



Distinguer/Expliciter. L'ontologie du Web comme ontologie "d'opérations"

Alexandre Monnin, Pierre Livet

► To cite this version:

Alexandre Monnin, Pierre Livet. Distinguer/Expliciter. L'ontologie du Web comme ontologie "d'opérations". *Intellectica - La revue de l'Association pour la Recherche sur les sciences de la Cognition (ARCo)*, 2014, Philosophie du Web et Ingénierie des connaissances, 1 (61), pp.59-104. 10.3406/intel.2014.1039 . hal-01079854v2

HAL Id: hal-01079854

<https://hal.science/hal-01079854v2>

Submitted on 14 Sep 2019

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.

Distinguer/Expliciter. L'ontologie du Web comme ontologie « d'opérations »

Alexandre MONNIN* & Pierre LIVET#

RÉSUMÉ. Prenant au sérieux la « métaphysique empirique » de l'architecture du Web, nous entendons saisir les entités qu'elle distingue. Il ne s'agit donc pas de reconduire ainsi des distinctions classiques. Le travail réalisé par les architectes du Web ressortit à l'ingénierie. Pour cette raison, nous posons les linéaments d'une approche processuelle de l'ontologie, intégrant en son sein les opérations de distinction et d'explicitation qui conduisent à faire émerger des types ontologiques. Cette réflexion est rapportée à plusieurs exemples abordés au cours de la première partie, des folksonomies aux ontologies informatiques, en passant par les principes architecturaux du Web. Enfin, nous prenons en compte la pluralité des modalités d'accès afin d'étendre à l'échelle du Web tout entier la réflexion sur la dynamique ontologique au cours de laquelle surgissent de nouvelles distinctions. Un pour rendre compte de la sémantique des URI est suggéré qui s'appuie sur la sémantique bi-dimensionnelle de Robert Stalnaker, étendue à la multi-dimensionnalité et repensée dans une optique dégagée de son ancrage exclusif dans la question du vrai.

Mots-clés : Philosophie du Web, ontologie, ontologies informatiques, Web, Web Sémantique, architecture du Web, URI, processus, opérations, distinction, explicitation.

ABSTRACT. Distinguishing/Making Explicit. Web Ontology as Ontology of Operations. By taking seriously the “empirical metaphysics” contained in the architecture of the Web, our goal is to grasp the kind of entities it came to distinguish. Yet, such distinctions are not understood as classical philosophical distinctions. Indeed, the work accomplished by the Web architects belongs to engineering. For that reason, we outline a processually-based approach to ontology. It embeds the very operations of distinction and explicitation through which ontological types are made to emerge. We then relate this first sketch to a variety of examples picked up from the first part, dealing with folksonomies, computer ontologies and the very principles behind the architecture of the Web. Therefrom, we take into account the plurality of the modes of access on the Web on a broader scale in order to extend our previous developments on the ontological dynamics through which new distinctions are made explicit. A model is laid down which draws from Robert Stalnaker's two-dimensional semantics, generalized on a multi-dimensional level and reshaped so as to escape the exclusive focus on the notion of truth.

Keywords: Philosophy of the Web, ontologies, computer ontologies, Web, Semantic Web, architecture of the Web, URI, process, operations, distinction, explicitation.

* WIMMICS (INRIA Sophia Antipolis / Laboratoire I3S) ; INRIA – Université Nice Sophia Antipolis (UNS) – CNRS : UMR7271. alexandre.monnin[at]web-and-philosophy.org.

Université d'Aix-Marseille. Pierre.Livet[at]univ-amu.fr.

« *What would it be for the very notion of distinction to be won, at a price, from a partially regular, partially turbulent, noisy and critical background – rather than for a formally first-order critical region to be defined on top, or hung from, a perfectly structured infinite silence?* »

(Smith, 1998, p. 333)

I – INTRODUCTION

La question ontologique a connu un renouveau important depuis une trentaine d'années, au sein de champs d'études distincts de la philosophie : Intelligence Artificielle, Ingénierie des Connaissances ou encore Web Sémantique. S'il y a indubitablement matière à se pencher à nouveaux frais sur cette question à partir des ontologies informatiques, nous avons cependant choisi, dans les développements qui suivent, de prendre nos distances vis-à-vis de ces débats afin de répondre à une autre question, concernant la nature de *l'ontologie du Web*¹. Avec la notion de ressource, c'est en effet la question de l'objet que le Web repose à nouveaux frais. L'enjeu porte donc sur la nature et les conditions d'émergence des types de cette « ontologie » – au sens où ce mot désigne avant tout une « théorie de l'objet ». Exigence qui conditionne d'éventuelles prises de positions ultérieures dans les débats tantôt mentionnés sur les ontologies. Passer outre cette étape préliminaire condamne d'ailleurs à ressasser implacablement les mêmes antagonismes². En prenant au sérieux la « métaphysique empirique » inhérente à l'architecture du Web, nous entendons donc saisir les entités qu'elle distingue.

C'est pourquoi la première partie de cet article retrace en ouverture les étapes au fil desquelles ont émergé les principes de l'architecture du Web. Préalable indispensable en vue de rapporter les évolutions ultérieures du Web 2.0 et du Web Sémantique à leur milieu technique d'origine et, partant, aux contraintes – techniques, conceptuelles, ontologiques – dont il s'avère porteur. Pour autant, il ne s'agit en aucun cas de reconduire ainsi des distinctions classiquement analytiques ou synthétiques *a priori*. Le travail réalisé par les architectes du Web ressortit à *l'ingénierie*. Pour cette raison, nous devons penser et thématiser dans notre enquête le *processus* même qui conduit à ces distinctions.

C'est là l'objet de la seconde partie, où sont détaillés, à rebours des approches ontologiques traditionnelles (en philosophie comme en informatique), que nous qualifions de « fréguennes », les linéaments d'une approche processuelle de l'ontologie, non plus dualiste mais triadique, intégrant en son sein les opérations de *distinctions* et d'*explicitation* qui conduisent à l'émergence de ses propres types. La réflexion menée est rapportée à plusieurs exemples, en guise de mise à l'épreuve, des folksonomies³ aux ontologies

¹ Cf. (Monnin, 2014).

² En guise d'illustration d'un débat engagé sur de mauvais rails, en ce qu'il n'est jamais parvenu à s'émanciper de positions antagonistes qui en constituent le point de départ mais également l'unique horizon, mentionnons les échanges entre Martin Mongin et Frédéric Nef au sujet des ontologies informatiques et de leur évaluation d'un point de vue philosophique (Mongin, 2006 ; Nef, 2007 ; Mongin, 2007).

³ Le mot « folksonomie », contraction de « *folk taxonomy* », a été forgé en 2004 par Thomas Vander Wal, architecte américain de l'information, pour désigner la mise en commun des fruits de l'activité

informatiques en passant par les principes architecturaux du Web, déjà abordés au cours de la première partie.

Enfin, dans un troisième temps, nous prenons en compte la pluralité des modalités d'accès, caractéristique du Web et pourtant sciemment laissée de côté dans la perspective du Web Sémantique, pour interroger la très grande variabilité des contenus associés à une ressource en vertu des principes établis précédemment. Ceci, afin d'étendre à l'échelle du Web tout entier la dynamique ontologique au cours de laquelle surgissent de nouvelles distinctions et, avec elles, de nouveaux types et de nouvelles entités. En parallèle, nous proposons un modèle destiné à rendre compte de la sémantique des URI dans une perspective compatible avec l'ontologie esquissée jusqu'à présent. Inspiré de la sémantique bi-dimensionnelle de Robert Stalnaker, il l'élève à la multi-dimensionnalité dans une optique dégagée de son ancrage dans la question du vrai, de manière à faire droit aux exigences spécifiques du Web.

II – L'ARCHITECTURE DU WEB AU PRISME DE SES DISTINCTIONS

Pour être en mesure de bien cerner les spécificités architecturales du Web, il convient dans un premier temps de déterminer avec précision le statut des entités – adresses, identifiants, etc. – à partir desquelles se conçoivent les évolutions qui ont conduit au Web social et au Web Sémantique. De ce point de vue, il ne s'agit surtout pas de restreindre la question de l'ontologie du Web aux seules ontologies du Web Sémantique. Pas plus qu'on ne saurait d'ailleurs réduire ce dernier à une simple extension des travaux menés en Intelligence Artificielle ou dans le domaine de l'ingénierie des connaissances. Une fois acclimatées au Web, les ontologies ne peuvent en effet plus s'envisager dans le simple prolongement théorique de disciplines qui ignoreraient tout de ce dernier. Nulle *tabula rasa* n'est ici de mise du fait des contraintes structurelles issues de son architecture. Ajoutons en outre que la manière dont les entités du Web ont été dégagées au fil du temps illustre l'idée développée dans cet article selon laquelle les types d'une ontologie sont le fruit d'opération de distinctions dont le résultat n'est pas donné à l'avance. Pour illustrer ce point, nous revenons dans ce qui suit sur l'évolution des différents standards consacrés aux identifiants du Web, culminant avec le style d'architecture REST qui définit la ressource comme le contrepoint desdits identifiants. Enfin, nous éclairons le statut des tags à l'aune des distinctions architecturales ébauchées précédemment.

II.1 - L'indétermination des distinctions en jeu : le Web entre UDI, URI, URL, URN, IRL et URC

Il est impossible ici de rappeler dans le détail l'histoire particulièrement riche des identifiants du Web. Considérés comme la brique de standards la plus

individuelle déployée sur les sites de tagging « collaboratifs » alors naissants, en particulier Flickr (autour du partage d'images) et Del.icio.us (autour du partage de « signets » – *social bookmarking*). Initialement considérées comme des alternatives aux classifications bibliothéconomiques comme aux structururations hiérarchiques de types taxonomiques ou logiques (les ontologies), les folksonomies ont néanmoins très vite fournit la matière à des couplages entre formes de conceptualisations distribuées et ascendantes (*bottom-up*), et formes plus structurées et descendantes (*top-down*). Sur cette question voir (Monnin, 2013b, IVe partie).

importante, c'est essentiellement par leur truchement que se pose la question de la nature de leurs corrélats (pages, documents, fichiers, ressources, etc.), auxquels nulle spécification n'est consacrée. En vérité, la nature de ces identifiants conditionne la manière dont le Web fait système : s'agit-il d'un hypertexte, d'un système de gestion de fichiers, ou tout autre chose encore ?

À l'origine, la couche de nommage du Web revêtait une importance toute particulière dans un contexte où celui-ci n'avait pas encore acquis sa prééminence actuelle vis-à-vis d'autres systèmes tels que Gopher, WAIS et autres Prospero. L'enjeu était alors de définir une syntaxe à même de permettre à l'ensemble de ces systèmes de communiquer par l'intermédiaire de leurs identifiants.

On peut faire remonter à 1992 les premières réflexions publiées à ce sujet par Tim Berners-Lee et ses collaborateurs. D'emblée, une préférence claire est accordée au choix d'un « nom logique », baptisé UDI⁴ (*Universal Document Identifier*) au détriment des adresses physiques. En cas de déplacement, l'ancienne adresse d'un objet n'est en effet d'aucune utilité pour le localiser car, nom « autonome » s'il en est, elle demeure comme « assignée à résidence » lors même que l'objet a pris son envol. Forger un nom, ayant une fonction explicitement distinguée d'une adresse, constitue dès lors la première étape pour établir des redirections menant au nouvel emplacement d'un objet. À côté du nommage pérenne, se dessine, on le note également, une forme d'adressage ayant pour but de pallier les déplacements physiques. Plus largement, le caractère pérenne du nommage lui-même demande à être assuré par une mise à l'écart préalable, dans le choix d'un nom logique, des détails qui l'attachent potentiellement à un « document » (format, longueur, etc.), le lestant d'une temporalité marquée du sceau de l'éphémère. Corrélativement, s'il y a bien des documents, ceux-ci sont d'abord saisis comme objets du réseau, objets d'une recherche ou d'une future consultation. Mais cette acception s'inscrit dans des bornes étroites puisqu'il est immédiatement précisé que : « *The “document” is the unit of retrieval and need not correspond to any unit of storage [...] We emphasize that this is the abstract view of the users, and these objects need not correspond to physical files on computers.* ». En d'autres termes, le contenu consulté est d'emblée découplé d'une unité de stockage « correspondante ». D'ailleurs, dès cette époque, rien n'interdit de mettre à profit les UDI pour identifier des services interactifs, autrement dit, moins des unités clairement circonscrites qu'une série de réponses variables générées en réponse à une requête ponctuelle.

L'année 1994 coïncide avec la naissance du W3C. À sa suite, une première vague de standardisation fut lancée. C'est de cette époque que remonte l'acronyme URI⁵ même s'il apparaît initialement dans quelques publications qui ne sont pas encore des standards de plein droit⁶. Par rapport aux UDI, le « D » de « document » a laissé la place au « R » de « ressource ». De simple note décrivant un rêve et s'achevant (de manière appropriée) sur un vibrant « *the dream is coming true* », nous sommes passés à l'ébauche d'une recommandation, toujours

⁴ (Berners-Lee *et al.*, 1992).

⁵ D'abord Universal Resource Identifier, puis Uniform Resource Identifier.

⁶ (Berners-Lee, 1994).

consultable sur le site de l'IETF⁷. Y sont également mentionnées les URL, pour désigner les URI faisant référence à des objets accessibles au moyen de protocoles existants (typiquement, le protocole Http pour les « http URI »). La dualité URI/URL n'est toutefois pas encore véritablement creusée. De même, si une place est accordée aux URN (*Uniform Resource Name*), ces noms « plus pérennes que les URL », nous sommes néanmoins renvoyés aux travaux développés à la même époque et en parallèle, par un autre groupe de travail de l'IETF, pour plus de précisions. À mesure que le Web gagnait en réalité, la problématique proprement syntaxique du nommage acquit une importance accrue dans un contexte où se posait avec de plus en plus d'acuité la question de sa coexistence avec d'autres systèmes d'information. Outre cet aspect, le texte de 1994 ajoute une précision supplémentaire, mettant en garde contre une possible confusion avec les chemins d'accès hiérarchisés des systèmes de fichiers Unix. Autant l'arborescence d'un système de fichiers simule par voie métaphorique un ordre spatial qui nous permet très efficacement de nous repérer, autant elle ne dévoile *nécessairement pas* la cartographie réelle des données physiquement stockées sur nos ordinateurs. Qui plus est, la métaphore hiérarchique opère avant tout au niveau local. Ce qui signifie qu'un transfert vers un système global décentralisé, tel que le Web, n'irait pas sans poser de nombreuses difficultés (en matière de pérennité de l'accès, de sécurité, etc.). D'où cette précision supplémentaire : la portion syntaxique d'une URI désignée comme un « chemin » (*path*, à ne surtout pas confondre avec un « chemin d'accès » de type Unix !) n'a pas de signification décodable, en droit, par un client, bien qu'un utilisateur puisse évidemment l'interpréter à ses risques et périls (ou ceux du gestionnaire du site concerné, en l'occurrence). Le client ne saurait donc en inférer la structure présidant au stockage des données dans une base située sur un serveur distant, à partir de la hiérarchie apparente, écrite à même l'URI. De ce point de vue, l'URI est donc, *en droit*, mais pas toujours *de fait*, « opaque » pour le client. De nombreuses URI sont en effet construites à partir des structures arborescentes inhérentes aux bases de données – précisément comme des chemins d'accès ! Rien ne l'interdit. C'est précisément pourquoi Tim Berners-Lee met ici en garde contre les conséquences de telles pratiques, qui conduisent à transférer au niveau global des contraintes locales, techniques, forcément fluentes en ce qu'invariablement elles imposeront des révisions allant dans le sens contraire de l'effort de pérennisation qu'exigent le recours aux identifiants du Web. Par rapport aux UDI, la demande (quasi morale) d'abstraction concerne désormais de nouvelles déterminations contingentes, liées non plus aux particularités d'un document mais bien aux systèmes techniques permettant d'en opérationnaliser l'accès⁸.

⁷ L'IETF (*Internet Engineering Task Force*) est l'organisme où sont publiées les RFC (*Requests For Comments*), document qui comprennent les principaux standards d'Internet (de même que la plupart de ceux qui touchent à l'architecture du Web).

⁸ Anticipant quelque peu sur la suite, relevons que ceci pose la question du nommage, non du point de vue sémantique mais bel et bien du point de vue de *l'écriture des identifiants*. Or, il n'existe aucun consensus en la matière, à peine quelques bonnes pratiques glanées çà et là. On peut y déceler une tension entre deux formes d'engagements qui elles-mêmes ne sont pas uniformes. D'une part un engagement que l'on dira « ontologique », touchant aussi bien au Web qu'au Web de données, où l'enjeu est de rendre les URI interprétables dans la double perspective d'une manipulation rendue plus aisée des schémas et d'une identification plus lisible des ressources. D'autre part, un engagement computationnel tirant dans deux

L'étape suivante est postérieure à la création du W3C⁹. Elle voit le standard concernant les URI se scinder en deux. De là proviennent les URL¹⁰, dont l'acronyme demeure aujourd'hui encore bien connu du grand public, et les URN¹¹. La raison derrière cette partition semble à première vue des plus simples. Le Web n'étant pas doté nativement d'un système de suivis des versions (*versioning*), une adresse permet d'accéder à des contenus fluents, sans qu'à chaque version ne soit assigné un identifiant unique. Outre les « adresses » vers ces contenus, d'autres pointeurs ont de ce fait été standardisés de façon à permettre une identification pérenne, selon un modèle en vigueur dans le monde de la bibliothéconomie. Tel est le rôle échu aux URN, noms propres pérennes utilisés pour identifier tous types d'objets. Contrairement au URL, associées en priorité au protocole Http, les standards consacrés aux URN ne tiennent aucunement compte des modalités d'accès aux contenus définies par les protocoles du Web. Ce volet est ainsi laissé à la discrétion des gestionnaires en charge des différents « schèmes » d'URN (organisation responsable des différentes familles d'URN telles <issn:>, <isbn:>, etc.). En complément à cette bipartition, il revient à de nouveaux pointeurs, les URC¹², de donner accès à des informations au sujet des objets identifiés par les URN. En définitive, la tripartition envisagée fut donc la suivante : aux URL revenait la charge de trouver où localiser des ressources accessibles ; aux URN, l'identification d'objets stables (parce qu'inaccessibles) ; aux URC, enfin, de fournir des « méta informations » au sujet des contenus et objets associés aux deux familles d'identifiant qui précèdent. Tel est, du moins, le point de vue des concepteurs des URN, qui dégagèrent de nouvelles entités en posant les distinctions permettant d'en rendre compte. Or, à l'épreuve du temps, celles-ci se sont avérées problématiques. L'opposition entre URN et URL tient à la possibilité de localiser les ressources, affaire des URL et non des URN. Pourtant, à y regarder de plus près, on constate que la RFC 1738, consacrée aux URL, s'inscrit dans la lignée des travaux antérieurs sur les UDI et les URI. L'adressage d'une

directions opposées : du côté d'une opacité complète ou, à l'inverse, d'une transparence laissant affleurer les détails techniques de l'implémentation sous-jacente. Pérenniser dans le premier cas, offrir des prises techniques pour des réutilisateurs avertis dans le second, tels sont les enjeux contradictoires de cette autre forme d'engagement. On peut en tirer le schéma suivant :

Engagement ontologique		Engagement computationnel	
Ressources Web	Types (Web de données)	Opaque	Transparent

⁹ Le *World Wide Web Consortium*, ou W3C, a été créé par Tim Berners-Lee en 1994. Il rassemble des centres de recherches et des industriels et assure la gouvernance technique du Web. On lui doit la mise au point des recommandations du Web Sémantique.

¹⁰ *Uniform Resource Locators*, (Berners-Lee *et al.*, 1994).

¹¹ *Uniform Resource Names*, (Sollins & Masinter, 1994).

¹² *Uniform Resource Characteristics* (ou *Citations*), (URI working group, 1994 ; Mealling, 1994 ; Daniel & Mealling, 1995 ; Hoffman & Daniel, 1995 ; Daniel, 1995 ; Fielding, 1995 ; Soergel, 1997). Les URC elles-mêmes témoignent de la difficulté à poser de nouvelles entités stabilisées. Les distinctions sur lesquelles elles s'appuient ont en effet été associées à des performances – ainsi qu'aux propriétés censées leur correspondre – que les URL ont "récupérées" par la suite, ôtant du même coup aux premières toute *raison d'être* (au sens fort du terme). Sur ce point, cf. (Monnin, 2013b, pp. 87-100).

ressource continue en effet d'être défini sur un mode abstrait. En outre, s'il est juste de découpler la ressource, entendue comme « unité d'information », de son « lieu » propre, d'un ancrage localisable, en revanche lui accoler invariablement un couple d'URN *et* d'URL contraste semble-t-il avec cette volonté. L'on est en effet conduit à se demander si une unité d'information aura besoin, *dans tous les cas*, d'un emplacement où résider (emplacement qui renvoie dès lors à l'adressage pris en charge au moyen des URL). Paradoxalement, la nécessité d'un adressage conçu sur le mode d'une localisation se lit moins dans la RFC 1738 (URL), que dans la RFC 1737 (URN). Elle semble avant tout nécessaire pour justifier, en miroir et par contraste, le caractère pérenne concédé aux URN ainsi qu'à leurs corrélats non-localisables. En réalité, les URL se situent dans un entre-deux à mi-chemin entre adresses physiques et purs noms propres, qui déstabilise le processus de distinction mis en œuvre du point de vue des URN (on notera d'ailleurs que la RFC 1738 ne mentionne à aucun moment les URN, se situant dans le sillage des URI, jamais en contrepoint des URN).

Pour comprendre la répartition des propriétés entre ces différents types de pointeurs, il convient de se tourner vers un autre standard, destiné à livrer une caractérisation générique des identifiants du Web : les IRL¹³. Celui-ci fixe les grands principes auxquels ceux-ci doivent se soumettre, apportant de précieuses précisions au sujet de la ressource, qui témoignent au passage de la difficulté à faire cohabiter URN et URL de concert :

« A resource can be many things. Besides the non-networked or non-electronic resources just mentioned, familiar examples are an electronic document, an image, a server (e.g., FTP, Gopher, Telnet, HTTP), or a collection of items (e.g., Gopher menu, FTP directory, HTML page). [...] Furthermore, the nature of certain potential resources, especially animate beings or physical objects with no electronic instantiation, makes network access meaningless in some cases; such resources have locators that would imply non-networked access, but again, access is not guaranteed. »¹⁴

Une ressource inclut donc des éléments présents sur le réseau – *ou non*. Ceci s'explique aisément si l'on s'avise que dans chaque cas l'accès n'est nullement garanti à l'échelle d'un système authentiquement décentralisé, mettant de ce fait à mal l'adéquation de la notion même d'adressage, s'il faut y voir un lien entre deux relata – par contraste avec un « simple » pointeur (« *A resource locator describes a location but never guarantees that access may be established.* »). Seule une autorité unique, située en surplomb, aurait la capacité de maintenir des liaisons fixes et assurées. Or, avec le Web, nous évoluons, rappelons-le, dans un univers décentralisé, raison pour laquelle aucune institution ne reçoit un tel mandat. Aussi, le type très particulier d'adressage que performant ces *locators* ne saurait par conséquent se confondre avec une fonction d'accès *garantie*. Quant à l'accès aux *networked objects* il ne suscite guère d'interrogations, au moins en apparence. Mais qu'en est-il à l'inverse des *non-networked objects*, également associés aux URL ? Pour l'heure, celui-ci semble par définition prescrit. Dans tous les cas de figures, par conséquent, l'accès n'est, quoi qu'il

¹³ *Internet Resource Locators*, (Kunze, 1995).

¹⁴ *Ibid.*

arrive, jamais garanti, et c'est bien ce trait, paradoxalement, qui est commun au URL et aux URN.

En fin de compte, la RFC 1736 ne livre aucune indication sur ce que signifie, pour une ressource, de se voir attribuer un lieu. D'ailleurs, certaines « n'existent pas encore ». Pourtant, IRL et URL seront censées leur fournir une localisation, ne serait-ce qu'« en attendant ». Pareil « lieu » excède visiblement les limites du réseau. Cette contradiction appelle donc un profond réexamen des distinctions posées jusqu'ici.

Ces difficultés signent toutefois moins l'abandon de la distinction URL/URN que de sa radicalisation, dont témoigne la cohabitation malaisée de RFC reposant sur des bases conceptuelles sensiblement différentes. Aussi ne désignera-t-elle plus désormais que deux rôles ou fonctions possibles des URI. La caractérisation moderne de ces dernières sera livrée dans deux RFC ultérieures, 2396 et 3986. Pourtant, ces deux standards n'en fournissent en aucun cas la version la plus complète, ni les raisons qui conduisirent à ces aménagements. Les examiner exige que l'on s'arrête quelques instants sur le « style d'architecture » du Web, en vertu duquel toutes les distinctions précédemment examinée subirent une complète remise à plat.

II.2 - REST et la stabilisation des grands types du Web

Il fallut en effet attendre le travail réalisé par Roy Fielding sur le style d'architecture REST¹⁵, autrement dit les grands principes guidant le design du Web, pour que celui-ci soit enfin « mis en conformité avec lui-même »¹⁶. Autrement dit, pour que ses instanciations s'accordent avec ses principes, étant entendu que les seconds furent progressivement dégagés à partir des premières. REST a ceci d'original que les ressources y occupent désormais le devant de la scène. Leur introduction constitue au demeurant l'une des innovations majeures de ce style d'architecture, consacrant le fait que le Web ne saurait être un simple système de consultation de fichiers ou de documents.

Pour expliquer une telle innovation, qui repose bien, en définitive, sur l'émergence d'une entité nouvelle, trois raisons principales sont avancées. *La première* tire argument du fait que certaines URI sont susceptibles d'identifier des services et ce qui résulte de leur utilisation ponctuelle. Une URI, dans ce cas précis, identifierait par exemple la possibilité d'interagir avec la machine et non le fruit ponctuel d'une telle interaction (éventuellement téléchargeable sous la forme d'un fichier). *La seconde* a trait aux limitations du Web. Les vertus de son design épuré sont aussi les vices de tout bon dispositif de gestion des fichiers : l'absence d'un quelconque système de « *versioning* »¹⁷, en particulier, s'avère rédhibitoire. Dans les termes des première RFC, le scénario envisagé était le

¹⁵ (Fielding, 2000) ; (Fielding & Taylor, 2002). Sur l'importance de REST, voir en particulier (Monnin, 2012 ; 2013a ; 2013b).

¹⁶ Sur cette expression, voir (Monnin, 2012).

¹⁷ Le *versioning*, ou « gestion de versions » en français, est un terme issu de l'informatique qui désigne l'enregistrement de toutes les versions d'un fichier (généralement un fichier texte contenant du code source) afin de permettre l'archivage et le suivi du travail de création d'un logiciel.

Acronymes	Documents	Propriétés	Corrélat
UDI	Articles (1992)	Nom logique, à ne pas confondre avec une adresse physique pour pallier les déplacements physiques du document ; rendu pérenne par la mise à l'écart des détails adventices	Objet ou document, unité de consultation plus que de stockage, unité sérielle en tant que série de réponse formulée à la suite d'une requête dotée d'un nom logique
URI	RFC 1630 (1994)	Cf. ci-dessus. Distinct d'un chemin d'accès local ¹⁸ ; de portée globale ; en droit opaque ; abstrait des contingences du système technique par lequel l'accès aux contenus est géré	Des objets accessibles si les URI sont également des URL
URL	RFC 1738 (1994)	Adresse non-physique	Des ressources (pas encore définies), identifiées de manière abstraite (des contenus accessibles du point de vue de la RFC 1739)
URN	RFC 1737 (1994)	Nom, identifiant	Ressource stable, inaccessible
IRL	RFC 1736 (1995)	Adresse (URL), Identifiant (URN), Description (URC)	Ressource – <i>networked</i> ou <i>non-networked</i>
URC	Différents drafts de standards	Méta-information	Liste d'identifiants – Document

Tableau 1

Récapitulatif des principales caractéristiques des localisateurs du Web jusqu'en 1995

¹⁸ De ce point de vue, l'UDI pourrait presque désigner la famille de tous les chemins d'accès. La famille de tous les chemins d'accès ? On pourrait en effet l'assimiler à un chemin d'accès abstrait, qui n'est pas censé refléter une hiérarchie (taxonomique, voire d'un simple empilement de dossiers) ni désigner une localisation précise mais garantir l'accès aux représentations d'une ressource. En ce sens, et à rebours même de la métaphore spatiale du chemin, il s'agit bien d'un « chemin d'accès » en un sens dé-spatialisé, essentiellement gagé sur la notion d'« accès ».

suivant : il s'agissait pour l'auteur d'un document de lui créer une adresse abstraite. Pourtant, remarque Fielding, en identifiant ainsi un « contenu transféré » sur le réseau (un fichier) par son adresse, même abstraite, chaque modification dudit contenu requerrait en toute logique une nouvelle adresse. Or, le Web n'a tout simplement pas été pensé pour gérer cette situation. Au lieu d'un défaut, il en va d'une décision de conception plus ou moins bien comprise initialement. *Enfin*, à certaines URI ne correspondent aucun documents, c'est notamment le cas lorsque celles-ci sont utilisées en guise d'URN, ou encore dans l'attente de la mise en service du dispositif destiné à générer la publication et l'accès aux contenus. Pour toutes ces raisons et d'autres encore, spécifiques au style REST, la ressource est telle que :

« [...] any concept that might be the target of an author's hypertext reference must fit within the definition of a resource. [...] More precisely, a resource R is a temporally varying membership function $M R(t)$, which for time t maps to a set of entities, or values, which are equivalent. The values in the set may be resource representations and/or resource identifiers. A resource can map to the empty set, which allows references to be made to a concept before any realization of that concept exists – a notion that was foreign to most hypertext systems prior to the Web. »¹⁹

Être, ou plus précisément, être un objet, ce n'est plus « être la valeur d'une variable liée », selon l'adage quinién fameux, mais quelque chose que l'on peut représenter par la fonction à laquelle une telle variable est associée et qui assure une correspondance, à un instant t , avec des entités – ou des valeurs – équivalentes. Insistons sur le mot « équivalentes ». Équivalente et non *identiques*. Équivalentes, à savoir valant pour cette correspondance à un instant t , et non identiques les unes aux autres.

Avec Quine, les objets se muaient en « *posits* ». *Posits* situés au même niveau que les entités dont une théorie ou un schème conceptuel supposent l'existence et hors desquels il n'y a pas de sens à les envisager (sauf à se situer à l'intérieur d'un schème conceptuel rival). Dans cette optique quiniénne, les individus ne reçoivent plus guère la priorité ontologique dont la tradition les affublait : sujet, suppôt, substance, essence, etc. Bien au contraire. Nulle chose désormais sans un mécanisme d'individuation préalable. Reste évidemment à délimiter la nature de cette procédure « sur » le Web. La transformation induite par le glissement de la variable vers la fonction engage selon nous une toute autre conception. En revanche, entre les deux, le point commun demeure l'accent indubitablement mis sur l'*individuation* :

« The definition of resource in REST is based on a simple premise: identifiers should change as infrequently as possible. Because the Web uses embedded identifiers rather than link servers, authors need an identifier that closely matches the semantics they intend by a hypermedia reference, allowing the reference to remain static even though the result of accessing that reference may change over time. REST accomplishes this by defining a resource to be the semantics of what the author

¹⁹ (Fielding & Taylor, 2002).

intends to identify, rather than the value corresponding to those semantics at the time the reference is created. It is then left to the author to ensure that the identifier chosen for a reference does indeed identify the intended semantics. »²⁰

Soulignons ce point crucial : la ressource est le contenu de ce qu'un auteur entend identifier. En complément, une URI permettra simplement d'y faire référence à l'échelle du Web – voire d'accéder à d'éventuelles « représentations-http » associées – le nom retenu pour désigner les contenus attachés à une ressource qui transitent sur le Web. Cette définition ne concède nul privilège à quelque procédure d'individuation que ce soit (pas plus la manière dont nous faisons référence dans le langage ordinaire qu'une autre). L'identité de la ressource ne s'ancre dans aucune thèse métaphysique ou langagière sous-jacente : l'ajointement d'une URI l'opérationnalise à soi seule. Tel était déjà l'enseignement de la RFC 3986 : le contenu de ce qu'un auteur entend identifier n'est nullement marqué, *a priori*, du sceau de l'immuable. En conséquence de quoi, unicité et stabilité de la ressource sont tout entières gagées, de l'extérieur, sur le recours aux identifiants les plus pérennes possibles : « *The definition of resource in REST is based on a simple premise: identifiers should change as infrequently as possible* »²¹. À ce prix, sa dynamique interne, au même titre que sa multiplicité, sont respectées.

Parmi les exemples mentionnés, un cas revient fréquemment. La possibilité qu'intervienne un croisement des *valeurs* associées à deux ressources différentes est soulignée pour appuyer la nécessité de ne pas s'en tenir à celles-ci. Ainsi, les valeurs de la-version-préférée-de-l'auteur-de-cet-article et le-papier-publié-dans-les-actes-de-la-conférence-X, ou, dans un autre domaine (celui de la gestion de versions du code-source informatique), la-dernière-version, la-révision-numéro-1.2.7 ou la-révision-incluse-avec-le-livrable-Orange, sont peut-être amenées à voir leurs représentations se croiser en un instant donné du temps. Cependant, ce ne sont précisément que des croisements ponctuels : l'identité se joue ici à un tout autre niveau.

Les avantages de cette définition sont de trois ordres. D'une part, elle permet d'opérer à un niveau suffisant de généralité pour ne pas avoir à distinguer artificiellement les sources d'information, les unes selon leur type, les autres selon leur implémentation. Ensuite, elle autorise un « couplage tardif » (*late binding*), ou, dirions-nous pour plus de précision, « distal », entre ressources et représentations, permettant notamment à la négociation de contenu²²

²⁰ *Ibid.*

²¹ *Ibid.*

²² La négociation de contenu (« conneg » en abrégé) est une fonctionnalité cruciale du protocole Http. Elle permet de spécifier les termes de la requête qu'un client adresse à un serveur en fonction de paramètres tels que le langage ou le format ; à charge ensuite pour les serveurs d'être en capacité de les émettre et pour les clients de les exprimer. Avec la négociation de contenu, il devient impossible de poser une relation fonctionnelle (1:1) entre *une* URI et *une* représentation car les paramètres en fonction desquels les représentations varient potentiellement l'interdisent. La négociation de contenu concourt à expliquer la montée en abstraction que représente la ressource au regard du document, du fichier ou encore de la page. Non seulement les représentations évoluent au fil du temps, mais qui plus est, les formes spécifiques qu'elles revêtent ponctuellement sont loin d'être fixes (d'ailleurs, si tant est que ces variations dépendent de paramètres activés à la volée et que les représentations soient générées sur le

d'intervenir très en aval, à la demande expresse d'un client (variations diachroniques). D'autre part, enfin, en donnant la possibilité de faire référence à un concept au lieu d'une représentation singulière, elle dispense d'avoir à s'astreindre à réécrire les liens à chaque changement de représentations au fil du temps (variations synchroniques).

Nonobstant ces raisons et leurs justifications, une architecture opérant sur des « concepts » a de quoi désarçonner. À plus forte raison si l'on considère qu'elle émane d'un ingénieur, dont la mission, dégager les grands principes du Web en même temps qu'il les implémente, ne laisse guère augurer de tels développements. Fielding lui-même a bien conscience du paradoxe sur lequel il fait fonds, le détaillant dans un paragraphe intitulé, avec une immense audace, « *Manipulating Shadows* » :

« 7.1.2 Manipulating Shadows. Defining resource such that a URI identifies a concept rather than a document leaves us with another question: how does a user access, manipulate, or transfer a concept such that they can get something useful when a hypertext link is selected? REST answers that question by defining the things that are manipulated to be representations of the identified resource, rather than the resource itself. An origin server maintains a mapping from resource identifiers to the set of representations corresponding to each resource. A resource is therefore manipulated by transferring representations through the generic interface defined by the resource identifier. »²³

La lecture des standards regorge de formulations faisant état d'actions sur la ressource, usant de raccourcis parfois trompeurs. Or, la ressource n'étant jamais, par définition, qu'un concept, « une ombre », elle ne saurait constituer l'unité opérationnelle ou le support causal manipulé par une machine (le client et sa requête, le serveur et sa réponse). Il ne s'agit donc jamais de l'atteindre directement – de la toucher de son index, elle qui, pourtant, est désignée par une URI. La formulation interpelle : « *An origin server maintains a mapping from resource identifiers to the set of representations corresponding to each resource.* » Concrètement, physiquement, la ressource n'est nulle part. D'un point de vue matériel, sur le serveur, l'articulation s'opère donc entre des médiateurs tout à fait tangibles : l'identifiant de la ressource (URI), l'ensemble des représentations valant pour cette dernière à un instant *t*, et l'ensemble des dispositifs qui auront permis de les générer. La ressource n'est autre que ce qui articule tous ces éléments ; autant que le résultat de leur articulation d'ailleurs, dans un mouvement de constitution réciproque sans cesse à renégocier.

REST en vient à distinguer en définitive trois entités : la ressource, l'état représentationnel de la ressource et sa représentation (l'acronyme REST, pour *REpresentational State Transfer*, ou « transfert d'état représentationnel », est le reflet de cette tripartition).

même mode, il n'y a tout simplement plus de sens à envisager une telle éventualité). Dans ces conditions, le critère opérant la synthèse des représentations ne sera rien d'autre que leur relative *fidélité* vis-à-vis de la ressource au nom de laquelle elles auront été servies.

²³ (Fielding & Taylor, 2002).

II.2.1 - La ressource

Selon la RFC 2396 (Berners-Lee *et al.*, 1998), une ressource peut être « n'importe quoi ». Roy Fielding l'a qualifiée d'« ombre » ou de « concept », posant de ce fait un abîme entre les ressources et les documents, fussent-ils numériques. Par définition, les ressources ne peuvent jamais être ni stockées ni directement accessibles, et ne sont manipulées que par le biais de leurs représentations.

II.2.2 - Les états d'une ressource

Les systèmes adhérant à REST sont dits « *stateless* » (sans états) du fait de l'absence de session maintenue sur le serveur. Cependant, bien que les ressources demeurent identiques à elles-mêmes (ou tout du moins, le *devraient*), on constate, au vu de leurs représentations, qu'elles n'en livrent pas moins des résultats éminemment variables. REST distingue donc une ressource de l'*état* dans lequel elle se trouve lorsqu'on l'interroge. Ce point fait écho à la distinction entre une fonction, son argument (ici la requête, *hic et nunc*) et son parcours de valeurs. Cette assimilation trouve cependant sa limite dans la mesure où toute fonction, nous l'avons rapidement évoqué, ne saurait se concevoir sans un formulaire adéquat. Autrement dit, sans une sémiotique associée ou une raison graphique. Sans écriture, pas de fonction. Or, sur quoi s'appuyer si l'on entend maintenir le caractère abstrait des ressources ? En réponse à cette interrogation, nous proposons de comprendre les ressources comme des règles²⁴, précisant ainsi l'assimilation opérée par Fielding entre ressources et concepts (il convient en effet de noter que, dans la littérature philosophique, les concepts sont eux-mêmes souvent traités comme des règles). Assimiler la ressource à une règle permet de mieux comprendre comment et pourquoi les états sont produits : par le *suivi* de la règle. Fondamentalement, une ressource permet de générer des états. Ces états sont ensuite rendus accessibles au fil du temps sous une forme matérielle adaptée au réseau (un flux de données) et, surtout, *fidèle* à la ressource, sans que ce point soit davantage formalisé que cela, laissant une immense latitude d'appréciation en la matière (en dépit de tentatives infructueuses autour des « définitions d'URI »²⁵). La difficulté de faire intervenir une symbolisation à cet endroit est en effet d'au moins deux ordres : a) toute « définition » a le défaut d'exiger de circonscrire le périmètre d'une ressource *a priori*, l'effort déployé dans un second temps en matière de calcul pouvant presque se comparer à un programme exécutant une suite d'instructions, l'exécution elle-même devenant presque accessoire ; b) d'autre part, la possibilité d'interpréter la règle à partir d'une définition que l'on suivrait pour calculer les représentations de la ressource pose la question de la nécessité de faire intervenir une nouvelle règle associée à ce symbole en vue de l'interpréter correctement, entraînant du même coup une régression à l'infini.

Naturellement, à parler de règles, certains cas deviennent pour le moins étranges. Tim Berners-Lee est-il une règle ? Bien sûr que non ! Mais une règle/ressource étant un moyen d'identifier Tim Berners-Lee, ce que l'on appelle « Tim Berners-Lee » dépendra toujours de la façon dont on individue ce

²⁴ Une règle est un « standard de correction » (Glock, 2003, p. 517) en fonction duquel il est possible d'exercer tout type d'activité (notamment, mais sans toutefois s'y limiter, le calcul).

²⁵ Sur cette question, voir en particulier (Rees, 2011).

« sujet ». Il peut s'agir soit du-fondateur-du-Web, du-directeur-du-W3C, d'un-homme-né-de-X-et-Y, ou simplement de « Tim Berners-Lee », quelle que soit la façon dont on détermine ce qui se tient derrière ce nom propre²⁶. Finalement, ce sont quatre ressources différentes ou, en d'autres termes, quatre objets différents, quatre manières différentes de se saisir de quelque chose ou de l'individuer. Au-delà de *ce que* l'on individue, la ressource, entendue de manière très générique, correspond à une *règle d'individuation*. Il est particulièrement important d'opérer la distinction entre ressources et objets au sens traditionnel du terme²⁷ car rien n'assure qu'une ressource « corresponde » exactement à « une chose réelle » dans le monde, simplement parce qu'elle a été publiée sur le Web. D'autant plus que l'objectif du Web Sémantique n'a jamais été de trouver le moyen de résoudre cette difficulté²⁸. Aux ressources ne doivent pas nécessairement correspondre des descriptions véridiques (comment l'établir techniquement et uniformément ? Comment l'établir *tout court* ?). En revanche, elles doivent avoir assez de contenu pour spécifier, comme l'expliquent Fielding et Taylor, ce qu'un auteur a l'intention d'identifier, et, dans le meilleur des cas, pour en calculer des « représentations ».

II.2.3 - Les états représentationnels d'une ressource

Les états restent abstraits. Ils partagent avec les ressources la caractéristique de ne pas être accessibles en tant que tels. Ce qui l'est, en revanche, c'est la *représentation* de l'état d'une ressource, son inscription sur un support physique.

II.3 - Le Web social : nouvelles distinctions et nouveaux types. L'exemple du tag à l'aune de l'architecture du Web

Élargissons maintenant notre réflexion menée, jusqu'à présent à partir de ce que nous apprend l'architecture du Web, en nous penchant sur une évolution typique du Web social. La notion de tag est en effet emblématique du Web dit 2.0, ayant été définie par Joshua Schachter, le créateur du site de *social bookmarking* del.icio.us (devenu plus tard delicious.com, après son rachat par Yahoo!), pour désigner l'utilisation de libellés de toutes natures (sans les restreindre à des mots clefs) destinés à qualifier des contenus sur le Web. On s'étonnera peut-être de voir succéder aux grands principes hérités du style d'architecture REST des considérations portant en apparence sur un objet de la couche applicative (la couche applicative de cette autre couche applicative – d'Internet – que constitue déjà le Web !)²⁹.

²⁶ Il est bien sûr envisageable d'en rester là. Non que le nom propre garantisse quoi que ce soit au sujet de l'objet à lui seul. En revanche, dans certaines situations, il se suffit à lui-même, *en lieu et place de l'objet*.

²⁷ Le sens historique du mot « objet », en philosophie, apparaîtra nettement plus proche de l'acception ici examinée du mot « ressource » que de son équivalent actuellement en circulation.

²⁸ Comme l'explique très joliment Larry Masinter, éditeur de plusieurs RFC très importantes : « *Naming is printing money* » (<http://www.slideshare.net/PhiloWeb/larry-masinter-philoweb>). Il convient donc de garder à l'esprit que le Web Sémantique n'a pas été conçu dans le but de distinguer la fausse monnaie de l'authentique (même si, comme nous le fait remarquer un *reviewer* que nous remercions, il n'est à cet égard, « ni nécessaire, ni *a fortiori* suffisant mais pressenti comme utile »).

²⁹ Notons toutefois que le même ordre de succession présidait déjà à la présentation donnée dans (Berners-Lee *et al.*, 2006), où la discussion autour des folksonomies et du tagging s'insérait au cœur d'un chapitre consacré aux grands principes de l'architecture du Web et du Web Sémantique.

Néanmoins, notons tout d'abord que le principe derrière le tagging (tel que nous allons l'examiner ici), est plus générique, qui, rapporté par exemple au langage OWL de création d'ontologies, se retrouve trait pour trait dans la liaison qui peut être faite à un haut niveau d'abstraction entre un littéral (*datatype*) et une URI via la relation `<owl:DatatypeProperty>` (par opposition à `<owl:ObjectProperty>`, associant des ressources identifiées chacune par une URI). Tout porte donc à croire qu'il s'agit bien d'un principe, la relation entre URI et littéraux, qui transcende telle ou telle implémentation donnée. L'intérêt du tagging fut néanmoins de permettre de le déployer sur une grande échelle, à tel point qu'il a pu, un temps, passer pour une alternative crédible aux outils de classifications plus formels que sont les taxonomies, les thésaurus ou encore les ontologies informatiques. Cette dimension d'appropriation nous intéresse tout particulièrement car elle permettra de mieux situer l'apport des « utilisateurs » dans la continuité des discussions ouvertes dans la seconde partie.

Les tags ont parfois été conçus à la manière de simples chaînes de caractères sises entre les balises ouvrantes et fermantes de l'élément HTML `<a>` qui sert à constituer des liens hypertextes. En cela, rien ne les distingue de n'importe quel élément HTML `<a>`, n'était-ce le fait qu'il est possible de les typer au moyen de microformats³⁰ (en particulier `rel="tag"`³¹, destiné à indiquer qu'un lien hypertextuel prend la valeur d'un tag).

Élément HTML `<a href="+ URI" + (microformat rel="tag")>+ libellé</a>`

Cette caractérisation, résumée de la façon suivante par Tantek Çelik, promoteur des microformats et créateur du wiki où la majorité d'entre eux ont été publiés et discutés, n'est pas sans poser de nombreuses difficultés :

By adding `rel="tag"` to a hyperlink, a page indicates that the destination of that hyperlink is an author-designated "tag" (or keyword/subject) for the current page. Note that a tag may just refer to a major portion of the current page (*i.e.* a blog post). *e.g.* by placing this link on a page, `<a href="http://technorati.com/tag/tech" rel="tag">tech</a>` the author indicates that the page (or some portion of the page) has the tag "tech". The linked page SHOULD exist, and it is the linked page, rather than the link text, that defines the tag. The last path component of the URL is the text of the tag, so `<a href="http://technorati.com/tag/tech" rel="tag">fish</a>` would indicate the tag "tech" rather than "fish".³²

En partant de cette définition, il appert en effet que taguer reviendrait simplement à choisir une URI plutôt qu'un libellé, à rebours des pratiques bien établies des utilisateurs.³³ En fait, dans l'exemple ci-dessus, Çelik renverse la

³⁰ Les microformats visent à enrichir les pages HTML au moyen de données structurées ajoutées aux attributs de ce langage.

³¹ <http://microformats.org/wiki/Rel-Tag>.

³² *Ibid.*

³³ C'est également une conception défendue par les créateurs de l'*Upper Tag Ontology*, (Ding *et al.*, 2008) : « *Tags are nothing more special than a typed hyperlink. We can use « rel » attribute to type hyperlinks* » En réalité, nombres de tags ne sont pris dans aucun lien hypertexte, à commencer par les machine tags de Flickr ou les hash tags de Twitter – du moins à l'origine, s'agissant de ces derniers.

perspective usuelle en focalisant son attention sur l'URI du tag « tech » présent sur le site Technorati.

```
<a href="http://technorati.com/tag/tech"
rel="tag">tech</a>
```

Ce qui précède est la formalisation d'un tag générique qui subsume ou agrège plusieurs ressources via leurs identifiants. Mais s'il existe bien une URI pour ce tag *générique*, c'est justement pour la raison très simple que des actions de tagging *individuelles* l'ont précédé. Celles-ci aboutissant, dans un second temps seulement, à la création d'une page recensant l'ensemble des liens tagués avec le libellé « tech » par tous les utilisateurs de cette plate-forme. D'ailleurs, dans certains cas, une telle page n'existe pas en tant que telle mais uniquement sous la forme d'une requête adressée à un moteur de recherche : l'URI correspondant n'identifie plus le-tag-générique-tech mais une requête générant des résultats à la volée³⁴. Éventuellement, d'ailleurs, en l'absence de tout tag individuel, auquel cas il est parfaitement illusoire de continuer à parler de tag *générique*. Utiliser le libellé « tech » en liaison avec l'URI du tag générique « tech » de Technorati n'a tout simplement aucun sens du point de vue d'un utilisateur dont l'action individuelle précède l'émergence d'un tel tag générique et la conditionne³⁵. On ne peut que le suivre si Çelik entend simplement affirmer que l'URI <http://technorati.com/tag/tech> identifie la ressource le-tag-générique-« tech »-sur-technorati, et non le-tag-« fish » (dans l'éventualité où le libellé associé contredirait l'URI qu'il qualifie à l'instar de fish).

Après tout, il a raison de donner la priorité à l'URI sur le libellé du lien hypertexte – ne serait-ce que parce qu'une ancre, à l'instar d'un lien, dépend constitutivement de la présence d'une URI. Cependant, cela n'a plus rien à voir avec la question du tagging : le choix d'un libellé, comme le tag singulier, précédant le tag générique et l'URI qui lui est associée. Les actions de tagging consistent ainsi, au minimum, en une ressource tagguée, un libellé et une relation qui les associe. Il en résulte des conditions d'identité extrêmement strictes, encadrant chaque acte de tagging singulier. Comprenons qu'ordinairement, via les interfaces les plus courantes du *social bookmarking*, une URI est associé à différents tags uniques procédant de l'ajout par un utilisateur d'un libellé, libellé associé à l'URI d'un tag générique, comme celle qu'évoquait justement Çelik (c'est d'ailleurs l'unique scénario correspondant à sa définition). Il n'est toutefois pas possible, dans pareille configuration, de retrouver, à même le tag, l'URI de la ressource annotée, pas plus que les trois informations suivantes :

- a) à quel type de ressource il est fait référence ;
- b) comment cette ressource est reliée à un tag ;

³⁴ Ce qui correspondrait par exemple à une autre ressource du type : les-résultats-d'une-recherche-sur-le-tag-« tech »-sur-technorati.

³⁵ Dans le cas contraire, on aura affaire à un référentiel (thésaurus ou assimilé), dont les concepts précèdent leur usage, et non à une folksonomie proprement dite.

c) ce que le libellé du tag signifie dans son usage actuel notamment en cas d'appariement entre un libellé et une URI du Web Sémantique, ce qui concerne tout particulièrement les scénarios de désambiguïsation³⁶.

En examinant des exemples de tagging sur une plate-forme emblématique telle que Delicious.com³⁷, force est de constater qu'ils ne répondent pas à ces critères drastiques. Sur Delicious, un tag communautaire identifie une ressource stable (un tag générique), autant qu'il livre accès, par l'intermédiaire de l'URI qui lui est adjointe, à ses représentations changeantes. Une URI comme

<http://delicious.com/alexandre_monnin/web>

a) *identifie* le tag-générique-« Web »-d'Alexandre-Monnin-sur-Delicious, autrement, dit le tag générique « Web » de sa folksonomie personnelle³⁸,

b) et donne accès à la liste constamment enrichie des ressources sélectionnées par cet utilisateur, partageant toutes un libellé en commun ; libellé ininterprété, ou plutôt dé-sémiotisé : réduit à une simple chaîne de caractères à partir de laquelle une URI du type <<http://technorati.com/tag/tech>> est forgée³⁹.

Il faudra donc comprendre comment passer d'une expression la plus précise possible du tag à une autre, beaucoup plus vague.

Ces préalables étant désormais posés, nous pouvons maintenant nous demander dans quelle mesure l'architecture du Web, ses grands types et les distinctions d'où ces derniers tirent leur origine, ouvrent de nouvelles perspectives pour l'ontologie.

III – L'ONTOLOGIE DU WEB : DISTINCTION, EXPLICITATION ET TRIADICITÉ

III.1 - Les ontologies et l'ontologie du Web

À première vue, ce que l'on nomme « ontologies » sur le Web s'inscrit dans la tradition ouverte en 1980 par John McCarthy⁴⁰, le fondateur de l'Intelligence Artificielle, avocat de la formalisation du sens commun par des moyens logiques avec Patrick J. Hayes. Empruntant à Quine, McCarthy fait explicitement mention de la définition quinienne de l'ontologie, définition extrêmement influente qui eut pour effet de replacer les questions ontologiques au centre des

³⁶ Le tagging sémantique repose en très large partie sur ce principe. Depuis MOAT (*Meaning of a Tag*, Passant *et al.*, 2009) et *Entity Describer* (Good & Kawa, *n.d.*), jusqu'à une application grand public comme Faviki (<http://www.faviki.com/pages/welcome/>) ou, dans un contexte culturel, l'expérimentation HDA-BO pour le Ministère de la Culture, à laquelle Alexandre Monnin a prêté son concours (<http://cblog.culture.fr/hda-lab%C2%A0-experimenter-le-tagging-semantique/>).

³⁷ Avant sa mise en vente par Yahoo! en 2011.

³⁸ La notion de folksonomie, personnelle ou collective, souffre de la même ambiguïté : en parlant de tous les tags, vise-t-on les occurrences (*token*) ou les tags génériques (*types*) ? Les deux acceptions sont possibles et la littérature en la matière ne se soucie guère de les préciser.

³⁹ De ce point de vue, on pourrait aller jusqu'à l'assimiler à un « objet-frontière personnel », dont on ne maîtrise plus les variations à l'échelon supérieur.

⁴⁰ (McCarthy, 1980).

réflexions dans l'espace analytique. Gare au contresens cependant, la définition donnée par Quine étant très largement idiosyncrasique. En effet, si l'ontologie pose la question « qu'y a-t-il ? » (*what there is*), elle se scinde désormais en deux disciplines : d'une part une théorie de l'existant, symbolisée par la maxime fameuse « être c'est être la valeur d'une variable liée » ; d'autre part une théorie des propriétés de ce qui existe. Sous couvert de poser deux disciplines distinctes, ontologie et « idéologie »⁴¹, Quine retrouve en fait, à rebours sans doute de ses intentions, deux des principales acceptions historiques de l'ontologie. La seconde remonte ainsi à la *Métaphysique* d'Aristote et sa reprise par la scolastique. Posant la question des propriétés des choses et des rapports de prédications qu'entretiennent le singulier et le général, elle donna lieu au Moyen-Âge à la célèbre querelle des universaux. La première correspond quant à elle au sens originel du mot « ontologie », issu d'auteurs appartenant à la scolastique tardive tels Goclenius et Lorhardus. Il désigne une partie de la métaphysique dévolue à la théorie de l'objet – le « quelque chose en général »⁴².

La reprise de l'ontologie sous l'impulsion de McCarthy prolongeait l'effort quinquennal déployé en vue d'élaborer un idiome canonique, véritable langage de la science construit au moyen de la logique du premier ordre, pour y exprimer ce qui relève chez l'auteur de « On what there is » aussi bien de l'ontologie que de l'idéologie. À ceci près que McCarthy, tout comme Hayes à sa suite, visaient davantage l'expression du sens commun que son élimination⁴³ (ce qui a pu conduire certains à qualifier ce courant de l'IA « d'aristotélicien »⁴⁴) ; tendance plus que jamais à l'œuvre dans le contexte du Web Sémantique. Ce dernier, dont les origines remontent à 1994, s'inscrit dans cette tradition tout en lui faisant subir une inflexion notable, l'accommodant partiellement aux grands principes architecturaux du Web. Ainsi, dans cette perspective, les ontologies sont-elles des combinaisons d'« adresses » ou d'« identifiants » (URL, URI) et de libellés, (noms d'objets ou de classes – leurs types – ressortissant à différents domaines), à l'aide de relations elles-mêmes dotées d'un localisateur (terme générique que nous utiliserons pour parler indifféremment d'une adresse ou d'un identifiant, à l'instar du mot « pointeur ») et d'un libellé, comme « a pour lien avec », « est un », voire « est une partie de ». De telles relations rendent possible l'établissement d'un réseau représentant l'ontologie d'un domaine (objets et propriétés/relations), tout comme l'articulation entre différents domaines (Munn & Smith, 2008).

⁴¹ (Quine, 1951) ; (Quine, 1983). Sur cette question, voir (Bourdeau, 2000 ; Decock, 2002) (une des questions abordées par ces deux ouvrages étant le rapport que la sémantique formelle et la théorie des modèles entretiennent avec l'ontologie).

⁴² On retrouve par exemple une telle dualité dans les deux livres de Frédéric Nef (Nef, 1998 ; 2006), consacrés pour l'un à l'objet quelconque, pour l'autre, aux propriétés des choses. Sur cette question de l'objet quelconque et ses sources médiévales, voir en particulier (Boulnois, 1999 ; Courtine, 1990).

⁴³ « [McCarthy 1959] proposed a program with "common sense" that would represent what it knows (mainly) by sentences in a suitable logical language. It would decide what to do by deducing a conclusion that it should perform a certain act. Performing the act would create a new situation, and it would again decide what to do. This requires representing both knowledge about the particular situation and general common sense knowledge as sentences of logic. » (McCarthy, 1980).

⁴⁴ Cf. (Bachimont, 1996).

En première approximation, il semble possible d'assimiler ces adresses/identifiants à des référents, et les noms/libellés à des sens. On pourrait donc rattacher cet usage du terme « ontologie » à la conception fré géenne⁴⁵, qui se construit sur une dualité entre des *référents* (par exemple la planète Vénus) et des *sens* associés à ces référents (par exemple, la qualification *d'étoile du soir* ou *d'étoile du matin*). L'ontologie par défaut dans certains cercles philosophiques et informatiques⁴⁶ est d'ailleurs, à l'heure actuelle, une ontologie de la *substance* (identifiée par un *nom propre*) et des *accidents* (identifiés par des *prédicats*), fusion abâtardie de l'ontologie aristotélicienne et de la logique fré géenne⁴⁷.

Cette distinction permet de reconstruire les trois positions dominantes à l'heure actuelle en métaphysique. Selon la première, la position aristotélicienne, les substances (particulières) constituent la cible d'une opération de référence, tandis que les propriétés (universelles comme particulières) sont les corrélats des prédicats. La suivante, la position nominaliste classique, pose que la référence est faite à des substrats particuliers, et la signification implique la contribution de propriétés particulières. Enfin, selon la dernière, d'inspiration « occamienne renouvelée » (ontologie de tropes⁴⁸), les référents sont identifiés par la compré sence de qualités particulières (de propriétés concrètes) et les significations liées par des similarités entre qualités ou propriétés particulières.

⁴⁵ Dans un célèbre article, intitulé *Sinn und Bedeutung*, Frege (1994) distinguait le *Sinn* (généralement traduit par le mot « sens ») de la *Bedeutung*, (« référence », « dénotation » ou « signification »), afin de rendre compte de l'hétérogénéité des modes d'accès à un même objet : « Il est naturel, écrivait-il, d'associer à un signe (nom, groupe de mots, caractères), outre ce qu'il désigne et qu'on pourrait appeler sa dénotation, ce que je voudrais appeler le sens du signe, où est contenu le mode de donation de l'objet » (*op. cit.*, p. 103). Deux expressions comme « L'étoile du matin » et « l'étoile soir » ont un même référent, Vénus, mais elles expriment des pensées différentes, différences qui correspondent à deux sens distincts : « La pensée ne peut donc pas être la dénotation de la proposition ; bien plutôt faut-il y voir le sens de la proposition » (*op. cit.*, p. 108). La nature du sens, comme du « mode de donation de l'objet », a entraîné aussi bien des rapprochements avec la phénoménologie husserlienne (Fisette, 1994) que des développements conduits par des auteurs travaillant sur les sémantiques cognitives. Il est toutefois éminemment périlleux d'attribuer ce type de considérations à Frege lui-même, dont les travaux demeurent marqués par un antipsychologisme foncier (l'équivoque porte alors sur ce qu'il entend par le mot « pensée », concept, qui, chez lui, ressortit précisément à tout autre chose qu'à l'ordre de la psychologie).

⁴⁶ Dont le représentant le plus connu est sans doute Barry Smith, qui ne considère plus appartenir à la communauté des philosophes mais des ingénieurs qui bâtissent des systèmes d'information autour d'ontologies de haut niveau (en particulier dans le domaine de la bio-informatique).

⁴⁷ Claude Imbert (Imbert, 1992 ; 1999) a largement insisté sur les restes de kantisme (voire d'aristotélisme) au cœur de certaines interprétations de la syntaxe fré géenne. Le fait que les ontologies actuelles soient à ce point pénétrées de la distinction fré géenne *Sinn/Bedeutung*, témoigne de la poursuite de l'entreprise ouverte par le « pacte apophantique », l'entrelacs de la logique et de l'ontologie.

⁴⁸ Les tropes, dans la métaphysique contemporaine, sont des qualités particulières (« ce rouge »). La théorie des tropes s'apparente au nominalisme en ceci qu'elle entend faire fi des universaux. Cet usage du mot « trope » a été introduit en 1953 par D.C. Williams (Williams, 1953).

Positions	<i>Bedeutung</i> -Référence (noms propres)	<i>Sinn</i> -Sens (prédicats)
Aristotélicienne	Substances	Propriétés
Nominaliste	Substrat particulier	Propriétés particulières
(Nominaliste) Tropiste	Compréhension de qualités particulières/propriétés concrètes	Similarités entre les qualités/propriétés de particuliers

Tableau 2

La distinction *Sinn/Bedeutung* et les positions ontologiques actuelles

III.2 - L'ontologie du web comme ontologie d'opérations (de distinctions)

Pourtant, dans la perspective que nous adoptons, le mode de structuration est plus fondamental que le genre des entités examinées. Cette approche, dynamique et relationnelle, demande à être prise en compte plus sérieusement par les ontologies philosophiques, généralement centrées sur des entités statiques. Dans cette perspective dynamique, les vieux dualismes de l'ontologie philosophique se réduisent à la dualité de deux *opérations* : a) *distinguer des items* – essayer de les identifier, même de manière relative, – et b) les *rendre*, ainsi que leurs articulations, *explicites* (par souci de brièveté, nous userons désormais des néologismes « explicite » et « explicitation » plutôt que « rendre explicite »).

L'explicitation est pour nous un processus qui débute avec une opération basique de distinction qui distingue des items précédemment non distingués. Mais en même temps, pour rendre explicites les articulations entre items distingués, elle use implicitement d'autres items, non encore distingués. La définition du type de ces entités est donc relative à l'étape du processus d'explicitation considérée. Parce que les distinctions sont de grain grossier à l'entame du processus, plusieurs genres d'entités demeurent indistingués à cette étape. On peut donc différencier trois stades. 1) Le type ontologique d'une entité qui n'est *pas distinguée* d'autres entités peut cependant *se manifester* dans une pratique ou un usage, donc en *restant implicite*. 2) Le type d'une entité est *distingué et demeure non explicite* quand elle a été captée mais que *les façons qu'elle a de se distinguer d'autres types d'entités* (qui lui serviraient de ce que nous appellerons des « distingueurs ») *ne sont pas elles-mêmes distinguées*. 3) Le type d'une entité est *explicite* quand *les façons qu'elle a de se distinguer d'autres entités ont elles-mêmes été distinguées*. Cette base triadique donne lieu à un processus récursif dès lors qu'il devient possible de 1) *distinguer des entités* en usant d'autres entités pour les distinguer, 2) de *rendre explicites ces distinctions* précédentes en *distinguant les distingueurs* au moyen de nouveaux distingueurs, 3) de rendre explicite ces distingueurs en les distinguant à l'aide de nouveaux distingueurs, et ainsi de suite.

À titre d'exemple (en repartant des types des ontologies classiques), les substances se peuvent distinguer *d'autres substances* : des substances particulières usent alors d'autres substances particulières en guise de distingueurs. Mais comment distinguer des substances particulières de *qualités* particulières ? Si nous affirmons que les substances en tant que telles se distinguent des qualités en tant que telles, usant par là-même des qualités en guise de distingueurs, cette réponse débouche automatiquement sur une autre question : quels traits des substances et des qualités permettent de les distinguer les unes des autres ? Nous pourrions répondre en mobilisant un nouveau type d'entités, le type des relations. Par exemple, si la « qualité » – et non la « substance » – constitue la catégorie de base, la relation de comprérence entre qualités (une qualité est comprérente à une autre si toutes deux coexistent dans un même faisceau (*bundle*) de qualités) construit le substrat à partir de ces qualités (au même niveau que les qualités et la relation). De même, si nous admettons une différence entre deux genres d'entités (substances et qualités), nous avons usé non seulement de la relation de différence mais aussi, à un autre niveau, des genres comme nouveau type d'entités – qu'on peut assimiler à des universaux.

Dès lors qu'à la première étape nous nous focalisons sur la substance, nous présupposons à minima d'autres substances pour être en mesure de distinguer cette substance initiale. Néanmoins, ces substances, utilisées en guise de distingueurs, ne sont pas elles-mêmes distinguées, elles demeurent indistinguées. À la suite de cette étape, nous avons envisagé trois possibilités : la *première*, qu'il n'y ait que des substances ; la *seconde*, que nous pourrions uniquement avoir des qualités et une relation – si ce que l'on appelle des substances résultent de la comprérence de qualités ; la *troisième*, que nous n'ayons que des relations – si nous considérons la comprérence et la similarité (ou la différence) comme plus fondamentales que leurs *relata*. Dans la foulée de cette seconde étape (visant par exemple à différencier les substances des qualités par exemple), nous pourrions aboutir à des genres (ou universaux), des relations ou les deux. Supposons qu'au cours d'une étape ultérieure nous ayons à rendre explicite en tant que distingueurs le ciment ou la connexion articulant une substance et ses qualités. Nous aurions alors à rendre explicite le distingueur de cette substance et de ces qualités, soit la relation entre la substance et les qualités. Si la qualité est un universel, la substance exemplifie cet universel, et la relation est une relation d'exemplification. Si la qualité est un particulier, nous avons besoin, outre la relation de comprérence, d'une relation d'instanciation, étant donné qu'une comprérence particulière est une instanciation d'une relation universelle de comprérence.

Si nous caractérisons les entités ontologiques que nous avons introduites en considérant les étapes au cours desquelles elles ont été rendues explicites, nous constatons alors que les particuliers, en tant que référents, sont des entités qui émergent au cours de la première étape. Les qualités, en tant que différant par leur type des substances, sont des entités ressortissant à la seconde étape. Les relations, distinguées des qualités et des substances, sont des entités qui apparaissent au cours de la troisième étape. Sans doute les genres ou les universaux, si nous y avons recours, sont-ils des entités qui ne font leur apparition qu'à l'étape suivante.

Un tel développement récursif permet d'intégrer d'emblée les processus dynamiques dans la sphère de l'ontologie, plutôt que de présupposer des entités statiques, définies une fois pour toutes. Au cours de son développement, ce processus charrie avec lui de nouveaux distingueurs, permettant de poser, en chemin, de nouvelles distinctions. Nous progressons ainsi pas à pas, du plus grossier vers le plus fin. En un sens, à la toute première étape, l'étape 0, nous avons uniquement affaire à des entités indistinguées. Plus tard, au cours de l'étape n° 1, nous avons des particuliers. Nous ne savons pas encore s'il s'agit de substances, de qualités ou de relations, ni même d'exemplification ou d'instanciation d'universaux : le type ontologique spécifique d'une entité est défini en conformité avec l'étape à laquelle il a été nécessaire de le distinguer d'autres types ontologiques. Jusqu'à cette étape, le type ontologique d'une entité peut être considéré comme un « type flottant », encore indéterminé au regard d'une pluralité de types encore possibles⁴⁹. Au cours de la seconde étape, nous distinguons, parmi les particuliers, qualités et substances. Les relations, s'il y en a, ne sont pas encore distinguées et ne le seront qu'au cours d'une troisième étape.

Dans ce qui suit, nous suggérons de concevoir l'ontologie du Web de la manière *dynamique* qui vient d'être décrite, en vertu d'un processus de distinctions de nature *triadique* plutôt que dualiste. Sur le Web, l'ontologie et les explicitations sont conditionnées par le processus de développement du réseau et de ce fait fonction des différentes étapes dudit processus, considéré sous un jour récursif. En ce sens, le Web Sémantique (Web 3.0) – le Web qui requiert des ontologies, pourrait conserver les traits du Web social (Web 2.0), dans lequel les folksonomies émergent du libre tagging, par les utilisateurs, d'informations et d'objets, et de la réunion, à différentes échelles selon des palliers successifs, de ces tags les uns avec les autres. La dynamique consistant à rendre explicites les types flottant de manière récursive semble fournir une propriété fondamentale de l'ontologie des Web 2.0 (social) et 3.0 (sémantique), et pourrait bien caractériser toute ontologie philosophique qui inclurait son propre processus d'explicitation.

III.3 - Au-delà des distinctions classiques : réseaux, ressources et tags

III.3.1 - Réseaux, liens et pointeurs

A priori, les distinctions des ontologies classiques apparaissent suffisantes pour reconstruire plusieurs concepts fondamentaux du Web, dont les URI. Les « adresses » pourraient en effet être vues comme les correspondants des référents et les noms, des « significations ». Les adresses semblent en effet avoir pour fonction essentielle de faire référence à des items. Quant aux noms (c'est également le cas des tags examinés plus bas), ils pointent apparemment vers des contenus signifiants, remplissant en première approximation le rôle dévolu aux significations dans l'ordre de la pensée ainsi qu'aux prédicats des formalismes

⁴⁹ Cette approche possède des affinités avec un certain pragmatisme actuel. Voir par exemple les travaux d'Antoine Hennion en sociologie, en particulier (Hennion, 2007), où ce dernier distingue, d'une manière sans doute encore trop dualiste « un monde externalisé, partagé, avec des entités autonomes » et un « monde internalisé, un monde de l'entre-deux, sans frontières claires, où aucun terme ne dispose *a priori* d'identité ou de propriétés arrêtées, où il s'agit de s'entre définir en participant activement à des opérations constitutives ».

logiques. Selon ce qui précède, adresses et noms sur le Web renvoient donc bien aux deux principales catégories fréggéennes⁵⁰:

Noms	Adresses
Significations	Référénts

Tableau 3
La distinction *Sinn/Bedeutung* et les URI

Seulement, une telle assimilation apparaît vite trop brutale. Nous pourrions croire que les adresses pointent vers des nœuds qui constituent les localisations de la structure du réseau. Le problème est que tout changement dans le réseau transforme la structure de ses relations, et par conséquent le rôle relationnel de ses nœuds. De nouveaux nœuds deviennent accessibles si l'on ajoute des liens. Étant donné qu'il n'existe pas de structure homogène de localisation, aucun espace homogène précédant les nœuds et les liens du réseau, la structure de celui-ci est modifiée à chaque fois qu'un nouveau nœud est ajouté – ou que de nouveaux nœuds et de nouveaux liens sont supprimés. L'ajout de pointeurs dans la représentation d'une ressource (traduction du lien hypertexte cliquable dans les termes de l'architecture du Web) modifie la structure du réseau que l'on peut établir et, partant, la situation relative (relationnelle) de chaque nœud. Ajouter de nouvelles significations revient dès lors à transformer le réseau et non simplement à y additionner par itération de nouveaux types. Ce qui, d'un point de vue fréggéen, pour lequel les significations demeurent stables et ne sont en aucun cas déterminées par les chemins inférentiels en constante évolution dans lesquels s'insèrent les termes conceptuels (contrepartie des significations dans le monde des processus concrets), n'est tout simplement pas envisageable. Au regard de la dynamique de construction et d'évolution d'un réseau intra-ontologique (entre les types d'une ontologie) ou inter-ontologiques (entre plusieurs ontologies et les types qui leurs sont associés)⁵¹, le type d'un item ne demeurera donc nullement fixe et sera susceptible de changer en fonction de ses interactions et d'éventuelles transformations consécutives de la structure du réseau.

Bien sûr, s'ajoute à cette dimension architecturale la réalité du Web, tramée par d'autres processus, quasi-métrologiques, visant à favoriser l'interconnexion pour augmenter la visibilité des contenus (songeons évidemment au PageRank de Google). Outre que le graphe du Web n'est pas connexe (il n'existe pas de chaîne permettant de relier l'un à l'autre deux nœuds quelconques), ce dernier comprenant des îlots « déconnectés », ses propriétés ont assurément une influence sur les répercussions qu'est susceptible d'entraîner une modification donnée. Ce qui ne change rien à la remise en cause du modèle fréggéen. En revanche, il convient d'ajouter à la vision holistique initiale une prise en compte des facteurs de connexité et de densité du graphe actuel, de même que les

⁵⁰ On trouvera un argumentaire plus détaillé de ce rapprochement, et sa critique, dans (Livet, 2012, pp. 398-400).

⁵¹ Pour emprunter un exemple issu du Web de données (ce constat valant néanmoins pour tout type de ressources).

processus ce qui le performant et qui diffèrent d'ailleurs en partie de ceux qui performant le graphe du Web de données. Le fait pour ce dernier d'être orienté, autrement dit que ses arcs aient un « sens » (soient notés (a, b) ou (b, a), contrairement aux arêtes qui sont de simples paires de sommets {a, b}), explique son caractère asymétrique. Ainsi, faire l'objet d'une désignation ou désigner soi-même aboutit à répartir de manière très inégale des propriétés émergentes telles que l'autorité, la centralité, la robustesse (tant au plan technique qu'économique), etc.

Du point de vue de l'évolution du réseau, comptent en premier lieu les interactions. Toutefois, à celles-ci s'ajoutent leurs relevés quasi-métrologiques, faisant apparaître de nouvelles asymétries (songeons à nouveau au PageRank de Google). Comptent également, à mesure que se développent ces interactions et que leurs relevés sont systématisés, la nature des termes ou *relata* ainsi constitués, et plus particulièrement, leur centralité. Ces éléments s'ajoutent à tous ceux que nous avons déjà examinés d'un point de vue strictement architectural. Le graphe du Web de données n'échappe pas à ces phénomènes additionnels, néanmoins leur impact reste encore à étudier plus en détail en ce qui le concerne.

III.3.2 - Ressources

Considérons d'autre part la relation qui va des localisateurs, identifiants ou adresses, à leurs référents – qui, dans le Web, sont des « ressources ». Ces ressources peuvent être constituées – au sens où elles les colligent – de « documents » de différents types (des textes, des images, des vidéos, de la musique, etc.), et leurs contenus se modifier dans le temps. Il est même possible qu'un site ne soit plus alimenté et que plus aucun contenu ne corresponde à une ressource donnée. De même, ce que l'on désigne génériquement au moyen du terme usuel d'« adresse » peut recouvrir différents opérateurs. Ainsi une URL est censée identifier la ressource par sa « localisation abstraite », comme potentiel mode d'accès ; alors qu'un URN est, lui, censé identifier la ressource comme objet, même (et surtout) en l'absence de toute possibilité d'y avoir accès.

En général, une URI s'appuie sur un mécanisme d'accès et donne les paramètres nécessaires pour retrouver le « lieu » (en fait, la possibilité technique d'accéder à un contenu, quel que soit son emplacement géographique) et le nom de la ressource. Mais elle ne garantit pas que quoi que ce soit sera donné par là-même (ni la ressource elle-même – en intégralité, ni ses représentations). Du coup, on semble en mesure d'utiliser un nom de référent sans référer vraiment, et un nom de concept (les contenus associés à la ressource) sans disposer du contenu conceptuel en question. Par ailleurs, un identifiant – ou une « localisation abstraite » – est aussi et avant tout le déclencheur d'un processus de recherche et si possible d'accès, et en ce sens, il désigne non seulement la ressource mais aussi le processus en question. Ce processus, est désigné de manière opératoire, alors que la relation à la ressource à laquelle il réfère reste encore en pointillés, le contenu de celle-ci pouvant varier.

En effet, donner accès ne revient pas simplement à pointer vers une ressource, donner accès consiste déjà à *user* d'une ressource. Si la ressource peut s'envisager à la manière d'une fonction, à l'instar de ce que propose Fielding, *user* d'une ressource, en un sens, revient à poser une *fonction de fonction*. Nous

sommes fort éloignés de la conception traditionnelle du référent. Ou, en tournant les choses autrement, nous découvrons qu'avoir accès à un référent exige bien davantage que le simple fait de pointer⁵². La comparaison avec une fonction s'impose avec une évidence d'autant plus grande qu'elle est bien *au cœur* de l'architecture du Web par l'entremise de REST.

En fin de compte, *deux interprétations des fonctions* s'opposent ici : une interprétation qui réintègre la substance⁵³ dans ses droits ; l'autre, qui l'en démet. La première est le fait des ontologies « classiques ». La fonction $F(x)$ y symbolise la qualification d'un sujet (x), une substance (x)* par un concept F qui lui-même renvoie à une propriété F^* . La seconde fait de la ressource une fonction $M_R(t)$ d'appartenance, ou plutôt d'agrégation des représentations coordonnées autour d'une ressource, variant au fil du temps, de la nature des requêtes et des modalités de les traiter. On pose alors une nouvelle distinction entre *deux fonctions* cette fois-ci : *pointer vers une ressource* et « *utiliser* » *cette ressource*, pour accéder à une représentation qualifiée⁵⁴ par ladite ressource.

Prenons la relation qui va dans l'autre sens, des éléments de la ressource (les représentations) aux différentes adresses. En changeant d'échelle, on peut concevoir que sans changer de contenu, une ressource puisse cependant changer de « site », et donc « d'adresse ». Elle peut aussi changer assez notablement de contenu sans cesser d'être la même ressource, un peu comme une personne peut bien changer de ville (d'adresse) mais aussi de profession (de rôle) voire de caractère (de contenu), sans cesser d'être la même personne⁵⁵. De plus, on souhaite pouvoir réutiliser ces éléments de contenu en leur donnant toutes les fonctions différentes que la logique rend possible⁵⁶ (ils doivent pouvoir passer de

⁵² Y compris en situation, dans le cas d'un déictique. Ce point est d'une importance toute particulière car il s'oppose aux deux extrêmes que constituent les conceptions philosophiques du nom propres et les scénarios envisagés à l'heure de l'Internet des objets, oscillant respectivement entre une distance et une proximité extrêmes (excessives) vis-à-vis du référent.

⁵³ Ou, selon l'interprétation tropiste, une compréhension de qualités.

⁵⁴ Cette conception de la qualification est analysée dans (Monnin, 2013b) au moyen des outils conceptuels élaborés dans (Livet & Nef, 2009). Elle correspond, au plan ontologique, au couplage d'un processus qualifiant (d'identification de la ressource – virtuelle – à partir de l'URI – actuelle), avec un second processus, actualisant le stade final, virtuel, du premier (ce second processus s'apparente à d'un processus de déréférencement, construisant l'accès aux représentations – actuelles – qualifiées par la ressource).

⁵⁵ Ceci n'exclut évidemment pas qu'au fil des opérations s'opèrent des changements plus profonds, aboutissant à d'autres modes d'individuation.

⁵⁶ Au plan logique, ce principe est au cœur de RDF, puisant ses sources dans Common Logic. Cf. (Monnin, 2013, pp. 182-183) : « RDF, à l'instar du Web, repose sur la possibilité d'associer des URI aux ressources, objets ou concepts (en fait de concepts, des prédicats polyadiques, des relations binaires). De la sorte, rien n'empêche une relation de figurer en position de sujet ou de prédicat dans un triplet, voire de s'appliquer à elle-même ($\langle S, P, O \rangle$, où S = la relation « est un », P = la relation « est un », O = le type « relation », soit $\langle \langle \text{« est un »} \rangle, \langle \text{« est un »} \rangle, \langle \text{« relation »} \rangle$). Pour ne pas violer les axiomes de la théorie standard des types la solution adoptée consiste à traiter les propriétés dyadiques tantôt comme des objets (types), tantôt selon leur extension (les objets qu'elles « contiennent »). ». Cette possibilité internalise à même des formalismes logiques des besoins évidents d'un point de vue technique à l'échelle du Web, auxquels ni OWL-DL ni OWL 2 ne satisfont pourtant : « des assertions logiques auxquelles on accède sous forme de documents RDF en déréférencant différentes URI (à différents « endroits » du Web), doivent, pour faire droit aux exigences du Web Sémantique, pouvoir être utilisées n'importe où. Ce qui est impossible avec la logique du premier ordre traditionnelle (TFOL) car, ce faisant, un même nom sera employé indifféremment pour désigner qui un objet, qui une propriété/relation, toutes choses que RDF ou

prédicat à sujet par exemple). On pourra donc utiliser un même nom pour remplir différentes fonctions, si bien qu'il contribuera à différentes opérations combinatoires ou de recherche de similarités tout en maintenant dans le même temps une relation avec un contenu stable.

III.3.3 - Tags

Enfin (et l'on entrevoit ici un nouvel usage des symboles renvoyant non plus seulement à des contenus ou des référents mais désormais à des processus) on peut adjoindre à un pointeurs les tags qui 1) donnent en raccourci une indication du contenu des ressources ; qui 2) peuvent indiquer en quoi ces ressources sont pertinentes pour telle ou telle enquête et à ce titre se comporter comme des étiquettes qui donnent du sens à des processus de mise en liaison entre diverses ressources ; et qui 3) peuvent enfin agir à la manière de signaux adressés à une communauté particulière d'utilisateurs du Web et désigner alors des éléments contextuels (Web social) et non plus seulement de contenu (architecture du Web).

Dans le cas précis du tagging, nous retrouvons d'ailleurs très directement la typologie triadique évoquée plus haut. La situation initiale implique, comme on l'a vu, trois éléments : un libellé, la ressource et l'URI de celle-ci, pas encore distingués. Pour ce faire, nous disposons à l'étape suivante de deux chaînes de caractères, la chaîne de caractères du libellé, qui le distingue des chaînes de caractères d'autres libellés, et l'URI de la ressource, la distinguant pareillement d'autres URI. La ressource, quant à elle, n'est pas encore distinguée. Par la suite, pour distinguer l'URI du tag, nous pouvons introduire une nouvelle entité, la relation⁵⁷. Elle-même n'est pas encore explicitée. En revanche, elle explicite le tag et distingue la ressource de l'URI. Enfin, l'introduction des actions de tagging (non explicitées) amène à expliciter cette relation. Ce qui exige, pour aller plus loin, d'introduire une ou plusieurs nouvelles relations, pas encore explicitées à cette étape, entre la relation initiale et ses types (les actes de tagging), eux-mêmes explicités désormais.

Il n'en demeure pas moins nécessaire d'exprimer le continuum allant de l'acte individuel au tag générique, en remontant ensuite du tag générique personnel au tag générique collectif. Le passage d'un niveau à l'autre est progressif, fruit d'une extension croissante des types d'opérations possibles, qui a la fois use de nouvelles distinctions et permet de se donner la liberté d'aller *ad libitum* dans chacune des directions. Par un processus inverse d'*indistinction*, il suffit ainsi de laisser progressivement de côté :

- a) les contraintes de cardinalité qui limitent le nombre de ressources assignées à une action de tagging ; et
- b) les relations entre libellés et ressources taguées ;

Common Logic autorisent, à l'inverse de TFOL. » (*op. cit.*, p. 289). C'est ce que Pat Hayes nomme le « principe de portabilité », élément constitutif de ce que devrait être selon lui une logique sensible aux spécificités du Web, et qu'il nomme une « Blogic ». Cf. l'introduction à ce volume. Le présent article entend rendre compte d'une dynamique dont les effets affleurent dans certains formalismes comme RDF ou Common Logic en délaissant toutefois la logique au profit de l'ontologie.

⁵⁷ C'est là l'une des principales caractéristiques de l'ontologie NiceTag élaborée par Alexandre Monnin, Freddy Limpens, David Laniado et Fabien Gandon, cf. <http://ns.inria.fr/nicetag/2010/09/09/voc>. Voir également (Monnin *et al.*, 2010).

pour que le fardeau de l'identité repose sur le seul signe, progressivement désémiotisé (réduit à une chaîne de caractères). Dans plusieurs schémas de tagging, en vertu de la propriété de l'ontologie SCOT⁵⁸, <scot:spelling Variant>, les distinctions orthographiques sont également susceptibles d'être ignorées, avec tous les risques que cela comporte (ex. : « Paris » et « paris » sont deux mots différents en français).

De ce mouvement résulte le passage du *tagging* à la *folksonomie* (personnelle puis collective), d'une action singulière clairement identifiée et individuée à des actes de tagging qui associent, pour des raisons non précisées (les relations faisant désormais défaut), une variété de ressources à un libellé. Il ne s'agit évidemment pas d'imposer l'ensemble de ces distinctions en toutes circonstances, sur un mode rigide. Il s'agit plutôt de les identifier, et d'en jouer par accroissement ou diminution en fonction de chaque situation⁵⁹.

III.3.4 - Distinguer et expliciter mais aussi « indistinguer » et « impliciter »

Au regard des exemples mobilisés, le trajet n'est donc pas unidirectionnel, allant toujours dans le sens d'une plus grande complexité. Si la complexité croissante des distinctions permet de spécifier avec un grain toujours plus fin le tag singulier, à l'inverse, pour en revenir à la folksonomie il est nécessaire de relâcher ces contraintes de manière à faire coexister simultanément différents niveaux ontologiques de grains variables. Il semble ainsi loisible d'enclencher un processus de développement ontologique non seulement de manière *ascendante*, en partant de distinctions grossières pour embrasser des distinctions toujours plus fines, mais également de manière *descendante*, en repartant, à rebours, des strates explicitées vers un niveau de plus grande indistinction – une manière de se montrer à nouveau plus libéral qu'à l'ordinaire vis-à-vis des catégories ontologiques, étant entendu que tous ces niveaux présentent une pertinence du point de vue ontologique.

C'est donc moins d'une ontologie d'entités *d'emblée distinguées* que nous éprouvons ici le besoin, que d'une ontologie opérationnelle de distinctions

⁵⁸ (Scerri *et al.*, 2008).

⁵⁹ Point d'autant plus crucial compte tenu du danger qui guette l'abondance excessive de distinctions. Voir en particulier l'exemple du langage de programmation 2-Lisp conçu par Brian Cantwell Smith (Smith, 1998, pp. 37-41), et la complexité induite par l'obligation de distinguer notamment les *chiffres* des *nombres* (les premiers étant l'expression des seconds). Contraintes beaucoup trop fortes dans la plupart des cas, insuffisantes dans les situations très spécifiques d'où elles étaient censées tirer leur justification. Smith en tire la morale suivante : « *It was soon clear that what we wanted, even if did not at the time know how to provide it, was a way of allowing distinctions to be made on the fly, as appropriate to the circumstances, in something of a type-coercive style – and also, tellingly, in a manner reminiscent of Heideggerian breakdown. Representational objects needed to become visible only when the use of them ceased to be transparent. Reason, moreover, argued against the conceit of ever being able to make all necessary distinctions in advance – i.e., against the presumption the original designer could foresee the finest-grain distinction anyone would ever need and thus supply the rest through a series of partitions or equivalent classes. Rather what was required was a sense of identity that would support dynamic, on-the-fly distinction or task-specific differentiation – including differentiation according to distinctions that had not even been imagined at a prior, safe, detached, “design” time. [...] If even arithmetic generated this much complexity that lends strong support to the idea that in more general situations it will be even more inadequate to treat objects as having stable purpose-independent identities without regard to the functions or regularities in which they participate.* »

successives et de distingueurs qui se trouvent progressivement qualifiés ; les distinctions qui distinguent des distinctions donnant lieu à des *explicitations* alors qu'à l'inverse, les opérations qui les renvoient dans l'implicite prennent les traits de ce que l'on nommera des « *indistinctions* » ou des « *implications* ». En définitive, le recours à une ontologie d'opérations autorise une meilleure prise en compte des distinctions locales, entées sur la réalité du Web, sans nécessiter l'import immotivé de catégories philosophiques non seulement stables mais aussi exogènes. Le point de départ retenu, dans le cas d'espèce qui nous concerne, découlant à la fois de l'analyse opérée par les architectes du Web dans les standards et des types qu'elle met en jeu⁶⁰. C'est l'autre enseignement important : il n'y a pas d'état initial absolu ou d'état zéro doté de primitive(s) originelle(s) mais des démarrages qui sont en fait des *redémarrages*, isolant certaines primitives dans des contextes plus ou moins locaux⁶¹.

À l'inverse d'une tentative visant à épuiser *a priori* la diversité des types, incarnée par l'analyse des systèmes philosophiques de Jules Vuillemin, systèmes universels car axiomatisés, multiples mais contradictoires, la dynamique bidirectionnelle du processus « d'explicitation » ne met en évidence aucun socle ontologique ultime reposant sur des entités statiques. Fort heureusement, d'ailleurs, car le Web lui-même n'est *pas* une « réalité ontologique ultime », bien plutôt un projet technique en constant devenir, à *faire*, dont l'architecture, pour demeurer stable, n'en est pas moins sujette à des agencements changeants, portés par des opérations distinguant (ou *indistinguant*) autant d'entités nouvelles. Celles-ci, à l'instar de la ressource, ne sont pas toujours appréhendées d'emblée de façon explicite. En atteste le rappel initial du processus de standardisation ayant abouti à faire émerger ces types ontologiques, où la

⁶⁰ Évidemment, le mot « analyse » comporte un danger immédiat du fait qu'il semble évacuer ce qui donne à cette situation ses traits spécifiques. Autrement dit, son caractère « performatif ». Les entités décrites dans les spécifications touchant à l'architecture du Web l'étant au moyen de standards, les seconds contribuent, selon des modalités toujours singulières, à performer les premiers (au sens de la « performation », terme emprunté à Michel Callon, que privilégie aujourd'hui la sociologie par rapport au concept de « performativité » venu d'Austin). Le défi consiste dès lors à rendre compte à la fois de ce qui diffère d'une posture philosophique traditionnelle, une fois réalisé le déplacement dans le monde de l'ingénierie, sans toutefois perdre de vue le caractère éminemment risqué – dénuée de toute garantie automatique – d'une telle opération (sous cet angle, le choix du mot « performativité » s'avère en effet particulièrement malheureux). S'ajoute à cela une difficulté supplémentaire, s'agissant cette fois-ci des critères permettant de reconnaître qu'un standard est « appliqué ». Ce point se fait tout particulièrement ressentir lorsque les contraintes qui découlent d'un standard s'exercent en dépit de sa relative méconnaissance (situation tout à fait analogue à celle que nous décrivons ici). Il est alors crucial de bien articuler dans la description ces deux aspects en apparence contradictoires pour tenter de comprendre ce qui les rend néanmoins compossibles.

⁶¹ Pensons au flux de l'expérience dont la description ouvre l'essai de William James intitulé « *The thing and its relations* » dans ses *Essays in Radical Empiricism*. Ce flux de l'expérience qui précède les distinctions est traversé de champs de forces, « de portions saillantes », sa pureté étant toute relative, ainsi que le note James lui-même. Néanmoins, l'établissement de distinctions ne fait pas disparaître le flux en tant que tel. Le face-à-face entre l'unicité de l'expérience pure et la diversité discrète des types abstraits, fixés depuis cette expérience pure, peut-être dépassé si l'on reconnaît a) que le flux n'est jamais amorphe (sa pureté est toute relative) et b) que les distinctions qui s'en extraient ne sont pas définitives au sens où elles en engendrent elles-mêmes de nouvelles. On dépasse ainsi l'alternative qui consiste à chercher à traduire un flux en devenir au moyen de catégories fixes. En réalité, les types distingués sont eux-mêmes en devenir en vertu des opérations récursives qu'ils suscitent autant qu'ils en résultent.

ressource figurait comme un type attendant moins d'être défini *a priori* que distingué et explicité. Enfin, les propriétés des entités considérées variant selon les agencements dans lesquels elles sont prises, il est tout à fait possible d'envisager des chaînes d'opérations divergentes à partir d'une « même » entité/primitive.

Une vision inspirée d'une certaine sociologie des usages manichéenne⁶² pourrait être tentée d'opposer à notre propos l'opposition classique entre *concepteurs* et *usagers*. La réflexion menée jusqu'à présent doit nous donner les moyens de ne pas entériner une alternative aussi brutale. Bien sûr, bâtir un réseau de localisateurs exigea des ingénieurs du Web qu'ils définissent un certain nombre d'entités ou de relations. Mais si, du point de vue de ces ingénieurs, ces types ont été rendus explicites, ce n'est pas ainsi qu'ils fonctionnent sur le Web – aux yeux des utilisateurs. Au contraire, ils demeurent une part cachée et indistinguée du Web, autorisant au passage une relance du processus de distinction et d'explicitation à partir de types flottants, dans de nouvelles directions jusque-là imprévues. Et ce, sans que les utilisateurs aient eu besoin d'avoir un accès préalable à cette part, à leurs yeux indistinguée, de l'architecture du Web.

La conception « anti-fétichiste » de l'usage (selon l'expression d'Antoine Hennion et Bruno Latour, (Hennion & Latour, 1993)) entend opposer aux prétentions des concepteurs, réputées centrées sur les objets, la « réalité » des pratiques. En nous penchant sur la « métaphysique empirique » des concepteurs du Web, ne prenons-nous pas le risque de la voir vidée de sa substance par une analyse ultérieure des usages dont les résultats ne sauraient aboutir aux mêmes conclusions que les nôtres ? Sans prétendre le moins du monde rendre compte de tous les détournements possibles ou de l'ensemble des « usages sociaux »⁶³, remarquons tout de même qu'en insistant sur la pluralité des niveaux distingués, plusieurs degrés d'appréhension, plus ou moins stabilisés, jamais définitifs, cohabitent. On notera qu'une partie de l'innovation ayant conduit au tagging et au *social bookmarking* a consisté dans un premier temps à faire émerger et à mobiliser sous forme d'application des entités (libellés, URL, etc.) maintenues à un certain niveau d'indistinction (Muxway) avant d'en généraliser le principe (Delicious). À partir de quoi, les ingénieurs du Web, au vu du succès rencontré par ces dispositifs, se sont *ensuite* évertués à les expliciter de manière à créer un pont entre le Web social et le Web Sémantique⁶⁴. Par-delà les détournements, « indistinguer » revient ainsi à innover, et ce doublement, par la relance ultérieure d'un nouveau cycle d'opérations de distinctions dont on ne sait *a*

⁶² « Quant aux sociologues des usages, nombre d'entre eux tendent à penser que ce qui se passe avant l'entrée en scène de l'usager ne les concerne pas dans les détails, et peut être appréhendé simplement en considérant que cette phase conduit à établir certaines caractéristiques techniques dont le sens et les effets seront redéfinis dans l'appropriation », (Mallard, 2011, pp. 259-260, in Denouël & Granjon, 2011).

⁶³ (Mallard, 2011), *op. cit.* Cf. également, pp. 259-260, note 2. Les usages sociaux s'opposeraient à « l'utilisabilité », plus servile, dont la sociologie de l'acteur-réseau ne parviendrait guère à s'extraire. Le cadre dont nous donnons ici l'esquisse nous semble à même d'échapper en grande partie à un tel dualisme.

⁶⁴ Outre NiceTag, on dénombre plus d'une demi-douzaine d'ontologies du tag (!) : celle de Newman, SCOT, SIOC, ES, UTO, Semdrop, Tagont, MOAT, Common Tag, Tagora, NAO (*Nepomuk Annotation Ontology*), MUTO (*Modular Unified Tagging Ontology*), Lexitags, etc.

priori où il mène. Dans ces conditions, réserver une place aux usages n'exige ni de passer les objets sous silence, ni de neutraliser la technique en guise de préalables indispensables, dans la mesure où l'ontologie examinée jusqu'à présent ne nous met nullement aux prises avec des « entités statiques, définies une fois pour toutes ». On cherchera alors vainement toutes traces de déterminisme, toute intention réalisée des concepteurs et, plus généralement, tout échelon fondamental⁶⁵.

Ici sans doute se marque l'écart le plus grand entre la philosophie et le Web, la quête philosophique des niveaux les plus élémentaires étant conçue sur le mode d'un mouvement de retour effectué en direction des entités les plus fondamentales, alors même que le mouvement concomitant, sur le Web, conduit à relâcher des contraintes formelles trop strictes pour ouvrir de nouvelles potentialités encore inexplorées, tant du point de vue des utilisateurs que des concepteurs.

IV – UNE ONTOLOGIE À LA MESURE D'UNE SÉMANTIQUE MULTI-DIMENSIONNELLE

On le voit, s'agissant de l'ontologie du Web, les choses sont ici plus complexes que ce que laisserait accroître toute focalisation exclusive sur les ontologies du Web Sémantique. Ces complexités nouvelles, entrevues au cours de la seconde partie, ont pu tenir à une double volonté des concepteurs et constructeurs successifs du Web : 1) de permettre à chacun de proposer des informations de manière libre, et 2) de donner à chacun la possibilité de retrouver chacun de ces contenus. Les deux exigences peuvent rentrer en conflit, lorsque chacun choisit à sa guise la forme associée à la présentation de son information et son mode d'accès. Le danger que tous ne puissent plus accéder aux informations est alors contenu par des régulations portant sur la constitution des adresses et sur les procédures d'accès. Le danger que l'on doive couler ses

⁶⁵ Brian Smith note l'apport des nouvelles pratiques débordant le cercle des informaticiens : « *Modern practice is bursting with possibility, as designers, playwrights, artists, journalists, musicians, educators, and the like are drawn into the act, along with the original scientists and engineers, and now also anthropologists, linguists, and sociologists. [...] it would be a mistake to think that these people are just users of computation. On the contrary they are participating in its invention – creating user interfaces, proposing architectures, rewriting the rules on what it is to publish, disrupting our understanding of identity.* » (Smith, 1998, pp. 359-360). L'allusion à la modification des règles de publication se justifie tout particulièrement au regard du tagging, n'étant rien d'autre qu'une tentative de compenser la perte de la dimension d'écriture inhérente au Read-Write Web des origines au moyen de formes d'annotations, très simples en apparence, qui se sont progressivement affinées au fil du temps. On comprend bien pourquoi la métaphysique sous-jacente à ces considérations se veut « participative » – toujours selon l'expression de Smith. Donner droit de cité à la métaphysique empirique des acteurs ne saurait donc se concevoir sans porter attention aux opérations qu'ils accomplissent (celles des architectes mais aussi celles des utilisateurs). La tâche des ingénieurs du Web consiste à bâtir des structures aptes à favoriser l'expansion d'une dynamique qui rappelle les enchaînements du processus (partiellement épistémique) d'explicitation entrevu jusqu'ici. Une fois ce point acquis, on pourra bien qualifier la position défendue jusqu'ici de « réalisme symétrique », à l'instar de Smith : « *I will call symmetrical realism – a construal of (non-naïve) realism that not only establishes some of the background assumptions or metaphysical preconditions on the existence of objects but places equally strong preconditions on the existence and nature of subjects, including on their epistemic achievements, with particular reference to the recalcitrant notion of objectivity.* [nous soulignons] » (Smith, 1998, p. 85).

informations dans un seul format rigide est conjuré par l'élasticité du lien entre un pointeur et le contenu d'une ressource, celui-ci étant susceptible de variations considérables.

Ce sont de bonnes intentions, mais peut-on mieux relier le programme d'une ontologie au sens philosophique du terme, avec les réquisits de cette ontologie dynamique que nous avons reconnue l'œuvre à travers le Web, et donc avec la combinaison de ces deux exigences et les problèmes complexes qu'elle soulève ? Montrons qu'une extension du schéma proposé par Robert Stalnaker sous le nom de « sémantique bi-dimensionnelle », extension à la multi-dimensionnalité, peut fournir un cadre adéquat pour commencer à penser la complexité du Web et les difficultés afférentes.

IV.1 - Quelle sémantique pour les URI ? Retour sur la sémantique bi-dimensionnelle

Au regard de cette ontologie singulière du Web, nous suggérons donc, pour rendre compte de la sémantique des URI et l'ontologie de la ressource, de repartir du schéma formel dont on dispose grâce à Stalnaker (1978). Schéma dont il s'est servi pour proposer une sémantique dite « bi-dimensionnelle » (*two-dimensional semantics*). Celle-ci a ouvert des discussions nombreuses et très techniques qui serviront avant tout de marchepied à nos propres analyses.

Ce schéma permet tout à la fois saisir des situations semblables à celles où le mot « eau » désigne H₂O sur Terre et une composition chimique XYZ sur Terre Jumelle, et celle où le même sens est accordé aux indexicaux (« je », « ici », « maintenant ») alors qu'ils peuvent désigner différentes personnes, lieux ou instants selon les situations.

Montrons comment cela fonctionne pour « je ». On construit une matrice. En *ligne* est indiqué le contexte – ou monde possible – dans lequel quelqu'un utilise « je » pour parler. Autrement dit, le contexte permettant de préciser *qui* est visé par « je » à ce même instant et dans cette situation (l'homologue du « sur Terre » dans le tableau). Dans chaque *colonne* est indiqué un contexte, ou monde possible – qui peut être identique à la situation initiale ou différent d'elle – dans lequel est visé au moyen d'une phrase le référent du terme « je » de la situation qui correspond à telle ou telle ligne (l'homologue de « vu de Terre », cf. tableau 5 *infra*).

Que ce soit Paul ou Jacques qui parlent de celui qui dit « je », c'est toujours Pierre qui est visé si Pierre était le « je » du contexte d'usage indiqué par sa ligne. En revanche, si Paul vise un « je » qui était utilisé dans un autre contexte d'usage, il ne s'agira plus de Pierre mais par exemple de Paul lui-même ou de Jacques. Le référent dépend donc du contexte d'usage, qui peut être celui où Paul ou Jacques utilisaient « je ».

Or, si les colonnes sont ordonnées correctement par rapport aux lignes, et si l'on se place sur les cases de la diagonale de la matrice, « je » va :

- pour le contexte d'usage ⟨Pierre, monde possible a_i , temps t_i ⟩ et les contextes où lui ou un autre vise le même « je » de ce contexte d'usage, désigner Pierre ;

- pour le contexte d'usage $\langle \text{Paul, monde possible } a_2, \text{ temps } t_2 \rangle$ et le contexte où l'on vise son « je », désigner *Paul* ;
- et de même, *mutatis mutandis*, pour Jacques.

Autrement dit, « je » va toujours désigner ce qu'on peut appeler « le locuteur de la phrase utilisant « je » dans la situation en question ». Si l'on note par « Vrai » les cas où ce dont on parle dans le contexte de visée rencontre bien le référent du « je » du contexte d'usage, on mettra un « Vrai » dans toutes les cases de la diagonale. Le référent de « je » est *sensible au contexte d'usage* mais le rôle sémantico-pragmatique de « je » est *toujours de désigner le bon référent si l'on se place dans le bon contexte d'usage*.

	Pierre	Paul	Jacques
$\langle \text{Pierre, } a_1, t_1 \rangle$	Pierre	Pierre	Pierre
$\langle \text{Paul, } a_2, t_2 \rangle$	Paul	Paul	Paul
$\langle \text{Jacques, } a_3, t_3 \rangle$	Jacques	Jacques	Jacques

Tableau 4
Matrice bidimensionnelle pour « je »

Plus généralement, la proposition « en ligne », dite C-proposition, est celle qui, pour un monde possible, nous donne un référent (et par là une valeur de vérité). Selon (Stalnaker, 2004), elle est également celle que nous utilisons dans la communication ordinaire en supposant que le monde en question est le monde actuel ou du moins le même monde à la fois pour le locuteur et pour le destinataire. On peut formuler la chose soit en disant qu'on a une fonction des mondes possibles vers des propositions [cf. colonnes], une proposition étant elle-même une fonction des mondes possibles vers des valeurs de vérité (par une sorte d'emboîtement) [cf. lignes] ; soit une seule fonction à partir de deux mondes possibles (celui du locuteur et celui de l'auditeur) vers une valeur de vérité.

	Vu de Terre	Vu de Terre Jumelle
$\langle \text{Sur Terre, } a, t \rangle$	H₂O	H ₂ O
$\langle \text{Sur Terre Jumelle, } a, t \rangle$	XYZ	XYZ

Tableau 5
Matrice bidimensionnelle pour « eau »⁶⁶

Ce dispositif résulte en apparence d'une simple combinatoire, si bien qu'on peut l'utiliser pour toute relation entre deux contextes. Il est également normatif, ou plus exactement soumis aux contraintes du rôle sémantique à l'origine de la

⁶⁶ Le long de l'axe vertical sont inscrits des mondes possibles centrés autour d'un agent a et d'un temps t inhérents à chaque monde.

définition de la relation en question, puisqu'il exige, pour qu'on puisse apposer la valeur « Vrai » dans chaque case de la diagonale (ce pourrait être aussi la valeur « succès », mieux adaptée au cas d'espèce du Web), que les colonnes et les lignes aient été bien ordonnées. Or cette ordonnance dépend des normes propres au rôle sémantico-pragmatique de la relation en question (dans l'exemple du Tableau 4, le rôle du locuteur de la phrase ; dans le Tableau 5, la composition d'une substance)⁶⁷.

IV.2 - Vers une sémantique multi-dimensionnelle

Si l'on considère seulement l'aspect combinatoire, il est évidemment formellement possible de ne pas se borner à deux dimensions, et de construire des relations et des matrices *multidimensionnelles*. C'est d'ailleurs nécessaire pour comprendre ce que fait le Web, la même « adresse » (identifiant ou localisateur) pouvant renvoyer à des contenus non seulement différents mais susceptibles qui plus est de varier dans le temps tout en ressortissant à une même référence, à une même *ressource*. On obtient ainsi un « vecteur d'information » qui est une fonction non pas de deux mondes possibles mais de n mondes (que cette multi-fonction soit déterminée selon des convergences ou des emboîtements). Il semble que la notion « d'accès », sur le Web là encore, corresponde à une telle multi-fonction. Elle est également adaptée aux réseaux dont les liens peuvent être différenciés au moyen de labels, à l'instar du Web (ce qui permet d'effectuer des recherches sur différentes ressources eu égard à un thème ou à une question, comme en témoigne l'exemple du tagging examiné plus haut, et, plus largement, toute qualification d'un pointeur/localisateur par un littéral, cas de figure typique des scénarios du Web social mais aussi du Web de données).

Dans le cas des contenus variés, voire fluctuants, d'une ressource, si l'on commence à deux dimensions, on pourra associer à la première, *en lignes*, les différentes composantes de cette ressource (au sens d'une représentation ponctuelle) à différents moments, et à la seconde, *en colonnes*, les requêtes qui utilisent le pointeur de la ressource. Le Web est supposé être construit de telle manière que chaque requête utilisant le bon localisateur tombe *au mieux* sur les différents composants de la même ressource (qui peuvent varier à différents moments).

Évidemment, il devient possible d'augmenter le nombre de dimensions, si l'on admet par exemple que l'on puisse arriver à une même ressource *via* des localisateurs variés. Cette variété elle-même constituera dès lors une dimension supplémentaire⁶⁸. On peut aussi utiliser différents tags, qui, aux yeux d'un utilisateur, seront plus évocateurs de tel ou tel aspect d'un composant du contenu de la ressource. Là encore, il faudra ajouter une ou plusieurs dimensions

⁶⁷ La sémantique bi-dimensionnelle généralisée suppose que chaque locuteur peut connaître *a priori* ce rôle, mais le formalisme matriciel implique seulement qu'on puisse trouver un ordre des lignes et des colonnes qui assure d'apposer « succès » dans chaque case de la diagonale, autrement dit qu'on puisse définir une condition d'auto-stabilité, une forme de point fixe. Une fois qu'on a trouvé cette stabilité, elle est notre repère mais il faut d'abord la trouver ; on n'en dispose pas *a priori*.

⁶⁸ Pour des raisons évidentes de lisibilité, nous ne représentons pas ces multiples dimensions sur un même schéma mais en multipliant leurs variétés.

supplémentaires. Ce qui, *in fine*, diffère de la sémantique usuelle. Celle-ci entend en effet arriver à *un* référent sous *un* concept et, par-là, à *une* valeur de vérité.

Par contraste, de façon à permettre d'accéder aux éléments d'une ressource, y compris à son état vide le cas échéant, ce qui compte sur le Web c'est seulement que *si l'on part d'une ressource*, on puisse disposer d'un faisceau de composants accessibles grâce à des identifiants, composants qui, sans être identiques les uns aux autres, *conviennent* bien à ladite ressource. Toujours dans ce sens « ressource vers adresses », on peut également ajouter une dimension supplémentaire en citant la possibilité que tel élément d'une ressource (une représentation ponctuelle) soit suffisamment spécifique en termes de contenus ou encore de formats, pour susciter un ou plusieurs tags non-pertinents pour qualifier d'autres éléments appartenant à la trajectoire de cette même ressource (bien que tenus pour fidèles à la même ressource, ils ne sont en effet nullement identiques aux précédents).

	Pointeur 1	Pointeur 2
⟨Paramètres de variation adressés 1, a , t ⟩	X	0
⟨Paramètres de variation adressés 2, a , t ⟩	0	X

Tableau 6
Ressources vers pointeurs⁶⁹ : se signaler

En partant de la ressource, la coordination recherchée s'accomplit sur le mode de la « qualification »⁷⁰. Des représentations sont ainsi qualifiées à partir du moment où l'on peut considérer à bon droit qu'elles s'intègrent au trajet d'une ressource (selon des critères susceptibles de varier en fonction des modalités suivant lesquelles des représentations sont générées, modalités elles-mêmes associées à des paramètres précis. Ces critères ne sont nullement livrés *a priori* : on retrouve ici le sens normatif, non-formalisable, de l'idée de régularité). Les composantes de la ressource sont dès lors évaluées en fonction de leur fidélité à celle-ci.

Les *lignes* ci-dessus représentent les paramètres ou modalités en vertu desquels les représentations se différencient les unes des autres dans un contexte a (tel ou tel *publisher* ou autorité) et à un instant t . Les *colonnes* représentent l'autre extrême du processus, à savoir les URI qui conditionnent l'accès aux représentations (qui donnent accès aux états représentationnels, modalisés selon certains paramètres). Cette schématisation correspond de manière quasi-triviale au paramétrage des entêtes Http (*Http Header Field*) d'un serveur, bien qu'elle se situe à un niveau d'abstraction qui la rend compatible avec d'autres solutions techniques.

La *ressource* s'apparente ici à un trajet inaccessible en tant que tel (donc, en un sens, « abstrait ») ; les *représentations* à des éléments concrets mais fluents ; les *URI* à des éléments concrets et stables⁷¹. Le sens « ressource vers adresse » part ainsi de l'abstrait pour aller au plus concret.

⁶⁹ Notons que les paramètres de variations ne sont pas tous nécessairement adressés.

⁷⁰ Au sens de (Livet & Nef, 2009), cf. *supra*.

⁷¹ Cette gradation est la suivante : Ressource (trajectoire abstraite) → Représentations (éléments concrets transitoires) → URI (identifiants concrets stables).

En outre, il convient de noter que tous ces schémas, qui en apparence conservent un caractère bi-dimensionnel (la multiplication des tableaux illustrant le caractère désormais multidimensionnel de l'analyse), devraient en réalité être élargis au moyen d'une représentation sous forme de *tenseurs*⁷². Aux matrices s'ajoute un axe temporel supplémentaire, dans la mesure où les paramètres de variations et les adresses associées sont susceptibles de changer au fil du temps.

Dans l'autre sens, comme l'exemplifie la négociation de contenu, on partira d'une URI à laquelle un ensemble d'autres URI seront reliées par redirection, afin, par exemple, de rendre accessibles des états représentationnels de ressources plus spécifiques. Ces représentations devant néanmoins être considérées comme qualifiées tant par les ressources (proximales) auxquelles elles « appartiennent » immédiatement, qu'avant redirection (ressources distales). L'URI de la Bible du roi James redirigera par exemple, en vertu de la négociation de contenu, vers d'autres URI identifiant tantôt la-Bible-du-roi-James-en-anglais, la-Bible-du-roi-James-en-anglais-et-HTML, voire la-Bible-du-roi-James-en-français (l'ensemble des ressources ainsi associées à la première contribuant à définir son identité, selon des trajectoires différentes ; ici, la Bible du roi James est *de facto* assimilée non à « un texte » mais à une « œuvre », ce qui diffère du tout au tout⁷³). Par conséquent, *si l'on part des pointeurs ou localisateurs*, on dispose d'un faisceau de pointeurs qui permettent d'activer différentes représentations d'une ressource (ou de ressources associées à la première) au gré de ses (leurs) variations spécifiques.

	(Composantes de la ressource x, t)	(Composantes de la ressource y, t)
Pointeur 1	x	0
Pointeur 2	0	x

Tableau 7
Pointeurs vers ressource : *déréférencer*

En partant cette fois-ci des adresses ou identifiants, sont visées les représentations auxquelles on accède à partir d'une URI. En vertu de la négociation de contenu, qui ne pose nulle obligation à ce que soit mise en place une redirection vers de nouvelles URI, cette relation ne saurait être fonctionnelle car elle peut déboucher, ponctuellement, sur une multitude de représentations.

« *Déréférencer* », c'est simplement assurer la communication par un faisceau dans le sens *pointeur vers ressource*. Il faudrait forger un terme pour le sens *ressource vers pointeur*, par exemple « *se signaler* ». Cela nous amène, comme

⁷² Si la ressource s'apparente à une règle, encore faut-il en souligner le caractère génératif et créateur. C'est d'ailleurs cette dimension créatrice que Whitehead associait aux tenseurs, qui jouent un rôle central dans sa métaphysique, comme de nombreux commentateurs l'ont d'ailleurs noté. Autant H2O et XYZ sont des référents stables, autant les éléments pertinents dans ces schémas dessinent une trajectoire dynamique dans ce que l'on appelle ici un « *faisceau* ».

⁷³ Cette notion fait directement écho aux modèles documentaires comme FRBR (*Functional Requirements for Bibliographic Records*). Toutefois, une discussion théorique sur les conditions d'individuation d'une œuvre appelle également, très vite, leur dépassement. La remarque qui précède ne préjuge donc nullement de notre adhésion à ces modèles conceptuels.

on l'a dit, au réseau qui peut relier des localisateurs et des ressources par l'intermédiaire d'une multitude de libellés formant *différents sous-réseaux*. Ce sont autant de relations de similarités différenciantes entre les contenus des ressources qui, lorsqu'elles permettent de situer suffisamment d'éléments les uns par rapport aux autres, jouent par-là même le rôle de dimensions de plein droit – à condition qu'elles présentent quelque indépendance entre elles et ne soient pas réductibles à des combinaisons d'autres similarités.

Tout ceci explique qu'il n'y ait ni fausseté ni négation dans cette affaire, et donc, à proprement parler, pas non plus de « vérité » puisque le « vide » en guise de résultat (de « déréférence » si l'on peut dire) conserve l'unicité du faisceau. Il semble alors que supposer une unité de sens ou de référent (comme « objet ») du côté de la ressource, dans le cadre du Web Sémantique, articulée autour d'un seul mode de représentation (à l'aide de langages de représentation des connaissances), en opposition à la diversité de modes d'accès⁷⁴, voire en l'offusquant purement et simplement, est en fait tout aussi dommageable et inexact que de supposer une stricte unité d'identifiant (*uniform*) du côté de l'adressage.

L'idéal du Web vise plutôt à garantir une grande liberté *du côté de la fonction d'accès* (le choix des pointeurs que l'on peut ensuite associer à des tags variés), et *du point de vue de la variété du contenu des ressources* (les modalités qu'empruntent les représentations : images, textes, vidéos, services interactifs, etc.). À condition de toujours demeurer en capacité de relier ces représentations *correctement* à la ressource, ce qui suppose de se trouver sur la zone de concordance des multiples dimensions (l'homologue de la diagonale) ; la valeur de vérité étant alors remplacée par la réussite ou la convergence des accès (leur appartenance à un même faisceau)⁷⁵. En un sens, le Web revient à admettre une

⁷⁴ Il ne s'agit évidemment pas, derrière la problématique des modes d'accès, de réintroduire le *Sinn* frégeén, où, à suivre la lecture de Claude Imbert, tout converge sur un mode kantien, la représentation de l'objet étant calculée selon un répertoire fini de catégories. (Monnin, 2013b) mobilise le concept d'instauration dû à Etienne Souriau afin d'éviter cet écueil. Il s'agirait plus précisément ici d'opérer un parallélisme entre le pluralisme des *modes d'instauration* de l'objet, d'une part, et la pluralité des *modes d'accès* à la ressource, d'autre part. Dans l'un et l'autre cas, la question du *vrai* se déplace bien du côté de la *réussite*, dont les critères ne sont nullement donnés d'avance.

Ceci permet également de proposer une alternative à l'idée trop simple de « réalisation physique » d'un objet informationnel. Ce patron de conception (*design pattern*), très utilisé dans l'ingénierie ontologique, a également été mobilisé lors des débats autour de la crise d'identité de Web afin d'introduire une distinction nette entre ressources « informationnelles » et « non-informationnelles » (soit, entre documents numériques et objets). Outre les difficultés liées à cette conception, discutées dans (Monnin, 2013b), on notera qu'elle s'appuie sur le présupposé d'un accès ponctuel aux objets *in toto* (leurs « réalisations »), présupposé éminemment discutable. L'interaction avec la représentation d'une ressource correspondant à un objet numérique n'est aucunement incompatible avec l'idée selon laquelle l'objet, dans sa plénitude, s'avère inaccessible. À l'inverse, elle est compatible avec le concept d'instauration dont elle souligne l'exigence (le succès d'une instauration se mesurant, dans ce cas précis, à l'aune de la possibilité d'interagir avec la *re-présentation* d'un objet – rendu ponctuellement présent sur fond d'absence). Sur ce point, voir (Smith, 1998).

⁷⁵ La *réussite* de chaque variation doit s'entendre dans la perspective de la *convergence* ou de la *cohérence* des variations entre elle, en dépit de leurs différences. Si l'on reprend l'idée selon laquelle la ressource est une règle, il convient de penser le déplacement que lui fait subir la *diagonale* conçue comme faisceau de localisateurs ou de composantes de ressources. Ce point n'est pas en contradiction avec l'idée précédemment avancée, simplement on insiste désormais sur une dimension potentiellement plus globale

sémantique où le contenu de la ressource demande d'explorer tous les « mondes possibles », soit, ramené à ce cadre précis, *tous les modes d'accès* – et ce dans les deux sens : des pointeurs vers la ressource et réciproquement. On passe ainsi d'une réflexion sur les mondes possibles ; usuelle dans l'orbe de la philosophie analytique, au pluralisme d'un monde unique, ce qui est tout à fait fidèle à l'esprit du Web – à condition d'éviter une conception de l'universalisme en vertu de laquelle l'unité et l'uniformité des entités visées serait *présupposées*.

IV.3 - Vers une sémantique multi-dimensionnelle à l'échelle du Web

Que devient, au gré de la multiplication des relations, la propriété d'auto-stabilité de la diagonale ? Pour certains couples de dimensions, il sera possible de définir une diagonale. Cela dit, on ne pourra évidemment pas le faire sur l'ensemble des dimensions. Il sera toutefois nécessaire que les différentes diagonales obtenues puissent former ce que l'on a commencé à nommer un « *faisceau* » (local, pour commencer).

En effet, pour continuer à parler d'une ressource et pour continuer à la trouver sur le Web, il faut bien que partant d'une « adresse », d'un label, ou d'une requête formulée à propos d'un ou plusieurs thèmes, on puisse toujours accéder à l'un des composants au moins de ladite ressource, et que ce composant soit en mesure de nous donner la situation de la ressource à laquelle il se rattache. Autrement dit, il faut que, partant de tout identifiant, tag ou libellé qui soit pertinent au regard de la ressource, on soit en mesure de converger vers celle-ci. Qu'intervienne, autrement dit, à un moment ou un autre dans le processus du Web, une *convergence* entre ces différentes voies. Et il faut aussi que les éléments qui participent d'une même ressource puissent se redistribuer en éventail (« *fan out* ») pour *rayonner* sur différents contenus. Pour que le Web ait du sens ou puisse servir d'échangeur pour des interactions sociales, il faut que tout contenu associé à une ressource renvoie à un identifiant, à un libellé ou à une requête, ce qui exige également d'associer convergence et rayonnement. Nous aboutissons donc à mettre en évidence deux processus, chacun partant d'une multiplicité, assurant une certaine unité et parvenant à une autre multiplicité, qui vont en sens inverse l'un de l'autre. Il n'est pas exigé qu'intervienne une correspondance terme à terme entre chacun des éléments de chacune de ces deux multiplicités mais à tout le moins que chaque processus fasse converger le faisceau à un moment donné, même si c'est pour ensuite le refaire diverger ou rayonner (et que cela marche ainsi dans les deux sens).

De manière significative, on peut également stratifier la notion de faisceau elle-même, sans oublier le *périmètre* de la convergence ou du rayonnement qu'elle implique, selon l'échelle adoptée. Ainsi s'agira-t-il, au niveau d'un même *publisher* (responsable d'un nom de domaine et détenteur de capacité de forger des URI qui en découle), de gérer l'accès aux représentations Http

alors que l'architecture du Web portait davantage sur un scénario centré sur la publication de contenus rendus accessibles par le truchement d'un *publisher* au titre d'une ressource. Le faisceau d'adresses/versions deviendrait dans ces conditions une métarègle, un *pattern* repérable à l'échelle du Web. La question de l'appartenance à un même faisceau se posant dès lors au-delà des variations contrôlées par un seul et même organisme.

(s'assurer qu'elles sont fidèles à une ressource par-delà leur variété, mettre en œuvre les redirections éventuelles requises en cas de négociation de contenu, etc.).

Dans l'éventualité où l'on aurait affaire à des *publishers* différents, on ne saurait compter sur un *faisceau unique de versions*. À titre d'exemple parlant, mentionnons les nombreux articles de presse publiés sur le portail Yahoo! et qui reprennent quasi à l'identique des contenus édités ailleurs. En dépit de cette coordination contractuelle entre des autorités distinctes, des différences importantes demeurent : les mises à jour et autres corrections ne sont pas, peu ou prou répercutées sur Yahoo!, les commentaires déposés par les lecteurs diffèrent, de même que leurs parcours ou les tags qu'ils posent (au moins en partie) et les réseaux qu'ils forment en définitive. En résumé, des coupes ponctuelles, bien qu'initialement quasi-identiques, s'inscrivent dans des trajectoires hétérogènes qui creusent des écarts de plus en plus grands. Le concept d'identité sur le Web doit donc être repensé du point de vue des rapports qu'entretiennent les trajectoires qualifiantes (les ressources) à l'égard de ce qu'elles qualifient, et non à partir des seules coupes ponctuelles au travers desquelles elles se donnent à voir ponctuellement (les représentations).

	Pointeurs de la ressource 1	Pointeurs de la ressource 2
⟨Faisceau 1, <i>a</i> , <i>t</i> ⟩	x	0
⟨Faisceau 2, <i>a</i> , <i>t</i> ⟩	0	x

Tableau 8

« Diagonale » de la ressource à l'échelle globale (faisceaux de composantes)

On notera que les cases de cette matrice équivalent chacune aux matrices présentées dans le Tableau 6 (il s'agit donc d'une matrice de matrices).

	Composantes de la ressource 1	Composantes de la ressource 2
⟨Faisceau 1, <i>a</i> , <i>t</i> ⟩	x	0
⟨Faisceau 2, <i>a</i> , <i>t</i> ⟩	0	x

Tableau 9

« Diagonale » de la ressource à l'échelle globale (faisceaux de pointeurs)

On notera que les cases de cette matrice équivalent chacune aux matrices présentées dans le Tableau 7 (il s'agit donc d'une matrice de matrices).

Quant au *faisceau de pointeurs* (identifiants, localisateurs), le déterminer témoigne de la difficulté que le Web Sémantique entend précisément résoudre : *comment s'assurer que des URI ressortissant à des noms de domaines – et de ce fait à des autorités – différents, renvoient bien au même objet ?* La solution proposée par défaut dans ce cadre est outillée au moyen de la relation `<owl:sameAs>`, destinée à poser de manière déclarative, par des moyens

manuels ou automatiques, que deux ressources sont « identiques ». En l'occurrence, cette solution s'appuie généralement sur une neutralisation implicite du problème en présupposant qu'existe un objet qui remplisse le rôle de référent unique – exigence en parfait accord avec l'emploi d'une relation d'identité au sens strict. Dans les faits, son usage épouse des nuances bien plus variées, comme l'ont montré les travaux de (Halpin *et al.*, 2010), en complet décalage avec des définitions trop strictes.

Les différentes échelles concernées semblent affublées de caractéristiques pour le moins hétérogènes. Distinguons ainsi :

- 1) l'idée d'un *faisceau de versions à l'échelle d'un publisher*, ce qui constitue déjà la problématique au cœur de l'architecture du Web, soit la coordination entre des URIs, les ressources qu'elles identifient et des représentations que ces dernières qualifient ;
- 2) l'idée d'un *faisceau de pointeurs à l'échelle du Web*, qui évoque plus immédiatement le Web Sémantique⁷⁶ (ce n'est donc pas un hasard si ce dernier point nous renvoie également aux URC) ;
- 3) un *faisceau de versions et de pointeurs à l'échelle du Web* ;
- 4) un *faisceau de versions, et/ou d'adresse, et/ou de tags, et/ou de libellés à l'échelle du Web*⁷⁷.

En fait, il n'y a pas de rupture complète entre chacun de ces niveaux. Songeons que les localisateurs, dans le cas de figure n° 3, peuvent donner une idée de la ressource liée à son contenu, à son *publisher*, voire, au niveau suivant (4), un éclairage sur les autres relations qu'elle entretient avec une communauté ou encore un sous-réseau, tous ces éléments étant compatibles avec les échelons 1 et 2. Inversement, même au premier échelon, on ne peut supposer une symétrie parfaite entre la structure du faisceau partant des ressources et celui qui part des pointeurs, à l'instar du niveau le plus global.

Typiquement, la coordination postulée est extrêmement forte lorsque l'identité entre ressources est notifiée par l'entremise de la relation `<owl:sameAs>`. On s'engage alors, en effet, à fournir bien davantage que des définitions identiques, en dépit de timides tentatives en ce sens qui ont toujours échoué – ce qui, au demeurant, ne serait guère contraignant. Évidemment, si l'on s'en tient aux seuls identifiants et à leur caractère *uniform*, l'étape de la coordination, la seule véritablement significative, s'en trouve escamotée. L'identité sur le Web reflète en effet un *pari* bien plus incertain, à savoir *que les trajets des variations associées à différentes ressources expriment une compatibilité vis-à-vis de ressources dont ils ne dépendent pas directement*.

⁷⁶ On l'a dit, l'identité entre URI est en effet le levier dont se sert le Web Sémantique à travers la relation `<owl:sameAs>`, pour affirmer l'identité entre deux ressources. Toutefois, ce scénario ne tient pas compte des composantes accessibles de la ressource. À raison, car la problématique de l'identité ne les concerne pas *stricto sensu*. Néanmoins, il nous semble important de poser la question de l'identité des ressources sans offusquer la place occupée par leurs composantes. En cela, l'approche que nous proposons est plus directement ancrée dans la réalité du Web qu'une approche essentiellement centrée sur les formalismes ontologiques du Web Sémantique, de RDF à OWL.

⁷⁷ Notons que le niveau 4 est tout aussi bien susceptible d'étendre que de limiter la portée du niveau 3 qui le précède.

Cette condition posée, il faudrait donc disposer d'une autre diagonale que celle considérée jusqu'ici car, à supposer que des faisceaux soient « suffisamment identiques », celle-ci ne saurait en toute logique marquer leur *pertinence* vis-à-vis des ressources auxquelles ils se rattachent directement (leur *fidélité*, oui, mais ceci suppose la mise en œuvre et l'entretien d'une coordination, critère auquel seul un *publisher* peut se soumettre). C'est pourquoi l'indicateur de pertinence recherché correspond aux « complémentaires » des diagonales entrevues jusqu'ici (le Tableau 10, ci-dessous, présente un indice de pertinence dans le sens *faisceaux vers versions* ; le Tableau 11, quant à lui, présente cet indice dans le sens inverse, *versions vers faisceaux*) :

	Pointeurs de la ressource 1	Pointeurs de la ressource 2
⟨Faisceau 1, <i>a</i> , <i>t</i> ⟩	X	0
⟨Faisceau 2, <i>a</i> , <i>t</i> ⟩	0	X

Tableau 10
Comparaison entre faisceaux de versions

	Composantes de la ressource 1	Composantes de la ressource 2
⟨Faisceau 1, <i>a</i> , <i>t</i> ⟩	X	0
⟨Faisceau 2, <i>a</i> , <i>t</i> ⟩	0	X

Tableau 11
Tableau inverse du Tableau 10 (à partir de faisceaux de pointeurs)

Nous étions partis d'une pure combinatoire, mais cette condition sur les faisceaux, qui est une condition plus générale que celle de l'auto-stabilité de la diagonale, fait ressurgir la normativité sémantico-pragmatique. *Pragmatiquement*, tout utilisateur d'une adresse et tout lanceur d'une requête doit pouvoir obtenir une réponse, fût-ce une réponse qui revient à dire : cette adresse ne correspond à rien, donc à aucune ressource. Ou encore : cette ressource n'a pas (n'a plus) de contenu (ce qui renvoie à des codes différents au niveau du protocole Http). *Sémantiquement*, les contenus doivent être appropriés d'une part aux pointeurs – ce qui les définit simplement comme faisant partie d'une même ressource – d'autre part aux labels ou thèmes des requêtes.

Mettons en évidence, pour finir, le rapport qu'entretient cette sémantique-pragmatique multidimensionnelle, assortie de ses deux faisceaux, *rayonnants* et *convergeants*, avec la perspective dynamique de l'ontologie que nous avons proposée et qui fait jouer un processus évolutif de *distinction* puis d'*explicitation* tout en ménageant également une place au processus inverse d'*implication*. Un pointeur distingue une ressource mais il n'explicite pas les éléments spécifiques de la ressource. Par la suite, des tags (notamment) peuvent être employés pour distinguer aussi bien des spécificités de la ressource que des

pointeurs. En d'autres termes, pour produire des formes d'explicitations. Cela reste dynamique et récursif, car si les spécificités des éléments de la ressource inspirent des tags, ces tags doivent à leur tour demeurer cohérents au regard de leur faisceau d'appartenance et permettre de retrouver la ressource, d'où des ajustements continuels. Si la ressource est un trajet, les éléments contextuels liés à une variation provisoire ou locale risquent d'être vite périmés (de même, plus l'on considère l'identité de la ressource à un niveau global, à l'échelle du Web par exemple, plus les nécessités d'ajuster seront fortes). Autre processus à prendre en compte : celui où des pointeurs associés à des contenus spécifiques vont permettre de relier différents faisceaux, construisant des réseaux propres à un usage, voire un cercle d'utilisateurs, de manière encore plus variable, et selon des ajustements continuels⁷⁸.

On n'en revient pas pour autant à la position *a priori* de la sémantique bi-dimensionnelle généralisée. En effet, si les architectes, programmeurs et *publishers* du Web doivent garantir 1) la correspondance entre pointeurs, ressources et composantes de la ressource (du point de vue de l'architecture du Web), 2) la pertinence des *liens* sémantiques (Web Sémantique) et, bien entendu, 3) celle des *liens* interactionnels (Web social), la coordination exigée dans tous les cas se stabilisera à mesure que les faisceaux assureront effectivement les convergences qui leur sont demandées. Les liens appelés à se renforcer sont ceux qui seront utilisés avec succès (l'utilisation et la mesure de cette utilisation étant dès lors indissociables). On peut supposer que de tels liens minimiseront les phénomènes d'équivocité et l'accès à des contenus sans rapports avec les attentes induites chez les utilisateurs. À l'inverse, certains liens surprenants peuvent faire jaillir de nouvelles connexions. Celles qui se renforcent sont alors celles qui sont empruntées – à l'instar des modèles d'apprentissage par renforcement des synapses activées. Nous pouvons découvrir dans les relations ainsi renforcées des liaisons catégorielles fondamentales, comme nous pouvons voir émerger des liens métaphoriques, voire des associations qui ont l'arbitraire, le succès et le caractère éphémère d'une mode. Il reste que tous ces liens qui émergent ont bel et bien une certaine normativité, fût-elle temporaire ou bien mieux enracinée.

On notera d'ailleurs, dans le recours aux matrices multidimensionnelles, un effet d'autoréférence du Web vis-à-vis de sa propre structure. Un réseau sous forme de graphe se définit par la matrice qui indique ceux de ses nœuds qui sont connectés. Un réseau qui utilise différents labels exige de définir différents types de connexions, donc de construire des matrices multidimensionnelles. La structure qui préside à la construction progressive du réseau du Web nous fournit ainsi, en fin de compte, le schéma de son ontologie.

V – CONCLUSION

Cette structure de réseau nous fournit un schéma, mais nous en fournit-elle réellement l'ontologie pleine et entière ? Peut-on encore, par exemple, ramener les différentes variétés ontologiques à quelques espèces de base ? Comme lorsque l'on distingue, en philosophie, des entités (ou substrats), des qualités (ou

⁷⁸ Un article du *Figaro* et son « équivalent » republié sur le portail Yahoo! renvoient potentiellement à des réseaux fortement différenciés.

propriétés) ou des relations. Rapportés au cas précis qui nous intéresse, les référents de base correspondraient aux substrats ; les contenus exigeraient quant à eux de poser en vis-à-vis les qualités et propriétés attachés à ces substrats ; alors que les relations permettraient de thématiser les renvois des premiers aux seconds – voire de décrire les processus d’opérations sur le réseau.

En un sens, le schéma présenté ne nous donne plus que des *relations*. Néanmoins, on pourrait également prétendre qu’il n’y a plus, sur le Web, au sens fort, que des symboles faisant office de *substrats* : des entités physiques ne valant que par leurs relations d’opposition binaire formant des graphes. On pourrait également soutenir qu’il ne s’y trouve que des *processus* (ou opérations). Ou encore, que tout est susceptible d’y jouer le rôle de contenu (de *propriétés* et/ou de *relations*), dans la mesure où les symboles se renvoient les uns aux autres non seulement à la manière de *référents* (lorsque des identifiants sont visés par d’autres identifiants) mais également en qualité de *propriétés* (à l’instar des labels associés aux nœuds du réseau : tags ou, en un sens plus générique, littéraux, parfois associés eux-mêmes à des URI, en particulier dans les scénarios de *social bookmarking*⁷⁹) ou de *relations* (les arcs typés, par exemple, qui relient deux ressources du Web Sémantique. Les ressources elles-mêmes n’existent que sur un mode relatif, eu égard aux opérations de « déréférence » ou aux recherches dont elles font l’objet – autant qu’elles les suscitent.

On pourrait cependant soutenir qu’une différence ontologique fondamentale intervient ici : entre les identifiants ou adresses, d’une part, et les contenus associés aux ressources, d’autre part. Différence qui en recouvre une autre, entre deux types d’opérations : pointer vers une référence (une ressource) et recouvrer ses contenus (déréférencer ses représentations). Le schéma multi-dimensionnel, avec ses faisceaux dans les deux sens, qui doivent tous deux satisfaire à la fois un besoin de *convergence* et de *rayonnement*, montre que les contenus jouent à leur tour le rôle de pointeurs, lors même que la multiplication des pointeurs eux-mêmes à l’échelle du Web assure la constitutions de nouveaux liens susceptibles de faire émerger de nouveaux contenus.

C’est bien parce que ces différents points de vue sont tous compatibles qu’il est nécessaire de les relier en vertu de la dynamique ontologique que nous avons évoquée. Reprenons : on lance une première opération ontologique telle que viser une « localisation » sur le Web et permettre de la retrouver. Ceci implique seulement d’avoir explicité l’opération d’« adressage » ou de référence ; nullement que l’entité ontologique ainsi introduite soit *ipso facto* réduite à un référent entendu comme un substrat qui serait la cible de cette visée. En effet, à ce stade, la distinction entre substrat, propriété ou relation, n’a pas encore été explicitée. On pose ensuite une ressource, en entendant par-là ce qui recouvre l’ensemble des contenus de la référence visée⁸⁰. On a donc rendu explicite la possibilité que la référence, une fois « déréféréncée », puisse multiplier les guises : formats, langues et – inévitablement – contenus. Ce qui implique sans doute d’introduire une nouvelle distinction ayant quelque affinité avec la

⁷⁹ Auquel cas, ces URI sont moins les URI du littéral lui-même, que le contenu changeant du tag générique forgé au gré de leur agrégation.

⁸⁰ Sans qu’interviennent ici les limitations génériques associées aux espèces, aux instances (*tokens*), etc.

différence entre substrats et propriétés, ne serait-ce que pour être en mesure de parler « *des* différents contenus d'*une* même ressource ». Mais la ressource joue aussi le rôle de référent, permettant de coupler de manière dynamique référence et contenus de la référence, ne se réduisant donc pas à un substrat. On utilise en outre des libellés et des labels pour relier entre elles des ressources pertinentes au regard d'un thème ou plus largement d'un usage quelconque, ce qui revient à rendre explicite la catégorie ontologique de relation. Mais ces relations elles-mêmes sont tout aussi bien des *contenus* (les URI des ressources constamment agrégées au moyen d'un tag générique) que des entités *référentes* (si l'on vise le tag générique lui-même au moyen de son URI ou, plus généralement, le lien entre deux ressources en considérant cette relation comme une troisième ressource à part entière, ce qui constitue rien de moins que le principe élémentaire du fonctionnement à base de triplets des formalismes du Web de données).

On voit comment procède la dynamique ontologique : on distingue une catégorie ontologique en effectuant une opération (par exemple, faire référence à un référent) mais *cela ne fixe pas pour autant le statut définitif de l'entité en question*. Elle demeure en capacité d'accueillir des déterminations d'autres types ontologiques, selon les besoins. Chaque nouveau développement ontologique exige de distinguer, et donc d'explicitier un nouveau type ontologique autorisant un nouvel usage, tout en laissant intacte la possibilité que l'entité visée sous ce nouveau type puisse également être utilisée sous le type précédent.

À chaque fois qu'on pose une nouvelle distinction – et, partant, une nouvelle explicitation – on mobilise une nouvelle opération. Cette opération n'est pas elle-même explicitée à ce stade, étant simplement utilisée. Aussi pour atteindre (indirectement !) une ressource en déréférençant son URI faut-il encore distinguer le référent de son contenu : ce qui implique une notion de qualité ou de propriété. L'opération de déréférence suppose dès lors une mise en relation du référent et des contenus mais le type ontologique de la relation n'est pas encore distingué et explicité à ce stade. De même, lorsqu'on explicitera la notion de lien, et donc de relation, on emploiera à cette fin les opérations de construction du lien, soit des processus. Mais on n'aura pas encore explicité ces processus eux-mêmes.

Inversement, une fois qu'on a distingué le contenu du référent, on n'a pas pour autant affaire à deux entités, dont la seconde serait incapable de jouer le rôle de la première parce qu'elle jouerait elle-même un rôle entièrement indépendant de celui-ci. Il en est de même pour les relations (qui retiennent en elles les fonctions de référence et de qualification par un contenu) et les processus opératoires (qui retiennent en eux les fonctions des relations, de qualification, etc.). Aussi la distinction et l'explicitation n'impliquent-elles nulle cloison étanche entre les types ontologiques distingués. Si la catégorie B a été distinguée de la catégorie A, elle peut toutefois hériter des fonctions de la catégorie A, dans la mesure où l'opération qui a permis de la distinguer exigeait de disposer desdites fonctions. Les distinctions se développent donc en s'appuyant les unes sur les autres, en quelque sorte par subdivision de fonctions emboîtées. Les types emboîtés font apparaître des distinctions mais ne sont pas pour autant isolables une fois pour toutes des fonctions des types emboîtant.

Dès lors, il est possible d'envisager le développement des différentes fonctions du Web comme la démonstration d'une dynamique ontologique à l'œuvre requérant des types ontologiques encore non-explicités pour faire surgir de nouvelles distinctions, et qui, lorsqu'on en vient à l'explicitation de ces types, *ré-implique* ces distinctions comme leur origine commune, à savoir le déploiement des types précédents.

RÉFÉRENCES

- Bachimont, B. (1996). *Herméneutique matérielle et Artéfacture : des machines qui pensent aux machines qui donnent à penser ; Critique du formalisme en intelligence artificielle*. Thèse de doctorat en épistémologie. Paris, École Polytechnique. Disponible sur : <http://www.utc.fr/~bachimon/Livresettheses.html>.
- Berners-Lee, T., Fielding, R., Irvine, U.C. & Masinter, L. (1998). *RFC 2396 – Uniform Resource Identifiers (URI): Generic Syntax*. Disponible sur : <http://tools.ietf.org/html/rfc2396>.
- Berners-Lee, T. (1994). *RFC 1630 – Universal Resource Identifiers in WWW: A Unifying Syntax for the Expression of Names and Addresses of Objects on the Network as used in the World-Wide Web*. Available at : <http://tools.ietf.org/html/rfc1630>.
- Berners-Lee, T., Groff, J.F. & Cailliau, R. (1992). *Universal Document Identifiers on the Network*. CERN.
- Berners-Lee, T., Masinter, L. & McCahill, M. (1994). *RFC 1738 – Uniform Resource Locators (URL)*. Available at : <http://tools.ietf.org/html/rfc1738>.
- Berners-Lee, T., Hall, W., Hendler, J.A., O'Hara, K., Shadbolt, N. & Weitzner, D.J. (2006). A framework for Web Science. *Foundations and Trends in Web Science, 1(1)*, Boston-Delft, Now Publishers Inc.
- Boulnois, O. (1999). *Être et représentation : Une généalogie de la métaphysique moderne à l'époque de Duns Scot*. Paris, Presses Universitaires de France.
- Bourdeau, M. (2000). *Locus Logicus*, Paris, L'Harmattan.
- Courtine, J.-F. (1990). *Suarez et le système de la métaphysique*. Paris, Presses Universitaires de France.
- Daniel, R. (1995). *An SGML-based URC Service*. Available at : <http://ftp.ics.uci.edu/pub/ietf/uri/draft-ietf-uri-urc-sgml-00.txt>.
- Daniel, R. & Mealling, M. (1995). *URC Scenarios and Requirements*. Available at : <http://ftp.ics.uci.edu/pub/ietf/uri/draft-ietf-uri-urc-req-01.txt>.
- Decock, L. (2002). *Trading Ontology for Ideology: The Interplay of Logic, Set Theory, and Semantics in Quine's Philosophy*. Dordrecht-Boston, Kluwer Academic Publishers.
- Denouël, J. & Granjon, F.(éds.) (2011). *Communiquer à l'ère numérique. Regards croisés sur la sociologie des usages*. 1e éd., Paris, Presses des Mines.
- Ding, Y. et al. (2008). Mediating and Analyzing Social Data. In R. Meersman & Z. Tari, (éds.) *On the Move to Meaningful Internet Systems: OTM 2008* (pp. 1355-1366. Berlin-Heidelberg, Springer. Lecture Notes in Computer Science. Disponible sur : http://link.springer.com/chapter/10.1007/978-3-540-88873-4_30.
- Fielding, R.T. (2000). *Architectural Styles and the Design of Network-based Software Architectures*. PhD Thesis. University of California, Irvine. Disponible sur : http://www.ics.uci.edu/%7Efielding/pubs/dissertation/fielding_dissertation.pdf.
- Fielding, R.T. (1995). *How Roy would Implement URNs and URCs Today*. Disponible sur : <http://ftp.ics.uci.edu/pub/ietf/uri/draft-ietf-uri-roy-urn-urc-00.txt>.
- Fielding, R.T. & Taylor, R.N. (2002). Principled design of the modern Web architecture. *ACM Transactions on Internet Technology (TOIT)*, 2(2), 115-150.
- Fisette, D. (1994). *Lecture frégéenne de la phénoménologie*. Combas, Éditions de l'Éclat.

- Frege, G. (1994). *Écrits logiques et philosophiques*. Paris, Éditions du Seuil.
- Glock, H.-J. (2003). *Dictionnaire Wittgenstein*. Paris, Éditions Gallimard.
- Good, B. & Kavas, E. *About the Entity Descriptor*. *entitydescriptor.org*. Disponible sur : <http://www.entitydescriptor.org/about.html>.
- Halpin, H. et al. (2010). When owl:sameAs Isn't the Same: An Analysis of Identity in Linked Data. In *ISWC 2010, Part I* (pp. 305-320). Berlin-Heidelberg, Springer, Lecture Notes in Computer Science.
- Hennion, A. (2007). *La passion musicale : Une sociologie de la médiation Édition revue et corrigée*. Paris, Éditions Métailié.
- Hennion, A. & Latour, B. (1993). Objet d'art, objet de science. Note sur les limites de l'anti-fétichisme. *Sociologie de l'art*, (6), 7-24.
- Hoffman, P.E. & Daniel, R. (1995). *Trivial URC Syntax: urc0*. Disponible sur : <http://ftp.ics.uci.edu/pub/ietf/uri/draft-ietf-uri-urc-trivial-00.txt>.
- Imbert, C. (1992). *Phénoménologie et langues formulaires*. Paris, Presses Universitaires de France.
- Imbert, C. (1999). *Pour une histoire de la logique: Un héritage platonicien*. Paris, Presses Universitaires de France.
- Kunze, J. (1995). *RFC 1736 – Functional Recommendations for Internet Resource Locators*. Disponible sur : <http://www.rfc-editor.org/rfc/rfc1736.txt>.
- Livet, P. (2012). Web Ontologies as Renewal of Classical Philosophical Ontology. *Metaphilosophy*, 43(4), 396-404.
- Livet, P. & Nef, F. (2009). *Les êtres sociaux : Processus et virtualité*. Paris, Hermann.
- McCarthy, J. (1980). Circumscription a Form of Nonmonotonic Reasoning. *Artificial Intelligence*, 13, 27-39.
- Mallard, A. (2011). Explorer les usages : un enjeu renouvelé pour l'innovation des TIC. In J. Denouël & F. Granjon (éds.) *Communiquer à l'ère numérique. Regards croisés sur la sociologie des usages*. Paris, Presses des Mines, pp. 253-282.
- Mealling, M. (1994). *Specification of Uniform Resource Characteristics*. Disponible sur : <http://ftp.ics.uci.edu/pub/ietf/uri/draft-ietf-uri-urc-00.txt>.
- Mongin, M. (2006). Qui sont les « nouveaux philosophes » analytiques ? – Quand la philosophie fricote avec le monde de l'ingénierie. *Esprit*, 12, (décembre 2006), 189-197.
- Mongin, M. (2007). Réponse de Martin Mongin. *Esprit*, 5, (mai 2007), 176-177.
- Monnin, A. (2012). *L'ingénierie philosophique comme design ontologique : retour sur l'émergence de la « ressource »*. *Réel-Virtuel*, 3. Disponible sur : <http://reelvirtuel.univ-paris1.fr/index.php?revue-en-ligne/3-monnin/>.
- Monnin, A. (2013a). *Les ressources, des ombres récalcitrantes*. *SociologieS*. Disponible sur : <http://sociologies.revues.org/4334>.
- Monnin, A. et al. (2010). Speech acts meet tagging: NiceTag ontology. In *Proceedings of the 6th International Conference on Semantic Systems. I-SEMANTICS '10*. New York, NY, ACM. Disponible sur : <http://doi.acm.org/10.1145/1839707.1839746>.
- Monnin, A. (2013b). *Vers une Philosophie du Web Le Web comme devenir-artefact de la philosophie (entre URIs, Tags, Ontologie(s) et Ressources)*. Thèse de doctorat, . Paris, Université Paris 1 Panthéon-Sorbonne.
- Monnin, A. (2014). The Web as Ontology: Web Architecture Between REST, Resources, and Rules. In H. Halpin & A. Monnin (éds.) *Philosophical engineering: toward a philosophy of the web*. Metaphilosophy. Oxford, Wiley-Blackwell.
- Nef, F. (1998). *L'objet quelconque : Recherches sur l'ontologie de l'objet*. Paris, Librairie philosophique J. Vrin.
- Nef, F. (2006). *Les propriétés des choses : Expérience et logique*. Paris, Librairie Philosophique Vrin.
- Nef, F. (2007). Qui sont les “nouveaux philosophes” analytiques ? II. *Esprit*, 5, (mai 2007), 174-176.

- Passant, A. *et al.* (2009). A URI is Worth a Thousand Tags: From Tagging to Linked Data with MOAT. *International Journal on Semantic Web & Information Systems*, 5(3), 71-94.
- Quine, W.V.O. (1951). Ontology and Ideology. *Philosophical Studies*, 2(1), 11-15.
- Quine, W.V.O. (1983). Ontology and Ideology Revisited. *The Journal of Philosophy*, 80(9), 499-502.
- Rees, J.A. (2011). *Providing and discovering definitions of URIs. Editor's Draft 25 June 2011*. Disponible sur : <http://www.w3.org/2001/tag/awwww/issue57/20110625/>.
- Scerri, S. *et al.* (2008). SCOT Ontology Specification (H. L. Kim & J. G. Breslin, édés.) Disponible sur : <http://scot-project.org/scot/>.
- Smith, B.C. (1998). On the Origin of Objects. Reprint. Cambridge, Mass., MIT Press.
- Soergel, D. (1997). Coverage of ASIS 1997 Annual Meeting. *Uniform Resource Identifiers, Metadata and What They Mean for Access to Networked Digital Resources*. *asis.org*. Disponible sur : <http://www.asis.org/Bulletin/Dec-97/am97extra/metadata.htm>.
- Sollins, K. & Masinter, L. (1994). *RFC 1737 – Functional Requirements for Uniform Resource Names*. Disponible sur : <http://tools.ietf.org/html/rfc1737>.
- URI working group (1994). *URN to URC resolution scenario*. Disponible sur : <http://ftp.ics.uci.edu/pub/ietf/uri/draft-ietf-uri-urn2urc-00.txt>.
- Williams, D.C. (1953). On the Elements of Being. *Review of Metaphysics*, 17. Disponible sur : <http://www.hist-analytic.org/WILLIAMS4.htm>.