.univ-rennes1.fr/>

Résumé. Cet article montre l'intérêt de combiner des méthodes numériques et symboliques pour obtenir une annotation sémantique des images IRM du cerveau humain. Il s'agit d'identifier des structures anatomiques du cortex cérébral humain, en utilisant conjointement des connaissances a priori de nature numérique et une ontologie des structures corticales du cerveau représentée en OWL DL, étendue par des règles SWRL. Ces connaissances symboliques a priori représentées dans des langages standards du Web deviennent non seulement partageables mais permettent aussi un raisonnement automatique qui aide l'utilisateur à la labellisation des structures anatomiques mises en évidence dans des images IRM du cerveau d'un individu donné.

1 Introduction

L'identification des structures anatomiques dans des images de résonance magnétique (IRM) du cerveau est un aspect important de la préparation d'une intervention chirurgicale en neurochirurgie, notamment lorsqu'une lésion est localisée sur le cortex cérébral. L'aide à la labellisation des structures corticales (gyri et sillon) entourant une lésion est particulièrement nécessaire pour déterminer la stratégie optimale de l'intervention. Il existe dans la littérature d'autres approches pour la labellisation automatique des structures anatomiques du cerveau, en particulier utilisant les SPAMs Collins et al. (1999) (Statistical Probability Anatomy Maps) qui sont des cartes probabilistes 3D des structures anatomiques du cerveau. Dans notre cas, l'information statistique est dérivée d'une base de données de 305 examens IRM réalisés chez des sujets normaux, et ré-alignés dans un système référentiel commun (nommé l'espace stéréotaxique). Les méthodes fondées sur l'utilisation d'atlas numériques comme les SPAMs ont un inconvénient. Elles ne sont pas robustes vis à vis des déformations causées par la présence d'une lésion dans le cerveau. Les méthodes symboliques, utilisant des connaissances a priori sur les relations topologiques entre les structures cérébrales peuvent constituer une alternative, ou un

complément. Cependant, afin d'avoir un processus de raisonnement efficace, il est nécessaire d'avoir suffisamment d'informations de départ pour pouvoir inférer de nouveaux faits. Dans ce papier une nouvelle méthode hybride d'étiquetage est proposée, dans laquelle les SPAMs sont utilisées pour fournir automatiquement des faits initiaux requis par notre méthode symbolique. L'objectif général du travail présenté ici est d'assister un neurochirurgien dans sa tâche d'identification, grâce à des connaissances explicites et partagées, sous la forme d'une ontologie des structures corticales du cerveau représentée en OWL DL OWL (2004), et d'une base de règles qui capturent les dépendances entre les propriétés de ces structures et les faits initiaux extraits des images et ceux fournis par les SPAMs.

2 Méthode

Notre méthode comprend deux étapes principales : la première est la segmentation du cerveau et l'extraction du graphe des sillons à partir d'un examen IRM. La deuxième étape qui est la plus importante dans notre étude est l'étiquetage des structures en présence dans la région d'intérêt sélectionnée par l'utilisateur. Ce papier se focalise sur cette deuxième étape. La figure 1 montre le schéma global de la méthode.

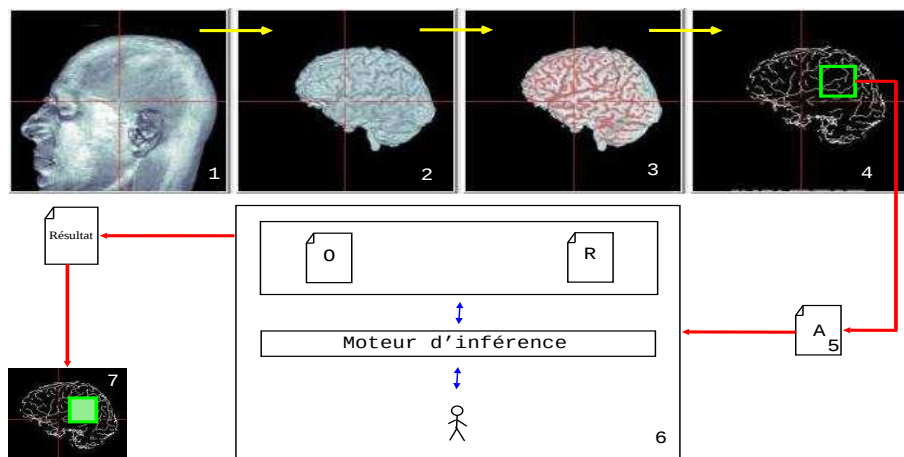


FIG. 1 – Le schéma général de la méthode

La méthode proposée utilise des connaissances a priori sous la forme d'une ontologie des structures corticales du cerveau, représentée en OWL DL, étendue avec des règles de Horn. Elle utilise aussi une base de faits extraits des images et complétés de faits obtenus à partir des SPAMs. La méthode met en jeu plusieurs étapes successives, représentées sur la figure 1. : (1) acquisition de l'IRM du patient, suivie de différents prétraitements, (2) segmentation du cerveau, (3) extraction des traces externes des sillons corticaux, (4) sélection par l'utilisateur d'une région d'intérêt (sous-ensemble) du graphe global, (5) extraction du sous-graphe correspondant à cette région, délimitation des surfaces formant des parties de gyri (patches) et description des relations entre ces patches et les traces externes des sillons (relations topolo-

giques, et d'orientation), cette phase n'est pas complètement automatique dans notre système. Ces derniers traitements conduisent à une base de faits $\mathcal{A}$ représentés dans un fichier OWL DL, dont notamment des faits émanant des SPAMs comme expliqué plus loin. L'étape (6) concerne le processus de raisonnement, basé sur l'ontologie du cerveau $\mathcal{O}$ (base de connaissances), la base des règles $\mathcal{R}$, et la base des faits $\mathcal{A}$, le moteur d'inférence, avec une certaine interaction avec l'utilisateur. Ce traitement infère les étiquettes pour les structures anatomiques extraites de la région choisie. La base de faits est alimentée essentiellement à partir de deux sources : les faits extraits directement des images (extraction de la région d'intérêt), et ceux extraits des SPAMs.

2.1 Extraction de la région d'intérêt

D'abord un sous-graphe de traces externes des sillons correspondant à une région d'intérêt que l'utilisateur sélectionne est automatiquement extrait du graphe global, comme montré sur la figure 2. Ici le graphe désigne l'ensemble des traces externes où chaque trace correspond à un tracé surfacique d'un sillon du cerveau.

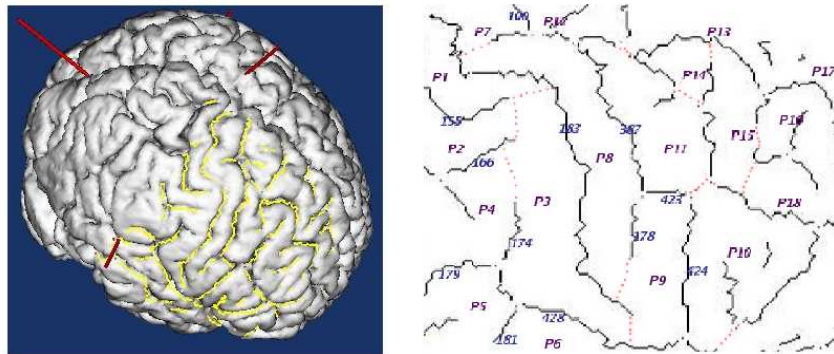


FIG. 2 – Première source de faits : l'extraction du graphe d'intérêt

Après l'extraction de la région d'intérêt, la délimitation des patches est effectuée. En effet, cette phase consiste à découper l'ensemble de la surface du sous-graphe extrait en plusieurs surfaces délimitées par les traces externes constituant le sous-graphe. Ces surfaces sont appelées ici "patches". Par exemple sur la figure 2 le patch P_9 est délimité par les traces 178, 423, 424 ... etc. Ensuite, les relations (topologiques et d'orientation) entre les traces externes et les patches sont extraites et décrites dans un fichier OWL DL. Un exemple des faits initiaux extraits est illustré par la figure 3.

2.2 Les SPAMs

Une SPAM (figure 4) est une image volumique 3D calculée dans l'espace stéréotaxique, et associée à une structure anatomique particulière, par exemple un gyrus donné. L'informa-

Annotation Sémantique d'Images IRM du Cerveau

```

<Patch rdf:ID="p9">
...
<isMAEBoundedBy>
  <AttributedEntity rdf:ID="AttEntity1">
    <entity>
      <SulcusSegment rdf:ID="178"/>
    </entity>
    <orientation>
      <Posterior rdf:ID="posteriorTo"/>
    </orientation>
    <orientation>
      <Unknown rdf:ID="unknown1"/>
    </orientation>
    <orientation>
      <Unknown rdf:ID="unknown2"/>
    </orientation>
    <MAEBounds rdf:resource="#p9"/>
  </AttributedEntity>
</isMAEBoundedBy>
...
</Patch>

```

```

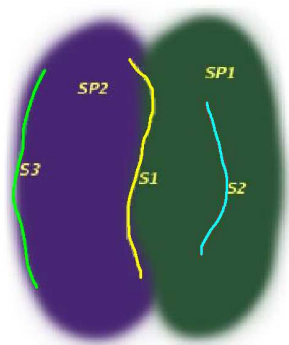
<owl:Class rdf:ID="Orientation">
  <rdfs:subClassOf>
    <owl:Class rdf:ID="Posterior"/>
  </rdfs:subClassOf>
</owl:Class>

<owl:ObjectProperty rdf:ID="isMAEBoundedBy">
  <inverseOf>
    <owl:ObjectProperty rdf:about="#MAEBounds">
      </inverseOf>
    </owl:ObjectProperty>
  </inverseOf>
</owl:ObjectProperty>

```

FIG. 3 – Exemple de faits initiaux de la Abox en OWL

tion en chaque point de ce volume renseigne sur la probabilité d'appartenance à la structure anatomique en question.



```

<SulcusSegment rdf:ID="183">
  <separates>
    <MAEPair rdf:ID="MAEPair41">
      <firstEntity>
        <AttributedEntity rdf:ID="entity41">
          <entity>
            <RightPreCentralGyrus rdf:ID="rPrCGyrus"/>
          </entity>
          <orientation>
            <Right rdf:ID="rightTo"/>
          </orientation>
          <orientation>
            <Posterior rdf:ID="posteriorTo"/>
          </orientation>
          <orientation>
            <Inferior rdf:ID="inferiorTo"/>
          </orientation>
        </AttributedEntity>
      </firstEntity>
      <secondEntity>
        ...
      </SulcusSegment>
    </MAEPair>
  </separates>
</SulcusSegment>

```

FIG. 4 – Les SPAMs

Nous utilisons les SPAMs comme suit (nous avons seulement considéré les SPAMs associées aux gyri les plus importants) : pour chaque trace externe (une trace externe est composée d'un ensemble de points) on transforme d'abord tous les points qui la composent dans l'espace stéréotaxique ; ceci est facile car la segmentation du cerveau effectuée précédemment avait nécessité le recalage dans l'espace stéréotaxique, de façon à permettre l'utilisation d'un masque du cerveau défini dans cet espace. On calcule ensuite la moyenne des probabilités d'appartenance de tous ces points à chaque SPAM, cette moyenne renseigne sur la position de la trace externe vis à vis de la SPAM. A partir de cette valeur on utilise quelques heuristiques : si la

SulcusSegment \ Gyri	Right-Precentral-Gyrus	Right-Postcentral-Gyrus	Right-Angular-Gyrus	RightSuperiorTemporalGyrus	...
ID = 183	0.384	0.186	0	0	.
ID = 178	0.218	0	0	0	.
ID = 155	0	0.477	0.105	0	.
ID = 298	0	0.038	0	0.076	.

FIG. 5 – Exemple de tableau de probabilités

probabilité dépasse un certain seuil, alors on considère que la trace externe est située à l'intérieur de la SPAM (cas du segment S_2 sur la figure 4). Si la probabilité est proche de zéro mais non nulle, alors on considère que le segment borde la SPAM (cas de S_1 et S_3).

Les heuristiques sont résumées comme suit : Pour chaque trace externe s_i on conserve ses deux probabilités les plus élevées p_{i1} et p_{i2} (supposons $p_{i1} > p_{i2}$) d'appartenir aux deux SPAMs correspondantes sp_1 et sp_2 . Soient MIN et MAX deux valeurs réelles qui sont déterminées de manière empirique.

- Si $(p_{i1} \in [MIN, MAX] \text{ et } p_{i2} \in [MIN, MAX]) \implies s_i \text{ separates object}(sp_1 : instance, sp_2 : instance)$, ici s_i est localisée entre sp_1 et sp_2 (cas de S_1 sur la figure 4).
- Si $(p_{i1} \in [MIN, MAX] \text{ et } p_{i2} < MIN) \implies s_i \text{ MAEBounds } sp_1 : instance$, i.e. elle borde l'une des deux SPAM cas de S_3 sur la figure 4)
- Si $(p_{i1} > MAX \text{ et } p_{i2} < MIN) \implies s_i \text{ isIn } sp_1 : instance$, i.e. elle est à l'intérieur de l'une des deux SPAM (cas de S_2 sur la figure 4)

Les résultats ainsi obtenus à partir de ces valeurs numériques fournissent des relations symboliques qui sont représentées dans un fichier OWL DL.

Pour déterminer les orientations des traces externes et des SPAMs les unes par rapport aux autres, par exemple dire que s_i borde sp_j antérieurement, on compare les coordonnées (x, y, z) du centre de la trace transformées dans l'espace stéréotaxique, aux coordonnées (x', y', z') du centre de la SPAM.

- Si $(x > x') \implies \text{Right}(s_i, sp_j) \text{ sinon } \text{Left}(s_i, sp_j)$
- Si $(y > y') \implies \text{Anterior}(s_i, sp_j) \text{ sinon } \text{Posterior}(s_i, sp_j)$
- Si $(z > z') \implies \text{Superior}(s_i, sp_j) \text{ sinon } \text{Inferior}(s_i, sp_j)$

Ainsi chaque entité a trois orientations : (Right ou Left), (Posterior ou Anterior), et (Superior ou Inferior). Tous les faits sont alors représentés en OWL DL comme montré sur la figure 3. Le centre d'une SPAM peut être un nuage de points dont l'intensité est la même et la plus élevée de la SPAM en question, par conséquent nous calculons l'orientation du centre de la trace externe par rapport au point le plus proche de ce nuage de points.

2.3 Ontologie et règles d'inférence

Les connaissances symboliques a priori sont représentées dans une ontologie qu'on appellera ontologie des structures corticales. Pour la représentation on utilise des langages standards du Web basés sur la logique du premier ordre qui permettent le partage et qui gardent les

Annotation Sémantique d'Images IRM du Cerveau

bonnes propriétés de complexité et décidabilité du raisonnement, à savoir OWL DL pour la représentation de l'ontologie, et SWRL pour les règles. Les noms des gyri utilisés dans l'ontologie coïncident avec ceux utilisés pour dénommer les SPAMs correspondantes. Afin d'enrichir l'ontologie, on l'étend avec des règles de Horn sans fonction. Les règles permettent de propager des relations et d'inférer de nouveaux faits à partir de ceux existant. L'ontologie et les règles ont été éditées avec Protégé Protégé (2006) (figure 6). Pour la conception de l'ontologie nous avons utilisé des atlas d'anatomie Ono et al. (1990) et d'autres travaux sur l'anatomie du cerveau humain Dameron (2003).

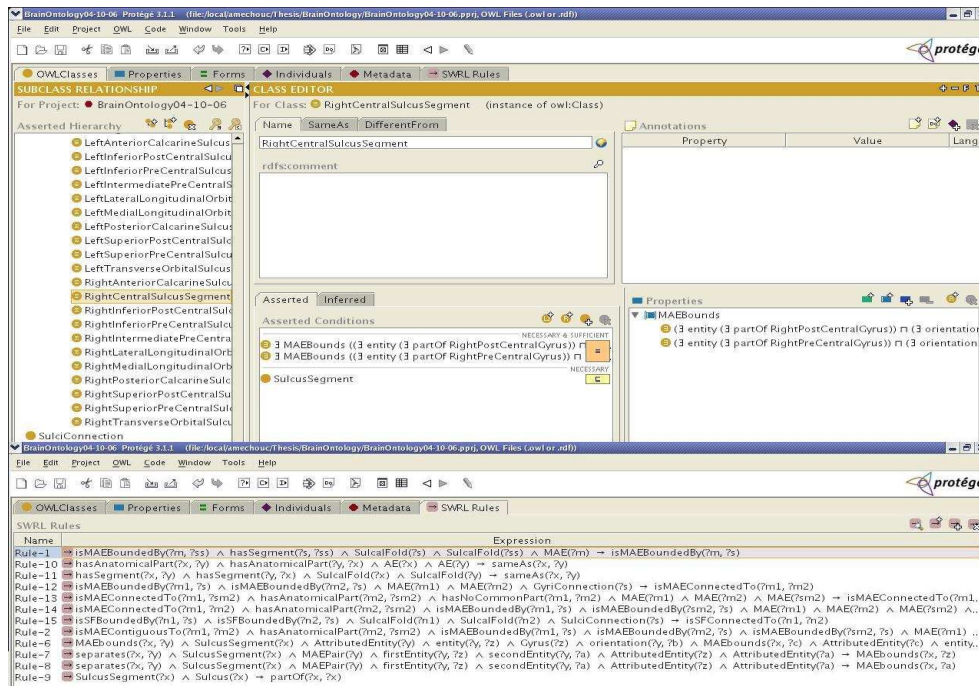


FIG. 6 – L'ontologie étendue avec les règles

La base de connaissance considérée globalement comprend trois parties :

- **TBox $\mathcal{O}$** : Elle fournit les définitions logiques des concepts (classes), des relations (propriétés), et des axiomes. Par exemple, les conditions nécessaires et suffisantes pour définir un segment du sillon central de l'hémisphère droit¹ sont : $RightCentralSulcusSegment \equiv ((\exists MAEBounds ((\exists entity (\exists partOf RightPostCentralGyrus)) \sqcap (\exists orientation Anterior))) \sqcup (\exists MAEBounds ((\exists entity RightPostCentralGyrus) \sqcap (\exists orientation Anterior)))) \sqcap ((\exists MAEBounds ((\exists entity (\exists partOf RightPreCentralGyrus)) \sqcap (\exists orientation Posterior))) \sqcup (\exists MAEBounds ((\exists entity RightPreCentralGyrus) \sqcap (\exists orientation Posterior))))$. En effet, un segment du sillon central est bordé antérieurement par le gyrus pré-central

¹Cette définition ne peut être considérée comme une définition universelle de $RightCentralSulcusSegment$, mais répond aux besoins de cette application qui focalise l'attention sur les relations topologiques

- ou une partie du gyrus pré-central, et postérieurement par le gyrus post-central ou une partie du gyrus post-central. Il faut noter que $MAEBounds = isMAEBoundedBy^{-1}$.
- **RBox $\mathcal{R}$** : Elle fournit des règles étendant l'ontologie pour augmenter son expressivité, par exemple la règle Rule-1 sur la figure 6 exprime la propagation de la bordure des parties au tout : $isMAEBoundedBy(?x, ?y) \wedge hasSegment(?z, ?y) \wedge SulcalFold(?z) \wedge SulcalFold(?y) \wedge MAE(?x) \rightarrow isMAEBoundedBy(?x, ?z)$. Ainsi, si on a une entité matérielle anatomique x qui est bordée par le sulcal fold y , et que y est une partie d'un autre sulcal fold z , alors on déduit que x est aussi bordé par z . Les règles d'inférence sont également utiles pour exprimer des requêtes. Par exemple, exprimer une requête qui recherche pour tous les patches p_i de la région d'intérêt étudiée, les instances de gyri dont ils font partie : $Q(?x_i, ..., ?x_n) \leftarrow \bigwedge_{i=1 \text{ to } n} (AE(?x_i)) \wedge hasPart(?x_i, p_i)$. L'extension des ontologies impose des restrictions Rosati (2005) à des fragments décidables communs aux DLs (logique des description) et les LPs (programmes logiques).
 - **Abox $\mathcal{A}$** : Elle contient les individus (instances de classes) et les instances des relations qui les relient. Par exemple $\langle Patch : P1 \rangle$ où $P1$ est une instance du concept Patch.

3 Le processus de raisonnement et d'étiquetage

Le processus d'étiquetage est organisé comme suit : on étiquette d'abord les patches en utilisant la règle suivante (on l'appellera $\mathcal{R}_{\mathcal{L}}$) :

- $MAEBounds(?x, ?y) \wedge SulcusSegment(?x) \wedge AttributedEntity(?y) \wedge entity(?y, ?z) \wedge Gyrus(?z) \wedge orientation(?y, ?b) \wedge MAEBounds(?x, ?c) \wedge AttributedEntity(?c) \wedge entity(?c, ?d) \wedge Patch(?d) \wedge orientation(?c, ?b) \wedge Orientation(?b) \rightarrow partOf(?d, ?z)$

En réalité, ce n'est pas exactement cette règle qui est utilisée. Etant donné que les faits proviennent de différentes sources, nous avons introduit le constructeur *sameAs* pour exprimer la même orientation. La règle utilisée est la suivante :

- $MAEBounds(?x, ?y) \wedge SulcusSegment(?x) \wedge AttributedEntity(?y) \wedge entity(?y, ?z) \wedge Gyrus(?z) \wedge orientation(?y, ?b) \wedge MAEBounds(?x, ?c) \wedge AttributedEntity(?c) \wedge entity(?c, ?d) \wedge Patch(?d) \wedge orientation(?c, ?bb) \wedge Orientation(?b) \wedge Orientation(?bb) \wedge sameAs(?b, ?bb) \rightarrow partOf(?d, ?z)$

Cette règle calcule la similarité entre les faits extraits directement de l'image, et ceux extraits des SPAMs ; ainsi, si on a un sulcus segment x qui borde un gyrus z avec une orientation b , et qu'en même temps il borde un patch d avec une même orientation b , alors on en déduit que le patch d fait partie du gyrus z . Evidemment plusieurs sulcus segments peuvent inférer la même chose pour le même patch. En revanche, même si c'est très rare, un même patch peut être affecté à plusieurs gyri. Ceci peut être dû à des incertitudes dans le calcul des orientations. Dans ce cas, pour décider à quel gyrus affecter le patch p on calcule pour chaque gyrus la $\sum p_i$ où les p_i sont les probabilités d'appartenance de p à ce gyrus, inférées par plusieurs sulcus segments. Ensuite on décide alors d'affecter le patch p au gyrus pour lequel la $\sum p_i$ est la plus élevée.

Les deux règles suivantes infèrent la bordure à partir de la séparation.

- $separates(?x, ?y) \wedge SulcusSegment(?x) \wedge MAEPair(?y) \wedge firstEntity(?y, ?z) \wedge secondEntity(?y, ?a) \wedge AttributedEntity(?z) \wedge AttributedEntity(?a) \rightarrow MAEBounds(?x, ?z)$
- $separates(?x, ?y) \wedge SulcusSegment(?x) \wedge MAEPair(?y) \wedge firstEntity(?y, ?z) \wedge secondEntity(?y, ?a) \wedge AttributedEntity(?z) \wedge AttributedEntity(?a) \rightarrow MAEBounds(?x, ?a)$

Il est clair que si un sulcus segment x sépare deux entités z et a , alors il les borde toutes les deux.

Après l'étiquetage des patches, on étiquette ensuite les sulci. Il devient possible d'étiqueter les segments en utilisant les définitions des sillons dans l'ontologie. Dès qu'un segment satisfait toutes les conditions nécessaires et suffisantes d'un sillon dans l'ontologie, il est automatiquement identifié comme instance de ce dernier.

La figure 7 résume le déroulement du processus d'étiquetage : les numéros figurant entre parenthèses dans le paragraphe suivant font références aux différentes étapes représentées sur la figure.

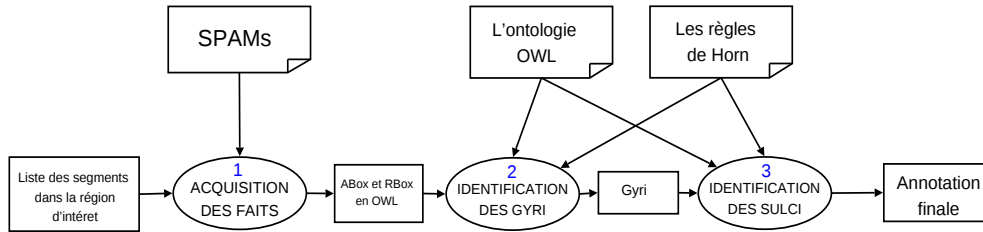


FIG. 7 – Processus d'étiquetage

(1) A partir de la liste des traces externes (segments ou sillons non nommés) et de la liste des SPAMs, on calcule la probabilité d'appartenance de chaque segment à chaque SPAM, on obtient ainsi un tableau comme celui de la figure 5.

Ce tableau est d'abord créé sous la forme d'un fichier XML, puis transformé en un fichier OWL. Ce dernier est alors fusionné avec l'ontologie OWL DL et les règles SWRL ainsi que les faits extraits des images et ceux extraits des SPAMs. A ce fichier se rajoute la règle d'inférence $\mathcal{R}_L$ qui est d'une importance majeure dans notre méthode comme expliqué précédemment. (2) Après avoir étiqueté les patches, (3) le raisonneur pourra utiliser les définitions des sillons décrites dans l'ontologie pour étiqueter les segments. Ainsi, tous les segments vérifiant les conditions nécessaires et suffisantes définies pour un sillon dans l'ontologie sont automatiquement identifiés comme instances de ce sillon. Le résultat final est constitué par l'ensemble des structures étiquetées dans la région d'intérêt.

Exemple : Supposons que l'on a un segment s_0 , et deux patches p_1 et p_2 . Soit $\text{separates}(s_0, ((\text{Patch} : p_1, \{\text{Anterior, Right, Superior}\}), (\text{Patch} : p_2, \{\text{Posterior, Left, Inferior}\})))$ un fait extrait de l'image par les outils numériques, et $\text{separates}(s_0, ((\text{RightPreCentralGyrus} : \text{pcgr}, \{\text{Anterior, Right, Superior}\}), (\text{RightPostCentralGyrus} : \text{pcgr}, \{\text{Posterior, Left, Inferior}\})))$ un fait extrait des SPAMs. En posant la requête introduite précédemment dans la section 2.3, le raisonneur, après application des règles, en particulier $\mathcal{R}_L$, va inférer : $\text{partOf}(p_1, \text{pcgr})$ et $\text{partOf}(p_2, \text{pcgr})$. Enfin, si on considère la définition de $\text{RightCentralSulcusSegment}$ donnée précédemment, alors le raisonneur pourra inférer que s_0 est une instance de $\text{RightCentralSulcusSegment}$.

4 Résultats expérimentaux

Nous avons effectué des tests sur des données cérébrales réelles. Nous avons choisi d'utiliser comme raisonneur KAON2 Kaon2 (2005) étant donné qu'il peut raisonner sur une ontologie étendue avec des règles ?. La figure 8 montre le résultat de l'étiquetage des patches de la région d'intérêt extraite de la figure 2.

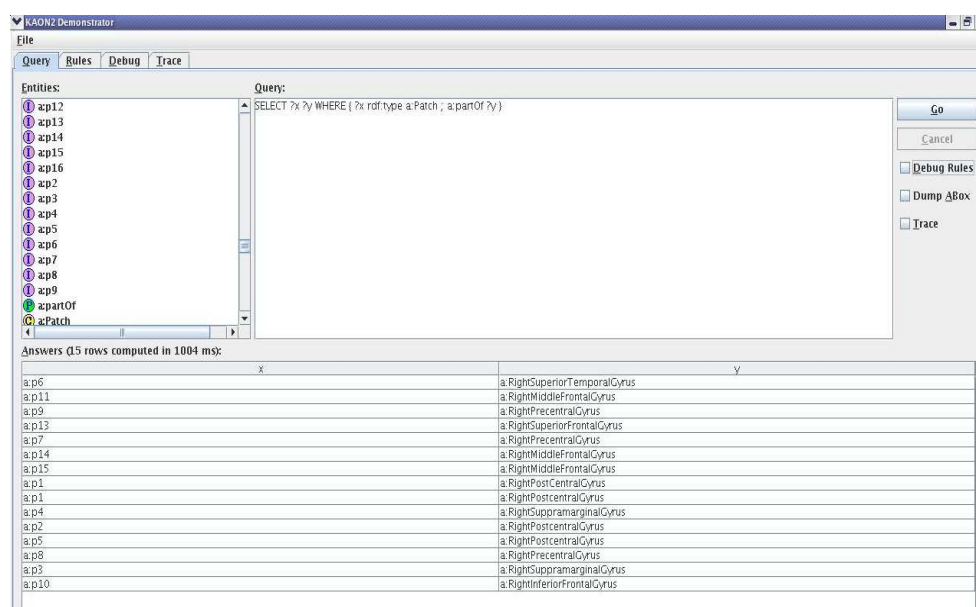


FIG. 8 – Etiquetage des patch avec KAON2

Ainsi, à partir de la définition du sillon central figurant dans l'ontologie, on peut savoir quels sillons en font partie. La figure 9 montre le résultat de la requête pour obtenir les parties du sillon central. Le résultat est le segment portant l'identifiant 183 qui est effectivement le sillon central.

Les tests n'ont pas été effectués sur une grande population de sujets en raison de la non automatisation complète du processus d'extraction des faits des images. Il est évident que la validation effective de cette approche passera par des tests sur plusieurs sujets (sains et pathologiques). Un cas pathologique dont nous disposons contient une lésion qui n'a pas provoqué de décalage des structures anatomiques. Les sujets présentant des lésions qui entraînent des décalages sur le cortex seront les plus intéressants dans les tests.

5 Discussion

L'originalité de la méthode proposée tient à plusieurs facteurs. Tout d'abord, elle tente d'utiliser de façon cohérente différentes formes d'a priori, des a priori de nature statistique

Annotation Sémantique d'Images IRM du Cerveau

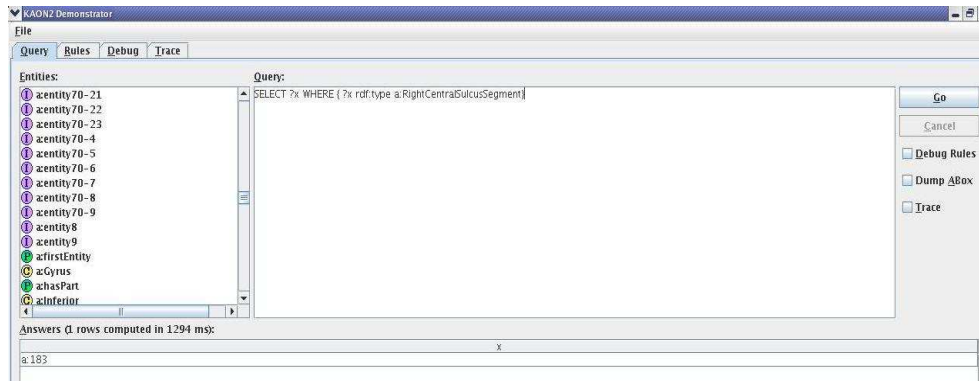


FIG. 9 – *Etiquetage des sillons avec KAON2*

comme des cartes statistiques 3D (SPAMs) des structures anatomiques, et des a priori de nature topologique - représentés sous forme symbolique au sein de l'ontologie des structures anatomiques. Les premiers fournissent des éléments d'identification de ces structures essentiellement fondés sur leur position par rapport à un espace de référence (en l'occurrence l'espace stéréotaxique). Ces éléments pourront donc éventuellement s'avérer faux en présence de grosses lésions susceptibles de déplacer localement certaines structures. En revanche ces déplacements seront toujours locaux, ce qui signifie que les informations fournies seront tout de même majoritairement valides. Le second type d'a priori, de nature symbolique, porte sur la topologie des structures et présente l'avantage d'être insensible à la présence de lésions. C'est précisément cette complémentarité que nous souhaitons exploiter, même si nous n'avons pas aujourd'hui de solution automatique pour la remise en cause de certains faits. Un second élément d'originalité concerne les connaissances représentées dans l'ontologie, et particulièrement le fait d'associer une ontologie fondée sur les logiques de description avec des règles. La motivation tient à la nécessité de pouvoir raisonner sur des parties des entités anatomiques - segments de sillons, patches qui sont des parties de gyrus - pour pouvoir les rattacher à des tout identifiables. En effet, il existe un large fossé entre les primitives qu'on peut extraire des images de façon automatique (par exemple des segments de sillons) et les structures anatomiques telles que les ont décrites les anatomistes, qui ont un caractère beaucoup plus macroscopique. C'est précisément le rôle de ces règles (ou de certaines d'entr'elles) que de modéliser les dépendances entre les relations partie-tout et les relations topologiques qui concernent les gyri et les sillons. Enfin un troisième niveau d'originalité concerne le côté général et donc ré-utilisable de l'ontologie. Notre approche a pour ambition de représenter les connaissances anatomiques de façon non-spécifique du problème à résoudre. Ceci suppose bien entendu des compromis. Ceci est particulièrement flagrant dans la définition des conditions nécessaires et suffisantes associées à chaque classe d'entité. Il est clair que les définitions que nous retenons ici sont principalement guidées par les besoins de l'application, qui, en l'occurrence, met fortement l'accent sur les aspects topologiques. Ceci pose la question fondamentale de la faisabilité de telles ontologies à caractère universel et des traitements à leur appliquer pour en dériver des ontologies applicatives, adaptées aux besoins spécifiques de telle ou telle application.

Notre souhait de proposer une ontologie ré-utilisable nous a conduits naturellement à choisir d'utiliser OWL. Le souhait de pouvoir bénéficier de bonnes propriétés au niveau calculatoire et utiliser les raisonneurs existants nous a conduits à choisir OWL DL et, pour ce qui concerne les règles, à nous limiter à des règles de Horn sans fonction. Ces choix nous ont amenés à faire plusieurs compromis, et contourner plusieurs limitations liées aux outils existants, par exemple :

- L'ontologie est représentée en OWL DL au lieu de OWL1.1 OWL1.1 (2006) et utilise des restrictions existentielles au lieu de restrictions de cardinalité quantifiées.
- KAON2 ne supportant pas encore le type nominal, nous avons défini des sous classes de Orientation, ex. ; Posterior, Anterior etc. avec des individus e.g., posteriorTo et utilisé des restrictions existentielles
- Toutes les relations n-aires sont transformées en relations binaires, par exemple nous avons défini une classe artificielle AttributedEntity qui a deux champs : une structure anatomique, et une collection de trois orientations pour représenter des relations attribuées comme dans FMA Golbreich et al. (2006).
- Les règles d'heuristiques en section 2.2 sont implémentées en C++ car KAON2 n'accepte pas encore les datatype.
- L'interface actuelle ne permet pas une bonne interaction avec l'utilisateur, elle peut être améliorée en rendant la méthode plus automatique et en offrant à l'utilisateur une interface conviviale où il agit seulement sur l'image et lance le service de raisonnement pour l'étiquetage (s'exécutant en arrière plan) d'une manière simple. La remise en cause des annotations sémantiques des images ne peut pas être faite d'une manière complètement automatique car on ne peut que visualiser les résultats d'inférence de KAON2 sans pouvoir les récupérer par des programmes.
- Enfin, les règles que nous avons définies sont des règles SWRL que KAON2 transforme en interne en règles DL-safe Motik et al. (2004). Pour cette raison, certaines solutions peuvent être manquées Golbreich (2006).

6 Conclusion

Cet article rapporte l'état actuel du développement d'un système hybride combinant des techniques symboliques et numériques pour la description des images IRM du cerveau. La méthode est basée sur une ontologie OWL DL étendue par des règles, et des faits provenant des outils numériques et de connaissances a priori sous la forme de cartes statistiques (SPAMs). Le système requiert un minimum d'interaction avec l'utilisateur. Jusqu'à présent, la méthode a été testée sur un ensemble limité d'images cérébrales, et qui ne présentaient pas de lésion. Il est important d'étendre les expérimentations à un plus grand nombre de cas, et notamment à des cas présentant une lésion, afin de tester sa robustesse, et sa valeur ajoutée par rapport à d'autres approches. Il est également nécessaire d'avoir un système plus automatique qui réduit les interactions avec l'utilisateur. Les travaux futurs consisteront également à voir comment surmonter les limitations rencontrées au niveau de la représentation des connaissances. Au delà de l'aide directe qu'elle peut fournir à un médecin dans le cadre - par exemple - du choix d'une stratégie opératoire, l'automatisation de l'annotation du contenu sémantique des images numériques présente des perspectives prometteuses pour la recherche de cas similaires pour l'aide à la décision, ou des études médicales statistiques sur les corrélations anatomie fonction

ou anatomie métabolisme au niveau du système nerveux central chez de grandes populations d'individus.

Références

- Collins, D. L., A. P. Zijdenbos, W. F. C. Baaré, et A. C. Evans (1999). INSECT: Improved cortical structure segmentation. *Proc. of the Annual Symposium on Information Processing in Medical Imaging IPMI99 LNCS 1613*, 210–223.
- Dameron, O. (2003). *Modélisation, représentation et partage de connaissances anatomiques sur le cortex cérébral*. Thèse de doctorat, Université de Rennes1.
- Golbreich, C. (2006). *RoW Workshop at WWW2006 1*, Web rules for Health Care and Life Sciences: use cases and requirements.
- Golbreich, C., S. Zhang, et O. Bodenreider (2006). The Foundational Model of Anatomy in OWL: Experience and Perspectives. *Journal of Web Semantics*, (4)3.
- Kaon2 (2005). KAON2, <http://kaon2.semanticweb.org/>.
- Motik, B., U. Sattler, et R. Studer (2004). Query answering for OWL DL with rules. *ISWC 2004, LNCS 3298 Springer*.
- Ono, M., S. Kubik, et C. Abernathy (1990). *Atlas of the Cerebral Sulci*. Thieme Medical Publishers, Inc.
- OWL (2004). OWL Web Ontology Language : Overview, <http://www.w3.org/tr/owl-features/>.
- OWL1.1 (2006). OWL 1.1 Web Ontology Language Syntax: <http://www-db.research.bell-labs.com/user/pfps/owl/syntax.html>.
- Protégé (2006). Protégé, <http://protege.stanford.edu>.
- Rosati, R. (2005). On the decidability and complexity of integrating ontologies and rules. *Journal of Web Semantics*.

7 Remerciements

Nous remercions Louis Collins de l'Institut Neurologique de Montréal qui nous a fourni les données SPAM, ainsi que le Conseil Régional de la Bretagne pour son soutien à ce travail.

Summary

This paper presents how an hybrid system combining symbolic and numerical techniques can be used for the description of brain MRI images. We show how reasoning supported by an OWL DL ontology extended with SWRL rules, and facts coming from numerical segmentation and statistical anatomical priors, could be efficient for assisting the labelling of some brain structures in MRI images.