

10



STRUMENTI

*Atti dei seminari di studio
Pisa, Scuola Normale Superiore
14 febbraio-12 dicembre 2006*

Seminari Signum 2006

a cura
di Simonetta Bassi



EDIZIONI
DELLA
NORMALE

Indice

Presentazione	7
Spoken dialog systems: from theory to technology GIUSEPPE RICCARDI, PAOLO BAGGIA	9
A Semiotic Metamodel for Bridging Lexical and Formal Semantics ALDO GANGEMI	21
Systems biology challenges computer science CORRADO PRIAMI, PAOLA QUAGLIA	49
Un approccio geometrico all'analisi dei testi letterari ROBERTO BASILI, PAOLO MAROCCO, ROBERTO GIGLIUCCI	61
Ubiquitous computing Challenges for the software engineer CARLO GHEZZI	81
Progettare interfacce web SOFIA POSTAI	93
Biblioteca digitale: esperienze e prospettive ANNA MARIA TAMMARO	101

Presentazione

Il secondo volume dei *Seminari* raccoglie alcuni risultati del lavoro di scambio e di confronto che si è svolto presso *Signum*, il laboratorio di ricerche informatiche per le discipline umanistiche della Scuola Normale Superiore di Pisa diretto da Michele Ciliberto. Sviluppando e allargando il programma già abbozzato nel primo volume questi contributi si propongono di affrontare a più livelli e in diversi settori della conoscenza il problema della rappresentazione e della trasmissione delle informazioni. A partire dai recenti progressi compiuti nell'ambito del riconoscimento automatico dei discorsi, i contributi si snodano presentando da una parte gli ultimi applicativi e le tecnologie di SDS (Spoken Dialog Systems), dall'altra metamodelli per la rappresentazione di informazioni semantiche e lessicali, sviluppando in questo caso una serie di riflessioni sull'utilizzo delle ontologie. La dinamicità e la complessità del linguaggio umano si affiancano alla simultaneità e alla attività cooperativa dei sistemi biologici, che a loro volta vengono analizzati attraverso specifici modelli computazionali per rappresentare le interconnessioni logiche dei processi delle strutture. Lo spazio in cui si distribuiscono le informazioni viene indagato anche attraverso modelli geometrici – dunque spaziali e non logici – e la ricchezza di informazioni che viene ricavata dall'intreccio dinamico di domini e di dati esige una potenza di calcolo che sempre maggiormente si trasferisce da una singola macchina ad un sistema integrato di calcolatori. Una nuvola di notizie, di dati, di oggetti si sta organizzando in un universo articolato, che per certi versi può essere assimilabile al mondo più noto e familiare della biblioteca, che da una parte raccoglie e concentra informazioni, ma dall'altra le ibrida e le moltiplica indefinitamente attraverso l'utilizzo e le domande del lettore.

Alla fine di questo lavoro un ringraziamento va a tutto il personale di *Signum* e del Centro Edizioni della Scuola Normale Superiore e in modo particolare a Francesca Di Dio che con la consueta competenza e dedizione ha curato la redazione del volumetto.

Spoken dialog systems: from theory to technology

ABSTRACT: Interacting with machines that listen, understand and react to human stimuli has been for many years the holy grail of computer scientists. In the last three decades scientists have made great contributions to the training, design and testing of conversational systems. In this paper we present the fundamentals of Spoken Dialog Systems (SDS) from Automatic Speech Recognition, to Spoken Language Understanding and to Text-to-Speech Synthesis. We give also an overview of the current technology and standards supporting speech technologies and applications.

INDEX TERMS: Automatic Speech Recognition, Spoken Language Understanding, Spoken Dialog Systems, W3C Standards.

1. *Introduction*

Interacting with machines that listen, understand and react to human stimuli has been for many years the holy grail of scientists across disciplines. First attempts to build such machines dates back to Kratzenstein's and Von Kempelen's (XVIII Century) mechanical speech synthesizer. However, it was not until first programmable computers were invented that computer scientists experimented with models of language understanding and human-interaction.

One of the debate in the scientific community has been primarily about (not) making anthropomorphic machines. Should we program the peripheral auditory processing of humans into computers? Should we parse language according to linguistic theories? Should computers

Part of this work has been published at the Toni Mian Workshop, Padua, 2007.

Giuseppe Riccardi is funded by Marie Curie Excellence Grant (ADAMACH 022593) and Marie Curie International Reintegration Grant (CASA 36556).

be emotional? The answer to these questions has been different at times depending on factors such as the knowledge about human speech/language processing as well as the availability of large databases and fast machines. Last but not least, in the last decade, some of the scientific contributions have contributed to generate technology used by millions of users in commercial applications.

In this paper we review the fundamental components of spoken dialog systems (SDS) and their theoretical motivation as well as point out current research directions. We decompose the *speech and language chain* into its building blocks. The Automatic Speech Recognizer (ASR) decodes the speech signal into word string hypotheses. The Spoken Language Understanding (SLU) maps word hypotheses into semantic interpretations. The Dialog Manager (DM) takes as input the semantic interpretations and instantiates a machine action. The Natural Language Generator (NLG) instantiates a linguistic realization from the machine action. In the last step of the speech chain, the Text-to-Speech (TTS) synthesizer verbalizes the linguistic realization provided by the NLG. In Fig. 1 show the information process pipeline model of human-machine communication.

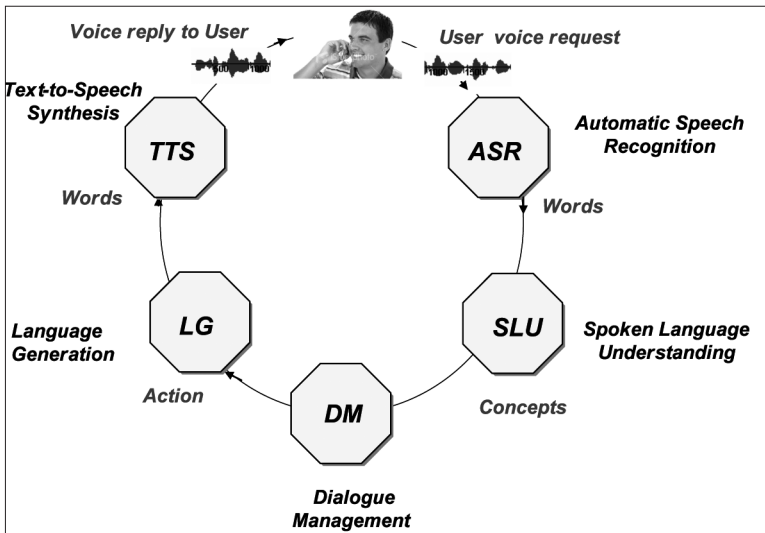


Fig. 1. The building blocks of a Spoken Dialog System.

In the next sections we review each component and the modeling framework. We review the international standards for spoken dialog systems and a spoken dialog prototype based on state-of-the-art spoken dialog technology.

2. Automatic Speech Recognition

In ASR the goal is to decode the word string that is *encoded* in the spoken utterance. The modern approach to ASR is to frame this problem in the noisy channel paradigm. The spoken utterance can be viewed as the noisy version of the word string and thus the decoding problem can be posed as:

$$\hat{W} = \arg \max_W P(W | X) = \arg \max_W \frac{P(X | W)P(W)}{P(X)} \quad (1)$$

where $\hat{W}$ is the decoded word sequence, X is the acoustic observation signal, W is the generic word sequence drawn from a vocabulary V . The term $P(A | W)$ is computed by estimating the *acoustic* model, while $P(W)$ is the stochastic language model (SLM). The denominator $P(A)$ is not relevant for the optimization problem in Eq. 1. In modern ASR the speech signal A is compressed into acoustic features that are required to be minimal (w.r.t. number), robust to noise and discriminative (e.g. w.r.t. to speech sounds). An effective approach to feature extraction is computed through Mel Frequency Cepstrum Coefficient (MFCC) analysis. The core of this procedure is the Mel filter bank which is motivated by human speech perception experiments [1]. Alternative approaches to feature extraction for ASR both in the frequency and time domain can be found in [2]. The probability $P(A | W)$ is estimated by the acoustic model which needs to account for how acoustic features are generated to produce sequence of speech units (e.g. vowels) and ultimately words. The most successful mathematical model to represent the generation of speech units within a stochastic framework is the hidden Markov model (HMM). HMMs define a double stochastic process with a 1) state transition probability matrix $\mathbf{A}=\{a_{ij}\}=\{P(s_j | s_i)\}$ over state pair space and an 2) output symbol probability distribution $\mathbf{B}=\{b_k\}=\{P(z | s_k)\}$ for each state s_k and symbol z in the dictionary of speech units [3]. HMMs allow us to de-

fine the search space of symbol sequences by designing the state topology. A review of HMMs and alternative models can be found in [4].

In Eq. 1 the term $P(W)$ is computed by estimating the probability of a generic word sequence $W=w_1, w_2, ..w_M$ where $w_i \in V$ and M is the number of words in the spoken utterance. In general, we do not know the *grammar* used by the human interacting with computers. There have been many attempts at formalizing grammars describing the space of allowed word sequences but with limited success in ASR, mostly due to coverage and brittleness issues. The most successful approach to modelling word sequences is purely statistical. The probability $P(W)$ is decomposed into n -gram probabilities:

$$P(W) = \prod_i P(w_i | w_1 \cdots w_{i-1}) \approx \prod_i P(w_i | h(w_1 \cdots w_{i-1})) \quad (2)$$

where the function $h(w_1 \cdots w_{i-1})$ maps the substring $w_1, w_2, ..w_{i-1}$ into an equivalent class. The most used equivalence class function is:

$$h(w_1 \cdots w_{i-1}) \rightarrow (w_{i-n+1} \cdots w_{i-1}) \text{ and } P(W) = \prod_i P(w_i | w_{i-n+1} \cdots w_{i-1}) \quad (3)$$

The n -gram probability $P(w_i | w_{i-n+1} \cdots w_{i-1})$ requires counts of n -tuples $C(w_{i-n+1} \cdots w_{i-1} w_i)$ estimated from large corpus of transcription of spoken utterances. Since most of n -tuples counts are zero, smoothing techniques are used to estimates a non-zero probability for all word n -tuples [5].

3. Spoken Language Understanding

ASR engine must have access to prior knowledge about the language to be recognized, in order to have good performance. If a dictation system or advanced SDS heavily relies on probabilistic knowledge to constrain the ASR decoding, such as the n -gram probabilities, in SDS it might be effective, if possible, to pre-define all the possible linguistic surface realization and interpretations that the user might say.

```

<grammar version="1.0" xml:lang="en-US"
  xmlns="http://www.w3.org/2001/06/grammar"
  tag-format="semantics/1.0" root="fromto">

  <rule id="fromto" scope="public">
    from <ruleref uri="#city"/>
      <tag>out.fromcity=rules.latest();</tag>
    to <ruleref uri="#city"/>
      <tag>out.tocity= rules.latest();</tag>
  </rule>
  <rule id="city">
    <one-of>
      <item>Brussels<tag>out="BRU"</tag></item>
      <item>Paris<tag>out="CDG"</tag></item>
      <item>Rome<tag>out="FCO"</tag></item>
    </one-of>
  </rule>
</grammar>

```

Fig. 2. Simple SRGS grammar with SISR script

The most common kind of speech grammars are Context Free Grammars (CFG, type 2 of the Chomsky hierarchy [6]) which can be integrated in the speech decoding process in a tractable, efficient and well performing algorithm. Relevant work has been done by World Wide Web Consortium (W3C) in the definition of a standard syntax for grammars: SRGS (Speech Recognition Grammar Specification Version 1.0) [7], which is largely adopted by the speech engine products.

A second improvement from W3C has been the definition of a standard for *semantic interpretation*, which is the component of a speech grammar devoted to the generation of a semantic result. In many architectures the ASR result is a string of recognized words or multiple results (e.g. N-bests or hypothesis lattices). SISR, which stands for ‘Semantic Interpretation for Speech Recognition Specification Version 1.0’ [8], offers two alternative ways to encode semantic interpretation in SRGS grammars, the first is called *literal semantics*, which allows to transliterate a recognized word for instance to return a single value for alternative pronunciation (‘coke’ for ‘coca cola’).

In alternative SISR 1.0 allows *script semantics*, which is based on ECMAScript/JavaScript (ECMA-327, the computationally constrained version of ECMA-262). With script semantics the ASR engine can return complex semantic results, apply validations to increase or

reduce the reliability of the recognized result, or encode the result to be suitable for the speech application. The grammar in Figure 1 includes SISR script semantics inside `<tag>` elements, so that for ‘from Rome to Paris’ returns the following ECMAScript object `{fromcity: "FCO", tocity: "CDG"}`. While CFG formalism is powerful in SLU we aim at understanding conversational speech, where we need to apply robust parsing techniques for mapping words to concepts [9,10].

4. *Dialog Models*

The dialog manager implements the control algorithm that takes the semantic interpretation hypotheses from machine’s SLU and computes which action the machine should take. For example if the SLU has hypothesized a user request for *flight status*, the DM might generate a machine request for *which flight*. The DM controls the type of actions (linguistic vs non-linguistic), timing of actions (when to perform them) as well as their sequences (strategies). In most DMs the underlying assumptions, is that the machine is expected to help the user perform a pre-defined task (e.g. travel information and planning). Task complexity has a direct impact on the complexity of the DM control algorithm. The simplest approach to DM is to encode the control via a set of rules or finite state automata [11]. While the latter approach is simple as well as limited, more recently there have been promising approaches to stochastic modelling of DM control [12].

5. *Natural Language Generation*

Natural language generation is the process of generating natural language output from DM actions. Other source of knowledge can affect such process (e.g. flight databases). NLG is the inverse of SLU with the most notable difference that the input is not-linguistic. The NLG algorithm is expected to select what-to-say and how-to-say. Since there are many linguistic choices for a given input, NLG algorithm decides what is the best amongst many alternatives. NLG in Spoken Dialog Systems is currently limited to selection from predefined textual realization or template filling at best. However in text processing there have been early attempts to design complete NLG systems [11].

6. Speech Synthesis

The main technique in modern Text-to-Speech (TTS) is the *Unit Selection* concatenation technique [13,14]. The basic idea is to extend the concatenation units from diphones to wider units from a large speech database. The result is highly intelligible and natural compared with early attempts in the '70s, being them parametric approaches (formant or articulatory) or diphone concatenation [15]. Such advances were possible also due to increased availability of CPU power and memory.

The TTS architecture [16] is composed of two main modules: the *Text Analyzer*, whose goal is to convert the input text into a phonetic and prosodic representation; and the *Speech Synthesizer*, whose goal is to find the optimal sequence of units, then to concatenate them to obtain the synthesizer output. Text analysis relies on language knowledge, lexicons and rules; Speech Synthesis relies on speech processing techniques to find the optimal sequence of segments and smoothly concatenate them. Acoustic databases are designed to provide high phonetic/prosodic coverage of the intended language (or application domain).

One of the research topic is to generate output with emotional control [17,18]. Loquendo TTS voices are produced with a list of 'expressive cues', which enable TTS users to enliven their voice prompts. This is a concrete attempt in the direction of expressive synthetic speech. Another important area is the multilinguality, where the TTS should be able to read multiple languages, even mixed in the same sentence. A common approach is to record voices from bilingual or polyglot speakers and allow the TTS engine to select the proper language. A more general approach is based on Phoneme Mapping [19], where each word/sentence is transcribed according to its language, and then the phonemes are mapped into similar ones according to articulatory features in the target language. Phonemes that do not belong to the native phonological system of the voice must be replaced by the most similar sounds available. Moreover there is increased interest in exploiting advanced techniques developed for ASR into TTS, see [20].

The input to a TTS engine can either be a text or a markup language, such as SSML (Speech Synthesis Markup Language) [21], which is a W3C Recommendation since 2004 and describes elements to control and improve TTS rendering.

7. Standards for Speech Applications

Today commercial SDS are based on VoiceXML 2.0/2.1 [22, 23], released by the W3C Voice Browser (VBWG) in 2004. VoiceXML markup describes in a declarative way the dialog interaction, where a VoiceXML platform downloads the documents via Hypertext Transfer Protocol (HTTP) from an application server. This allows to decouple the dialog logic creation, that can be tightly coupled with back-end DBs, from the VoiceXML platform and to re-use the web architecture to produce and deploy speech applications. For these reasons the VoiceXML is today pervasive into commercial speech applications ranging from simple touch-tone driven Interactive Voice Response (IVR) to self-service information systems, to customer care.

A second area of interest for the application of speech technologies is the multimodal environment, which is pushed by the increasing presence of mobile devices, kiosks, and laptop. In a multimodal applications the DM must be able to integrate and coordinate different input/output modalities, such as speech, gesture, audio, video, etc. To promote standards for multimodal W3C launched in 2002 the Multimodal Interaction working group, in [24] the proposals of a Multimodal architecture and other related standards (EMMA [25] for rich annotation of input and InkML [26]). The first generation of standards (VoiceXML 2.0/2.1) is now a pervasive reality, but a new generation is under development, such as SCXML [27] which allows a DM to control different modalities and is going to become the build block of future speech and multimodal applications.

8. Spoken Dialog Systems

In this section we describe the architecture of an SDS based on speech grammars and VXML standard platform used in our lab at University of Trento. In Figure 4 the architecture of the Spoken Dialogue System is supported by the VXML platform. The DM engages the user in a conversation to drive the user through the completion of its request. The DM processes the semantic representation of the utterance returned by the VXML platform and query the following resources to generate a machine action. *Dialog History*: the database where the information about past spoken dialog events (e.g. users' transcriptions, semantic interpretations, machine actions etc.) is stored. The *Domain*

Concept Ontology (DCO): A hierarchical description of the set of domain concepts. In Fig. 3 we give an excerpt of such ontology in the University help-desk domain:

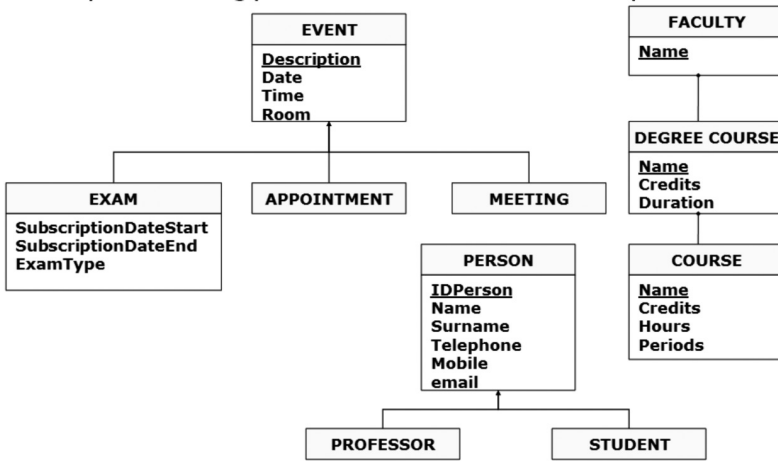


Fig. 3. Domain Concept Ontology Fragment

The *Action Ontology (AO)*: A set of predicates that operate on one or more concepts. The *Action Constraints (AC)*: a set of pre-conditions that are necessary to activate any of the predicate in the Action Ontology. *Strategy Planner (SP)*: this module generates a sequence of actions to be taken based on the input from AO, AC and Dialog History. The DM matches the current semantic interpretation against the DCO and retrieves the possible machine actions from the AO. This set is validated by matching the action constraints present a given dialog turn. The set of available machine actions are then used to build a strategy plan via the SP module. The key resource of the DM control is the dialog history which is instrumental to take decisions in the context of current or past dialog interactions. A more detailed discussion of our data-centric approach to Dialog Management can be found in [28]. When the DM has selected the *best* machine action, the NLG module is invoked and then its output verbalized by the platform TTS. In the current implementation the NLG is driven by templates or text tables indexed by machine actions. The DM iterates the previously described steps until the call reaches a final state (e.g., in the case of a routing application, the call is transferred to a live agent, an IVR or the caller hangs up or the caller reaches the goal).

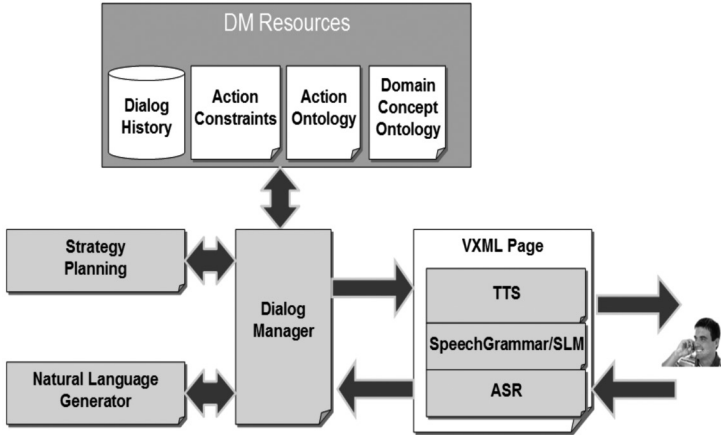


Fig. 4. Spoken Dialog System Architecture

In our architecture we have decoupled the data representation (e.g. Dialog History, DCO, etc.) from the control (DM) so that it can support multiple DM models (e.g. rules, finite state automata, stochastic models) and adapt to new in-domain contexts or multiple domains (e.g. different DCO).

9. Conclusion

We have described the fundamental information processes of spoken dialog systems. Current mathematical models rely on statistical modeling of data from automatic speech recognition to text-to-speech synthesis. Formal models of concept and domain knowledge is required to complete the design of spoken dialog systems. In the last decade some of the research achievements have spurred successful technological and commercial applications. The gap between research and industry has been filled, partly, by the standardization efforts within the World Wide Web Consortium.

10. Acknowledgements

We would like to thank Pierluigi Roberti and Sebastian Varges for their contribution to the development of the spoken dialog prototype and LOQUENDO for providing the VoiceXML VoxNauta™ platform to carry out the experiments.

11. Bibliography

- [1] S.S. STEVENS, J. VOLKMAN, E.B. NEWMAN, *A Scale for the Measurement of the Psychological Magnitude of Pitch*, in «Journal of the Acoustical Society of America», VIII, 1937, pp. 185-190.
- [2] X. HUANG, A. ACERO, H.W. HON, *Spoken Language Processing*, Prentice Hall 2001.
- [3] L.R. RABINER, *Tutorial on Hidden Markov Models and Selected Applications in Speech Recognition*, in «Proc. IEEE», LXXVII, 2, 1989, pp. 257-286.
- [4] J.A. BILMES, *What HMMs Can Do*, in «IEICE Trans. Inf and Syst.», ELXXXIX-D(3), 2006.
- [5] J. GOODMAN, *A Bit of Progress in Language Modeling*, in «Computer Speech and Language», XV, 4, 2001, pp. 403-434.
- [6] N. CHOMSKY, *On Certain Formal Properties of Grammars*, «Information and Control», II, 1959, pp. 137-169.
- [7] A. HUNT, S. MCGLASHAN, *Speech Recognition Grammar Specification Version 1.0*. W3C Recommendation, Mar. 2004.
- [8] L. VAN TICHELEN, D. BURKE, *Semantic Interpretation for Speech Recognition (SISR) Version 1.0*. W3C Recommendation, Apr. 2007.
- [9] A.L. GORIN, G. RICCARDI, J.H. WRIGHT, *How May I Help You?*, in «Speech Communication», XXIII, 1997, pp. 113-127.
- [10] R. DE MORI *et al.*, *Spoken Language Understanding: A Survey*, to appear on «IEEE Signal Processing Magazine», 2008.
- [11] D. JURAFSKY, J. MARTIN, *Speech and Language Processing*, Prentice Hall 2000.
- [12] E. LEVIN, R. PIERACCINI, W. ECKERT, *A Stochastic Model of Human-Machine Interaction for Learning Dialog Strategies*, in «IEEE Trans. SAP», VIII, 2000.
- [13] M. BALESTRI *et al.*, *Choose the Best to Modify the Least: a New Generation Concatenative Synthesis System*, in «Proc. of EUROSPEECH-99», V, 1999, pp. 2291-2294.
- [14] A. HUNT, A.W. BLACK, *Unit Selection in a Concatenative Speech Synthesis Using a Large Speech Database*, in «Proc. of ICASSP», 1996, pp. 373-376.
- [15] D. KLATT, *Review of Text-to-Speech Conversion for English*, in «Journal of the Acoustical Society of America», LXXXII, 1987, pp. 737-793.
- [16] *Multilingual Text-to-Speech Synthesis: The Bell Labs Approach*, ed. by R. Sproat, Kluwer Academic Publishers 1997.
- [17] E. EIDE *et al.*, *Towards Synthesizing Expressive Speech*, in «Text to Speech Synthesis: New Paradigms and Advances», Prentice Hall 2004, pp. 219-248.
- [18] E. ZOVATO *et al.*, *Towards Emotional Speech Synthesis: a Rule Based Approach*, in «Proc. 5th ISCA Speech Synthesis Work», 2004.

- [19] L. BADINO, C. BAROLO, S. QUAZZA, *Language Independent Phoneme Mapping for Foreign TTS*, «Proc. ISCA Speech Synthesis Work», 2004, pp. 217-218.
- [20] M. OSTENDORF, I. BULYKO, *The Use of Speech Recognition Technology in Speech Synthesis*, in «Text-to-Speech Synthesis: New Paradigms and Advances», Prentice Hall 2004, pp. 109-133.
- [21] D. BURNETT *et al.*, *Speech Synthesis Markup Language (SSML) Version 1.0*. W3C Recommendation, Sept. 2004.
- [22] S. MCGLASHAN *et al.*, *Voice Extensible Markup Language (VoiceXML) Version 2.0*. W3C Recommendation, Mar. 2004.
- [23] M. OSHRY *et al.*, *Voice Extensible Markup Language (VoiceXML) 2.1*. W3C Recommendation, Jun. 2007.
- [24] W3C Multimodal Interaction : <<http://www.w3.org/2002/mmi>>.
- [25] M. JOHNSTON *et al.*, *EMMA: Extensible MultiModal Annotation Markup Language*. W3C Candidate Recommendation, Dec. 2007.
- [26] Y-M. CHEE *et al.*, *Ink Markup Language (InkML)*. W3C Working Draft, Oct. 2006.
- [27] J. BARNETT *et al.*, *State Chart XML (SCXML): State Machine Notation for Control Abstraction*. W3C Working Draft, Feb. 2007.
- [28] S. VARGES, G. RICCARDI, *A Data-Centric Architecture for Data-Driven Spoken Dialog Systems*. Proc. IEEE ASRU Workshop, Kyoto 2007.

GIUSEPPE RICCARDI, PAOLO BAGGIA

A Semiotic Metamodel for Bridging Lexical and Formal Semantics

1. Introduction

The amount of lexical resources that are developed either as long-term repositories, or as short-term products of NLP techniques, is growing significantly, posing the problem of understanding their commonalities and their potential for reusability and interoperability. A major concern for reusability and interoperability is the ability to control, both intellectually and computationally, the semantics of the data that are contained in a lexical resource. Lexical data semantics is usually left implicit, either because there is not a shared agreement on how to represent that semantics, or because developers of lexical resources do not customarily employ formal methods.

This article presents *Semion*, a metamodel¹ for bridging lexical and formal semantics, thus allowing the representation of relations (e.g. annotation, reuse, mapping, transformation, etc.), which can hold between the elements of a lexicon (either established, like WordNet, or created for local purposes), and the elements of a typed logical theory, e.g. an ontology.

Semion is based on a standard version of semiotics, and is applied to define metamodels for lexical resources, such as a FrameNet (BAKER, FILLMORE, LOWE 1998) metamodel (*OntoFrameNet*)², and to introduce a formal method for lexical information integration.

Elsewhere (GANGEMI 2010), Semion has been founded on an ontology of *schemata* (LAKOFF 1987, JOHNSON 1987, LANGACKER 1990, TALMY 2003), where schemata are considered as *invariances that emerge from the co-evolution of organisms and environment, and that*

¹ A metamodel, broadly speaking, is a model that describes constructs and rules needed to create specific models.

² <<http://www.ontologydesignpatterns.org/ont/ofn/ofntb.owl>>.

are exemplified by neurobiological, cognitive, linguistic, and social constructs. That foundation allows to liaise lexical and formal semantics with cognitive linguistics and construction grammar. However, it is not described here for space reasons.

While specific relations between individual ontologies and lexica are addressed in literature quite often, e.g. (GANGEMI, GUARINO, MASOLO, OLTRAMARI 2003; BUITELAAR, CHOI, GANGEMI, HUANG AND OLTRAMARI 2007; DE LUCA, EUL, NÜRNBERGER, 2007; SCHEFCZYK, BAKER, NARAYANAN 2008; HUANG, CALZOLARI, GANGEMI, LENCI, OLTRAMARI PREVOT 2008), it is far less usual to propose a metamodel to formally represent the liaison between lexical and formal semantics. Metamodels have been proposed to abridge different lexical resources, starting with OLIF (McCORMICK, LIESKE CULUM, 2004), and recently with reference to lexical semantics, as in LMF (FRANCOPOULO, GEORGE, CALZOLARI, MONACHINI, BEL, PET, SORIA 2006), where an attempt has been made to informally align some lexica under a same metamodel.³ Some steps towards linking lexica and ontologies have been also made in order to manage lexicon reuse in ontologies (PAZIENZA, STELLATO 2006), multi-linguality in ontologies (PETERS, MONTIEL-PONSODA, AGUADO DE CEA, GÓMEZ-PÉREZ 2007), as well as to make cookbook-like transforms between syntactic patterns and formal constructs (CIMIANO, HAASE, HEROLD, MANTEL, BUITELAAR 2007).

Semion abstracts from individual interfaces, lexical standards or specific transformation methods, by providing a semiotic ontology that covers both the intuitive semantics of different lexica, and formal semantics. Formal semantics is a theory of meaning based on a formal language and on its interpretation given by assigning a denotation (e.g. a set extension) to each non-logical construct in that language. Semion is an appropriate metamodel, because its constructs are intuitive enough in order to align the underlying assumptions of different lexical resources, but they can also be made formal by applying a *transformation pattern*⁴ to a formal semantic construct.

A notable advantage of this intermediate layer is that any interface or translation method can refer to a unique intermediate level, without worrying about the intended conceptualization of the data

³ <http://lirics.loria.fr/doc_pub/ExtendedExamplesOfLexiconsUsingLMF29August05.pdf>.

models used in the original lexical resources, or about how to access them. Moreover, developers of lexical resources can continue developing their resources without changing their workflow in order to stay tuned with e.g. semantic web applications, which require different data models.

This article demonstrates the application of Semion by formalizing the metamodels needed to integrate e.g. *verb classes* from VerbNet (KIPPER, DANG, PALMER 2000), *frames* from FrameNet (Fillmore, Kay, O'Connor 1988), *synsets* from WordNet (FELLBAUM 1998), etc., and showing sample customizable translations to formal entities like *frames* (MINSKY 1975; BRACHMAN 1977), *intensional relations* (GUARINO 1998), *content ontology design patterns* (GANGEMI 2005), etc.

Semion is compliant with the intuitive social meaning underlying lexical notions, while retaining the possibility of mapping them to formal notions, which grants desirable computational properties to lexical resources, specially in order to foster the achievement of a generalized *information integration*, which is a main challenge e.g. for 'Web Science' (BERNERS-LEE, HALL, HENDLER, SHADBOLT, WEITZNER 2006) and the *Linking Open Data* W3C project.⁴

Semion is built on top of the c.DnS ontology, which is described in section 1.2. In section 1.3, Semion is introduced; in section 1.4, Semion-based metamodeling is exemplified with reference to WordNet, FrameNet, and VerbNet; in section 1.5, mapping and transformation patterns are exemplified.

2. The c.DnS ontology

c.DnS is represented in both first-order logic and OWL-DL (BECHHOFFER, HARMELEN, HENDLER, HORROCKS, MCGUINNES, PATEL-SCHNEIDER, STEIN 2004), the Web Ontology Language in its description logic variety.⁵ Here only the first-order version of c.DnS is summarized as a foundation for Semion.

⁴ <<http://esw.w3.org/topic/SweoIG/TaskForces/CommunityProjects/LinkingOpenData>>.

⁵ The OWL-DL ontologies presented here can be downloaded from: <<http://ontologydesignpatterns.org/ont/cdns/index.html>>.

The core structure for the *c.DnS* ontology is represented as a relation with arity=8 (see [GANGEMI 2008] for an axiomatization):

$$c.DnS(d,s,c,e,a,k,i,t) \rightarrow D(d) \wedge S(s) \wedge C(c) \wedge E(e) \wedge A(a) \wedge K(k) \wedge I(i) \wedge T(t) \quad (1.1)$$

where *D* is the class of *Descriptions*, *S* is the class of *Situations*, *C* is the class of *Concepts*, *E* is the class of *Entities*, *A* is the class of *Social Agents*, *K* is the class of *Collections*, *I* is the class of *Information Objects*, and *T* is the class of *Time intervals*.

c.DnS classes are structured as follows: *E* is the class of everything that is assumed to exist in some domain of interest, for any possible world. *E* is partitioned in the class *SE* of ‘schematic entities’, i.e. entities that are axiomatized in *c.DnS* (*D*, *S*, *C*, *A*, *K*, *I*), and the class *SE* of ‘non-schematic entities’, which are not characterized in *c.DnS* (e.g. *T*). Schematic entities include concepts, roles, relationships, information, organizations, rules, plans, groups, etc. Examples of non-schematic entities include time intervals, events, physical objects, spatial coordinates, and whatever is not considered as a schematic entity by a modeller.

G is another subclass of *E*, and includes either schematic or non-schematic entities. Its definition is: *any entity that is described by a description in c.DnS*.⁶

In intuitive terms, *c.DnS* classes allow to model how a social agent, as a member of a certain community, singles out a situation at a certain time, by using information objects that express a descriptive relation that assigns concepts to entities within that situation. In other words, these classes express the constructivist assumption according to which, in order to contextualize entities and concepts, one needs to take into account the viewpoint, for which the concept is defined or used, the situation that the viewpoint ‘carves out’ from the observable environment, the entities that are in the setting of said situation, the social agents who share the viewpoint, the community, of which these agents are members, the information object by which the viewpoint is expressed, and the time-span, at which the viewpoint is applied to an environment..

⁶ When an entity is described in *c.DnS*, it gets a ‘unity criterion’: a property that makes that entity an individual, distinct from any other one. For example topological self-connexity, perceptual saliency, functional role in a system are typical unity criteria.

The key notion in *c.DnS* is *satisfiability* of a description within a situation. Situations (*circumstantial* contexts) select a set of entities and their relations as being relevant from the viewpoint of a description (*conceptual* context). In mainstream terms, a situation is the context in which a set of entities *count as* the concepts in the "concepts" in the context of a "description". The *countsAs* relation (SEARLE 1995), originally defined as holding between an entity, a concept, and a generic context, is then revised in order to allow for two types of contexts, which are orderly paired to entities and concepts.

For example, the relation:

$$\text{counts As}(\text{John}, \text{Student}, \text{University}) \quad (1.2)$$

saying that John counts as a student in a university context, can be refined in *c.DnS* as:

$$c.DnS([\text{John}, \text{JohnAtUniversity}], [\text{Student}, \text{UniversityRules}]) \quad (1.3)$$

i.e. that $\text{John} \in E$, in the circumstantial context of a university ($\text{JohnAtUniversity} \in S$), is a student ($\text{Student} \in C$) according to the conceptual context of the rules of that university ($\text{UniversityRules} \in D$).

The other classes in *c.DnS* represent two additional context types: informational, and social.

Informational contexts are the ones encompassing the information objects that are used to express descriptions and concepts, for example the sentence:

$$\text{John goes to the university} \quad (1.4)$$

in the context of a family conversation about John respecting course duties gets appropriate circumstantial and conceptual contexts, while a context like this resume of a TV episode, in which John is a policeman, does not:

John goes to the university and while undergoing an MRI, they discover that he has something metallic inside him which is preventing the MRI scan (1.5)

Social contexts like communities, groups, etc., are the ones encompassing agents that conceptualize entities. E.g. the online community of death metal fans could fit the conceptual context of John as a university student, while a local group of knitted lace shawl makers is far less typical.

c.DnS provides four context types (conceptual, circumstantial, informational, and social), which are summarized in the class diagram of Figure 1.1. The classes involved in each context type can be considered as the representation of ‘naturalized’ counterparts of typical abstract notions such as concept, relation, collection, information, etc. This naturalization operates by firstly *reifying* the abstract notion, and then assuming that it is something that has a lifecycle in the social world. For example, the concept ‘table’ is an abstract entity, which is formally represented as a class, but we can conceive of a story of that concept, its adoption, social function and changes in different cultures, etc. The entity that has that story is the *naturalization* of the abstract concept of a table.

Figure 1.1 also shows some relations between classes in the four context types, which are actually some relevant projections of the *c.DnS* relation. Here the ones that are needed for Semion are listed. For a more complete axiomatization, and technical details on how *c.DnS* is applied in domain ontology projects, see (GANGEMI 2008).

The following is the signature of the projections from Figure 1.1:

$$\Pi_{g.cdns} = \{\text{defines}, \text{usesConcept}, \text{satisfies}, \text{classifies}, \text{about}, \text{describes}, \text{conceptualizes}, \text{expresses}, \text{memberOf}, \text{isSettingFor}, \text{covers}, \text{characterizes}, \text{interpretedAs}\} \quad (1.6)$$

The rationale for the introduction of projections is such that each projection implies the full *c.DnS* relation, according to the axiom schema in 1.7.

$$\pi(x_1 \dots x_{n \geq 2} \mid x_i \in \{d, s, c, e, a, k, i, t\}) \rightarrow c.DnS(d, s, c, e, a, k, i, t) \quad (1.7)$$

Descriptions are schematic entities that naturalize (the reified intension of) *n*-ary, even polymorphic, relations; for example, the *give*(*x*,*y*,*z*,*t*) relation (some *x* gives some *y* to some *z* at time *t*) can be reified as *D*(*giving*). The axioms of the original relation, e.g. domain restrictions, are reified accordingly, by using the *defines* or *usesConcept* relations. For example,

$$\text{defines}(\text{giving}, \text{donor}) \quad (1.8)$$

$$\text{usesConcept}(\text{giving}, \text{timespan}) \quad (1.9)$$

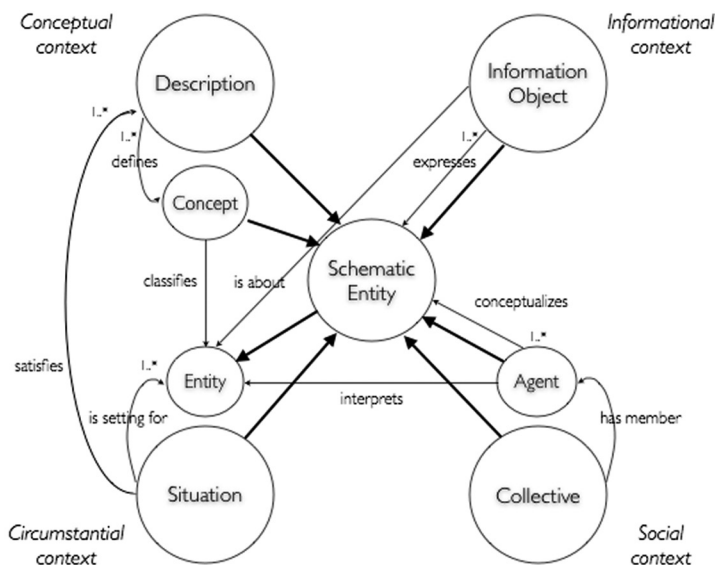


Fig. 1.1 The contextual bindings for the representation in C.DnS (following the OWL version of the ontology). Ovals denotes classes, bold arrows denote subclass relations, regular arrows denote relation holding between members of the linked classes. The cardinality of a relation and its inverse is by default $0..^*$, except when indicated explicitly.

In *c.DnS*, descriptions must be *conceptualized* by social agents, and *expressed* by some information object, i.e. they should be communicable (MASOLO, VIEU, BOTTAZZI, CATENACCI, FERRARIO, GANGEMI, GUARINO 2004). Examples of descriptions include theories, regulations, plans, diagnoses, projects, designs, techniques, social practices, etc. Descriptions can *describe* entities. For example,

$$\text{unifies}(\text{giving}, / \text{donor collection} / \quad (1.10)$$

$$\text{describes}(\text{giving}, / \text{my recent birthday gift} / , t_1) \quad (1.11)$$

Descriptions, as any schematic entity, can be *specialized* (the reification of the formal subsumption relation) and *instantiated* (the reification of the formal inclusion relation) by other descriptions.

Situations are schematic entities that naturalize (reified) instances of *n*-ary relations; for example, the relationship implicit in the sentence:

Mario gave a bicycle to Anna on Sunday (1.12)

can be formalized as:

give(Mario, bicycle, Anna, Sunday) (1.13)

and can be reified as:

S(/Mario gave a bicycle to Anna on Sunday/) (1.14)

Similarly to conceptual axioms for descriptions, the assertional axioms for situations need also to be reified accordingly, typically as elementary situations that are part of the complete situation, e.g., if the assertional relation axiom: *receives*(Amelie, puppet) is reified as the assertional class axiom: *S*(/Anna receives a puppet/), the following holds:

hasPart(/Mario gave a bicycle to Anna on Sunday/,
/Anna receives a bicycle/) (1.15)

In *c.DnS*, situations must *satisfy* a description and are *settingsFor* entities, e.g.:

satisfies(
/Mario gave a bicycle to Anna on Sunday/, *giving*) (1.16)

settingFor(
/Mario gave a bicycle to Anna on Sunday/, Mario) (1.17)

Examples of situations include facts, plan executions, legal cases, diagnostic cases, attempted projects, technical actions. *Situation classes* project n-ary relation extensions into class extensions. For example, the *give*(*x,y,z,t*) relation can be projected as the situation class *Giving* $\subseteq S$, so that the following holds:

Giving(/Mario gave a bicycle to Anna on Sunday/) (1.18)

Concepts are schematic entities that naturalize (the reified intension of) classes; for example, the *Person*(*x*) class can be reified as *C*(*person*). Concepts are *defined* or *used* in descriptions, for example in order to reify the domains of n-ary relations. The axiom:

$$give(x, y, z, t) \rightarrow person(x) \quad (1.19)$$

can be reified as

$$\begin{array}{l} D(giving) \\ C(person) \\ defines(giving, person) \end{array} \quad (1.20)$$

Concepts typically *classify* entities, e.g.

$$classifies(person, Mario) \quad (1.21)$$

and can *cover* or *characterize* collections, e.g.:

$$cover(person, personCollection) \quad (1.22)$$

$$characterizes(person, Italians) \quad (1.23)$$

Collections are schematic entities that naturalize the reified extension of classes; for example, the $\{x_1 \dots x_n\}$ extension of class *Person* can be reified as $K(personCollection)$, so that:

$$\begin{array}{l} \forall(x)(Person(x) \rightarrow \\ (memberOf(x, personCollection, t_1))) \end{array} \quad (1.24)$$

Collections are *coveredBy* or *characterizedBy* concepts, and can have members, e.g.

$$memberOf(Mario, personCollection, t) \quad (1.25)$$

Collections capture the common sense intuition underlying groups, teams, collections, collectives, associations, etc.

Social agents are schematic entities that personify other entities within the social realm: corporations, institutions, organizations, social relata of natural persons. For example, the natural person *Mario* can be personified as *(MarioAsLegalPerson)*. Social agents must be *introducedBy* descriptions, for example by legal constitutive rules (SEARLE 1995); social agents are also able to *conceptualize* descriptions.

Information objects are schematic entities that naturalize units of information: the character *Q*, the German word *Sturm*, the symbol $\otimes$, the text of *Dante's Comedy*, the image of *Francis Bacon's Study from In-*

nocent X , etc. Information objects *express* a schematic entity ($se \in SE$): a description, a concept, a situation, a collection, another information object, or even a social agent. For example, *expresses*(*Sturm*, *Storm*). Information objects can also be *about* other entities, typically situations; for example,

$$\begin{aligned} & \textit{about}(\textit{"Mario gave a bicycle to Anna on Sunday"}, \\ & \quad / \textit{Mario gave a bicycle to Anna on Sunday}/) \end{aligned} \quad (1.26)$$

Based on these projections, the axiom 1.27 formalizes the basic idea of an entity that is given a unity criterion by being *described* by a description 1.27.

$$G(x) \leftrightarrow E(x) \wedge \exists(y, t)(D(y) \wedge \textit{describes}(y, x, t)) \quad (1.27)$$

1.2.1 Metamodeling with $c.DnS$

In this section, the value of $c.DnS$ for metamodeling is justified with reference to semiotics and formal semantics, i.e. where the bridging between lexical and formal semantics is to be expected.

Semiotics The expressive power of $c.DnS$ lies at the level where semiotic activity of cognitive systems occurs: where agents encode expressions that have a meaning in order to denote or construct reference entities in the world. From this perspective (PEIRCE 1958; JAKOBSON 1990; Eco 1975), $c.DnS$ informational context matches the *expression layer*, circumstantial context matches the *reference layer*, social context matches the *interpreter layer*, while all contexts together, and especially the conceptual one, match the *meaning layer*. For example, one can represent statements such as

$$\text{When João says he's rich, he means he has a lot of friends} \quad (1.28)$$

which applies both referential and metalinguistic functions (JAKOBSON 1990) in a same speech act (SEARLE 1969); João is an agent in the social context, and uses contextualized information objects ('rich') that have a contextualized meaning ('having a lot of friends') and contextualized circumstances (João's linguistic act and his situation of having a lot of friends). This is the case for most linguistic acts that are implicit in lexica, thesauri, explanatory texts, web tags, etc.

As a semiotic-oriented ontology, $c.DnS$ can be used to align and integrate models that have heterogeneous semantics and (implicitly) encode different linguistic acts, e.g. different lexical models:

WordNet (FELLBAUM 1998), FrameNet (RUPPENHOFER, ELLSWORTH, PETRUCK, JOHNSON, SCHEFFCZYK 2006), VerbNet (KIPPER, DANG, PALMER 2000); different theories of meaning: frame semantics (FILMORE 1982), semiotic theory (e.g. Eco 1975), formal languages such as OWL (W3C 2004); different texts: explanatory, metalinguistic; tagging (e.g. in Web2.0 applications, [GRUBER 2007]) vs. topic assignment (e.g. in subject hierarchies, [WELTY 1999]), etc.

A collection of examples for this integration task is collected under the LMM (Lexical MetaModel) umbrella⁷ (PICCA, GANGEMI, GLIOZZO 2008), as a formal infrastructure for ‘extreme information integration’ over the Web. In section 1.3, c.DnS is extended with a semiotic vocabulary specifically.

Formal semantics From a formal perspective, the basic intuition of *c.DnS* can be interpreted in the context of a procedure of logical reification:⁸ a *description* can be understood as the reification of 1) $\rho \in T$, with ρ being an intensional relation of any arity, either mono- or polymorphic, and T being an ontology (a typed logical theory), and 2) the axioms $\alpha_{1\dots n}$ that characterize ρ (i.e. the sub-theory $T_\rho \subseteq T$ with $\alpha_{1\dots n} \in T_\rho$).

On the other hand, *situations* result from reifications of (1) each of the individuals $r_{1\dots n} \in R^I$, R^I being the extensional interpretation of ρ , and of (2) the assertions $a_{1\dots n}$ that characterize r_i in accordance with the extensional interpretation of the axioms $\alpha_{1\dots n} \in T_\rho$. A *situation class* is consequently the reification of the set $\{r_1, \dots, r_n\}$ where $r_i \in R^I$.⁹

So equipped, *c.DnS* is able to formally represent any lexical knowledge base. The most complex lexical resource model is FrameNet’s, therefore the ability to represent FrameNet is a good benchmark for *c.DnS* being appropriate as a foundation for a lexical metamodel. FrameNet representation in *c.DnS* is ensured by the assumption that frames, schemata from cognitive linguistics, patterns from knowledge engineering, etc. can all be considered as *n*-ary relations, with typed arguments (either mandatory or optional), qualified cardinalities, etc.

⁷ <<http://ontologydesignpatterns.org/ont/lmm/opensource2lmm.owl>>.

⁸ See also (MASOLO, VIEU, BOTTAZZI, CATENACCI, FERRARIO, GANGEMI, GUARINO 2004) for an alternative, but compatible axiomatization of a part of *c.DnS*.

⁹ Notice that reification is used here in two different senses, as pinpointed in (GALTON 1995): type-reification of classes to individuals, metaclasses to classes, etc., versus token-reification of tuples to individuals, sets of tuples to classes, etc.

For example,

$$\text{Desiring}(x,y,e) \rightarrow \text{Agent}(x) \wedge \text{Agent}(y) \wedge \text{Event}(e) \quad (1.29)$$

An occurrence of a frame is straightforwardly treated as an instance of an n-ary relation, e.g.:

$$\text{Desiring}(\text{Susan}, \text{Marko}, \text{ListeningToHer}) \quad (1.30)$$

The logical representation of frames as n-ary intensional relations is elegant and clear, but hardly manageable by automated reasoners on large knowledge bases. A hard design problem is constituted by the polymorphism of many n-ary relations, which can vary in number of the arguments that can be taken by the relation. For example, the same frame *Desiring* can be assumed with four arguments:

$$\text{Desiring}(x,y,e,t) \rightarrow \text{Agent}(x) \wedge \text{Agent}(y) \wedge \text{Event}(e) \wedge \text{Time}(t) \quad (1.31)$$

This problem was originally evidenced by Davidson with reference to a logic of events (DAVIDSON 1967).

A formal semantics for frames that is also computationally manageable has been provided by description logics (BAADER, CALVANESE, MCGUINNESS, NARDI AND PATEL-SCHNEIDER 2003). Due to their limited expressive power – e.g. they can only represent relations with arity=2 – that is balanced by desirable computational complexity properties, description logics represent frames as classes, with *roles* (binary relations) that link a class to the types of the arguments of the original n-ary relation. Those types are classes as well, so that a graph of frames emerges out of this semantics. The example in 1.31 can be reengineered in DL as follows:

$$T \sqsubseteq \forall R_1. \text{Agent}, T \sqsubseteq \forall R_1^-. \text{Desiring} \quad (1.32)$$

$$T \sqsubseteq \forall R_2. \text{Agent}, T \sqsubseteq \forall R_2^-. \text{Desiring} \quad (1.33)$$

$$T \sqsubseteq \forall R_3. \text{Event}, T \sqsubseteq \forall R_3^-. \text{Desiring} \quad (1.34)$$

$$T \sqsubseteq \forall R_4. \text{Time}, T \sqsubseteq \forall R_4^-. \text{Desiring} \quad (1.35)$$

$$\text{Desiring} \sqsubseteq (=1R_1 \sqcap =1R_2 \sqcap =1R_3 \sqcap =1R_4) \quad (1.36)$$

The computational features of description logics make them a reasonable choice to formally represent linguistic frames, and this is the approach adopted by (SCHEFFCZYK, BAKER, NARAYANAN 2008). On

the other hand, even the description logic solution is limited in formalizing the metalevel conceptualization of frames and schemata. For example, the intended semantics of FrameNet relations between frames, between frame elements, and between frames and frame elements, lexical units, and lexemes is hardly representable in a description logic.

Frames can be subframes of others, can have multiple linguistic units that realize them, multiple lexemes that lexicalize those units, can have frame elements that are core or peripheral, words can evoke frames, etc. Logically speaking, these are second-order relations, and cannot be rebuilt into regular description logic semantics, which is basically first-order. However, recent advancements in higher-order description logics (DE GIACOMO, LENZERINI, ROSATI, 2008) are very promising in order to represent the full range of frame-related relations. See also section 1.4.1.

3. A semiotic metamodel

This section introduces *Semion*, an ontology that represents a semiotic pattern that dates back at least to Peirce (PEIRCE 1958) and Saussure (SAUSSURE 1906).

Peirce used a peculiar terminology, and the versioning of his theory is not trivial. *Semion* encodes a pattern that basically conveys his mainstream ideas: meaning as a role, indirectness of reference, and the dialogic nature of thinking. These ideas have slowly found their way into the literature, and can be formalized by using *c.DnS* as a backbone.

Expressions are information objects used to express a meaning in context at some time. *c.DnS* has contextualization as a primitive assumption, therefore each extension of it assumes a multi-faceted contextualization as depicted in Figure 1.1. The Expression class (that Peirce called «representamen», and Saussure «signifiant») is minimally axiomatized by assuming that an expression is an information object that expresses some schematic entity at some time, and is about some entity at that time:

$$\begin{aligned} \text{Exp}(e) =_{df} & I(i) \wedge \exists(i, e, t) \\ (SE(se) \wedge \text{expresses}(e, se, t) \wedge E(x) \wedge T(t) \wedge \text{isAbout}(e, x, t)) \end{aligned} \quad (1.37)$$

Meanings are schematic entities that are expressed by an expression in context at some time. The Meaning class (that Peirce called «interpretant», and Saussure «signifié») is minimally axiomatized by assuming that a meaning is a schematic entity that is expressed by some information object at some time, and allows the interpretation of some entity at that time:

$$Mea(m) =_{df} SE(m) \wedge \exists(i,e,t)(I(i) \wedge expresses(i,m,t) \wedge E(e) \wedge T(t) \wedge interpretedAs(e,m,t)) \quad (1.38)$$

References are entities, which an expression is about at some time. The Reference class (that Peirce called «object») is minimally axiomatized by assuming that a reference is an entity, which an information object is about at some time, and which is interpreted according to a schematic entity at that time:

$$Ref(r) =_{df} E(r) \wedge \exists(i,se,t)(I(i) \wedge isAbout(i,r,t) \wedge SE(se) \wedge T(t) \wedge interpretedAs(r,se,t)) \quad (1.39)$$

Interpreters are agents, which conceptualize a meaning at some time in an ideal dialogic context with other agents. The Interpreter class is axiomatized by assuming that an interpreter is an agent, which conceptualizes a schematic entity in the context of a situation at some time, which also involves another agent:

$$Int(a) =_{df} A(a) \wedge \exists(se,s,t,a_1)(SE(se) \wedge conceptualized(a,se,t) \wedge S(s) \wedge T(t) \wedge settingFor(s,a) \wedge settingFor(s,se) \wedge settingFor(s,t) \wedge A(a_1) \wedge settingFor(s,a_1)) \quad (1.40)$$

The situation of an interpreter conceptualizing a meaning evoked by an expression, in a context involving another interpreter conceptualizing the same expression, is called here *linguistic act* (*LingAct*). It is related to the notion of *speech act* from (SEARLE 1969), and to the notion of *social act* from (REINACH 1983). In linguistic acts, the interpretive activity of agents in the loop. The *LingAct* class is axiomatized by assuming that a linguistic act is a situation, in which two agents conceptualize two meanings for a same expression at two given time spans, and referring to two entities (the two agents, meanings, time

spans and entities resp. are not necessarily different).¹⁰ Before introducing the class of linguistic acts, the maximal Semion relation is shown, which leverages the c.DnS ‘construction principle’ (GANGEMI 2008) and the previous definitions:

$$\begin{aligned}
 \text{semion}(a, e, m, r, l, t, a_1, m_1, r_1, t_1) =_{df} & (\text{Int}(a) \wedge \text{Exp}(e) \wedge \\
 & \text{Mea}(m) \wedge \text{Ref}(r) \wedge S(l) \wedge T(t) \wedge A(a_1) \wedge \text{Mea}(m_1) \wedge T(t_1) \wedge \\
 & \text{Ref}(r_1) \wedge \text{conceptualized}(a, m, t) \wedge \text{expresses}(e, m, t) \wedge \\
 & \text{interpretedAs}(r, m, t) \wedge \text{conceptualizes}(a_1, m_1, t_1) \wedge \\
 & \text{isAbout}(e, r, t) \wedge \text{expresses}(e, m_1, t_1) \wedge \text{isAbout}(e, r_1, t_1) \wedge \\
 & \text{settingFor}(l, a) \wedge \text{settingFor}(l, a_1) \wedge \text{settingFor}(l, t) \wedge \\
 & \text{settingFor}(l, t_1) \wedge \text{settingFor}(l, m) \wedge \text{settingFor}(l, m_1) \wedge \\
 & \text{settingFor}(l, r) \wedge \text{settingFor}(l, r_1) \wedge \text{settingFor}(l, e) \wedge
 \end{aligned} \quad (1.41)$$

So characterized, an instance of the Semion relation is an occurrence of the semiotic pattern in a community of agents that share some common knowledge.

Now, the LingAct class can be introduced directly by assuming the Semion relation:

$$\begin{aligned}
 \text{LingAct}(l) =_{df} & S(l) \wedge \exists (a, e, m, r, t, a_1, m_1, r_1, t_1) \\
 & (\text{semion}(a, e, m, r, l, t, a_1, m_1, r_1, t_1))
 \end{aligned} \quad (1.42)$$

The Semion approach is pragmatic, in the spirit of Peirce’s: a meaning can be *any* schematic entity, including expressions, concepts, descriptions, collections, or situations, or even a non-schematic entity. Therefore, any meaning is easily representable by specializing the axiom 1.38. For example, the act performed by lexicographers, by which expressions have other expressions as their meanings specializes ($\text{Mea} \subseteq \text{SE}$) as ($\text{Mea} \subseteq \text{I}$). Cognitive theories of meaning, which defend the individual dimension of meaning, can be represented by specializing the axiom 1.38 as ($\text{Mea} \subseteq \text{E}$). Frame semantics (section 1.4.1) can be represented by specializing 1.38 as ($\text{Mea} \subseteq \text{D}$). Extensional formal semantics can be represented by specializing 1.38 as $\text{Mea} \subseteq \text{K}$, etc.

¹⁰ In the dialogic view of semiotics, even an interpreter alone has an ‘internal conversation’.

Moreover, indirectness of reference can be defended or not in some theory of meaning, but in Semion, any such theory can be represented: if some form of conceptualism is taken, the *isAbout* relation can be used with a dependence on a meaning as a mediator; in some form of referentialism, it can be applied directly.

Semion-based models, as exemplified in section 1.4.1, can be transformed into (formal) ontologies by applying a transformation pattern. Since *c.DnS* leverages logical reification, its de-reification is already a transformation pattern; whenever a customization is needed, different patterns can be defined and applied, and the formal choices made are then explicitly represented. See section 1.4.3 for an example.

Since any kind of linguistic act (for example, explanatory text, lexicographic metalanguage, document tagging or indexing, etc.) can be represented as an instantiation of the *LingAct* class, the coverage of Semion is very broad, and ready to apply within any information integration task, for example over the Semantic Web by using its OWL version.¹¹

4. *Applying Semion to lexical resources*

In this section, the application of Semion to some lexical resources is exemplified. Section 1.4.1 describes a part of the FrameNet meta-model based on Semion, and how it allows to create a formal version of FrameNet, called OntoFrameNet. In section 1.4.2 the same procedure is applied to VerbNet and WordNet. In section 1.4.3 Semion is applied to mapping and transformation examples.

4.1. *OntoFrameNet*

FrameNet is a lexical knowledge base, which consists of a set of *frames*, which have proper *frame elements* and *lexical units*, expressed by *lexemes*. Frame elements are unique to their frame, and can be optional. An occurrence of a frame consists in some piece of text whose words can be normalized as lexemes from a lexical unit of a frame, and which have semantic roles dictated by the elements of that frame. A frame can occur with all its roles filled, or not. Frames can be *lexicalized* or not. The non-lexicalized ones typically encode *schemata* from

¹¹ <<http://www.ontologydesignpatterns.org/ont/cdns/semion.owl>>.

cognitive linguistics. Frames, as well as frame elements, are related between them, e.g. through the *subframe* relation. FrameNet contains more information, related to parts of speech, *semantic types* assigned to frames, elements, and lexical units, as well as other metadata.

A complete reengineering of FrameNet (version 1.2) as a *c.DnS* extension can be found in the OWL version of OntoFrameNet.¹² OntoFrameNet is based on *c.DnS*, therefore all FrameNet data are put in the same domain of quantification, by using a reified higher-order approach. This transformation allows to preserve the original schema of the knowledge base, without any loss of information. A limitation is that the automated reasoning over OntoFrameNet occurs mainly at the OWL ABox level, the assertional part of an ontology, (BAADER, CALVANESE, MCGUINNESS, NARDI, PATEL-SCHNEIDER 2003). Anyway by using the full set of semiotic ontologies in the LMM umbrella (PICCA, GANGEMI, GLIOZZO 2008), a custom translation of selected parts of OntoFrameNet to a TBox can be performed with an explicit semantics (cf. section 1.4.3 below).

The backbone of FrameNet is the notion of a *Frame*. As the authors pragmatically state (RUPPENHOFER, ELLSWORTH, PETRUCK, JOHNSON, SCHEFFCZYK 2006): «with enough time to make a truly in-depth analysis of the data, and enough data to make an exhaustive account of the language, then undoubtedly each lexical unit could be given its own unique description in terms of the frames and/or subframes which it evokes. The situation is, in a sense, worse than the question suggests: it isn't that every word has its own frame, but every sense of every word (i.e., every lexical unit) has its own frame. It's a matter of granularity. Instead, we are sorting lexical units into groups in the hope that they permit parallel analyses in terms of certain basic semantic roles, i.e., the frame elements that we have assigned to the frame. This allows us (1) to make the sorts of generalizations that should be helpful to the users mentioned above and (2) to provide semantically annotated sentences that can exemplify paraphrase relations within given semantic domains».

In practice, FrameNet is trying to find schematic invariances in the conceptual structures of linguistic usage, in order to reduce the complexity of expliciting all the schemata applicable to each word sense. This hypothesis is compatible with the cognitive linguistics paradigm,

¹² <<http://www.ontologydesignpattern.org/ont/ofn/ofntb.owl>>.

with intensional relations in formal semantics, as well as with the *c.DnS* reified relational ontology. The core *OntoFrameNet* metamodel consists of the following relation (1.43):

$$\begin{aligned} & \text{FrameNetRel}(f, fe, st, lu, l) \rightarrow \\ & \text{Frame}(f) \wedge \text{FE}(fe) \wedge \text{hasFe}(f, fe) \wedge \\ & \text{SemanticType}(st) \wedge \text{hasSemType}(fe, st) \wedge \\ & \text{LexicalUnit}(lu) \wedge \text{hasLU}(f, lu) \wedge \\ & \text{Lexemel}(l) \wedge \text{hasLexemel}(lu, l) \end{aligned} \quad (1.43)$$

For example:

$$\begin{aligned} & \text{FrameNetRel}(\text{F_Desiring}, \text{FE_Event_3363}, \text{StateOfAffairs}, \\ & \text{LU_desire.v_6413}, \text{LEX_desire_10357}) \end{aligned} \quad (1.44)$$

The projections of the frame relation characterize its arguments further: for each frame there are one or more elements, but each element is unique to one frame. For each frame there are one or more lexical units (senses), but each unit is unique to a frame. For each frame, frame element or lexical unit there should be a semantic type (in the core relation, only the type of the frame element is mandatory). Moreover, several relations create further ordering between frames, and between frame elements. Unfortunately, *FrameNet* data are not complete: for example, many frame elements are still missing a semantic type.

In *OntoFrameNet*, besides formalizing the metamodel and creating inverse projections where needed, some additions have been implemented; for example, a *generic* frame element has been created for sets of frame elements with the same name: in this way it is possible to run more sophisticated queries in order to measure frame distance (e.g. finding those sharing two generic frame elements except *Space* and *Time*). Moreover, situations corresponding to occurrences of frames in the interaction with the environment, as expressed by textual sentences (e.g. those annotated in *PropBank* with frames and frame elements), have been given room in a newly created class (*FrameOccurrence*).

The alignment is summarized as follows, and a shortened axiomatization is presented in axioms from 1.45 to 1.55. *Frames* are aligned as meanings in *Semion*, and since frames have a relational structure (as conceptual contexts), they are more specifically aligned as descriptions (reified intensional relations).

The relations between frames have been aligned consequently: the frame *inheritance* relation as *c.DnS* specialization, the *subframe* relation as proper part, etc.

The *evokes* relation between lexemes or lexical units, and frames is aligned to the *evokes_r* relation.

Frame elements are ‘FEin’ a frame, and are aligned as meanings, and as concepts (uniquely) defined in a frame.

Lexical units are aligned as meanings, and as descriptions, expressed by a specific collection of lexemes, which is also a reportable construction.

Lexemes are aligned as expressions.

Occurrences of frames are aligned as situations that are expressed (*evoked_r*) by expressions (as sentences).

$$\text{Frame} \subseteq (\text{Mea} \cap D) \quad (1.45)$$

$$\text{inheritsFrom}(f_1, f_2) \rightarrow \text{specializes}(f_1, f_2) \wedge \text{Frame}(f_1 \wedge \text{Frame}(f_2)) \quad (1.46)$$

$$\text{isSubFrameOf}(f_1, f_2) \rightarrow \text{isProperPartOf}(f_1, f_2) \wedge \text{Frame}(f_1 \wedge \text{Frame}(f_2)) \quad (1.47)$$

$$\text{evokes}(x, y) \rightarrow \text{evokes_r}(x, y) \quad (1.48)$$

$$\text{hasFE}(f, fe) \rightarrow \text{defines}(f, fe) \wedge F(f) \wedge FE(fe) \quad (1.49)$$

$$FE(fe) =_{df} \text{Mea}(fe) \wedge C(fe) \wedge \exists!(f) (\text{Frame}(f) \wedge \text{defines}(f, fe)) \quad (1.50)$$

$$\text{GenFE}(gfe) =_{df} \text{Mea}(gfe) \wedge C(gfe) \wedge \exists(fe) (\text{FE}(fe) = \wedge \text{specializes}(fe, gfe) \wedge \neg \exists(f) (\text{Frame}(f) \wedge \text{defines}(f, gfe))) \quad (1.51)$$

$$\text{LexicalUnit}(lu) \rightarrow \text{Mea}(lu) \wedge D(lu) \wedge \exists(ag, l, t) (K(ag) \wedge \text{Lexemel}(l) \wedge \text{hasProperPart}(ag, l) \wedge \text{expresses}(ag, lu, t)) \quad (1.52)$$

$$\text{Lexeme}(le) \rightarrow \text{Exp}(le) \quad (1.53)$$

$$\text{SemanticType} \subseteq (\text{Mea} \cup C) \quad (1.54)$$

$$\text{FrameOccurrence}(fo) =_{df} S(fo) \wedge \exists(f, ag, lu, t) (\text{Frame}(f) \wedge K(ag) \wedge \text{satisfies}(fo, f) \wedge \text{Exp}(le) \wedge \text{hasProperPart}(ag, le) \wedge \text{evokesxr}(ag, fo, t)) \quad (1.55)$$

Superficially, the linguistic act involved in FrameNet consists in a *metalinguistic* function (JAKOBSON 1990), typical of lexica, dictionaries, etc., in which an agent assigns meanings to expressions, and the observable situation of the act is a linguistic situation. On the other

hand, frame semantics tries to reach out to language usage, not only to an abstract characterization of linguistic items; as a matter of fact, frame occurrences, as denoted by annotations made over the Prop-Bank corpus, are real world situations (explanatory, expressive, etc.), not metalinguistic ones. This hybrid nature of frame semantics distinguishes it from e.g. WordNet-based annotations of corpora, where real world situations cannot be denoted in a relational way.

While other lexical resources, such as WordNet and VerbNet, have not the same groundedness as FrameNet, nonetheless they are widely used and contain a lot of reusable content that can be combined effectively with FrameNet.

4.2. *OntoVerbNet and WordNet*

VerbNet (KIPPER, DANG, PALMER 2000), has a different metamodel from FrameNet; it is focused on verb syntax and semantics, rather than frame semantics, which abstracts out of parts of speech. A new metamodel (called OntoVerbNet) has been created, which is shown partly in the maximal semantic relation 1.56 (some names from the original relational database schema have been changed for readability):

$$\begin{aligned}
 & \text{VerbNetRel}(vn, fr, pr, ar, ca) \rightarrow \\
 & \text{VNClass}(vn) \wedge \text{VNFrame}(fr) \wedge \text{hasFrame}(vn, fr) \wedge \\
 & \text{Predicate}(pr) \wedge \text{hasPred}(fr, pr) \wedge \text{Argument}(ar) \wedge \\
 & \text{hasArg}(pr, ar) \wedge \text{Category}(ca) \wedge \text{hasType}(ar, ca) \quad (1.56) \\
 & \text{VNClass}(un) \rightarrow \exists(v)(\text{Verb}(v) \wedge \text{hasMember}(un, v)) \quad (1.57)
 \end{aligned}$$

For example:

$$\begin{aligned}
 & \text{VerbNetRel}(\text{battle 36.4}, fr_{8.1}, \text{social interaction}, \\
 & \text{Actor1}, \text{animate}) \quad (1.58)
 \end{aligned}$$

$$\text{hasMember}(\text{battle 36.4}, \text{argue}) \quad (1.59)$$

The OntoVerbNet interpretation over VerbNet represents different lexical semantics for each ‘verb class’, trying to catch the basic semantic structure of VerbNet, consisting of typed arguments holding for a predicate in a ‘frame’ that contributes to the complete semantics of a verb class; frames also have syntactic constructs applicable to the verb, to which the frame is applied.

In the example 1.58, the verb class *battle* has a frame (including both syntactic and semantic specifications), with some predicates (*social interaction*, *conflict*, *about*), each having arguments (e.g. *Actor1*), with a

category (e.g. *animate*). One or more verbs are member of the verb class. For each verb class, more than one frame for a verb class, predicate for a frame, and argument/category for a predicate can be asserted.

VerbNet relies on a small amount of primitives (about 100 predicates, 70 arguments and 40 categories in version 2.1) to account for the semantics of verbs. No assumption of uniqueness of arguments or predicates for a frame are made. The VerbNet approach is therefore closer to traditional linguistic theories, and it is not trivial to match it to FrameNet construction grammar. Here some alignment suggestions are shown¹³ which can help on achieving that task:

$$V \text{ NClass} \subseteq (Mea \cap D) \quad (1.60)$$

$$subClass(x, y) \rightarrow specializes(x, y) \wedge V \text{ NClass}(x) \wedge V \text{ NClass}(y) \quad (1.61)$$

$$V \text{ NFrame} \subseteq (Mea \cap D) \quad (1.62)$$

$$hasFrame(vn, fr) \rightarrow hasProperPart(vn, fr) \wedge V \text{ NClass}(vn) \wedge V \text{ NFrame}(fr) \quad (1.63)$$

$$Predicate \subseteq (Mea \cap D) \quad (1.64)$$

$$hasPredicate(fr, pr) \rightarrow hasProperPart(fr, pr) \wedge V \text{ NFrame}(fr) \wedge Predicate(pr) \quad (1.65)$$

$$Argument(ar) \subseteq (Mea \cap C) \quad (1.66)$$

$$Argument(ar) \rightarrow \exists (pr) (usesConcept(pr, ar)) \quad (1.67)$$

$$hasArg(pr, ar) \rightarrow usesConcept(pr, ar) \wedge Predicate(pr) \wedge Argument(ar) \quad (1.68)$$

$$Category \subseteq (Mea \cap C) \quad (1.69)$$

$$hasType(ar, ca) \rightarrow specializes(ar, ca) \wedge Argument(ar) \wedge Category(ca) \quad (1.70)$$

$$Verb \subseteq Exp \quad (1.71)$$

$$hasMember(vn, v) \rightarrow expresses(v, vn) \wedge Verb(v) \wedge V \text{ NClass}(vn) \quad (1.72)$$

In practice, a VerbNet predicate is comparable to a FrameNet frame, but it is not unique to a verb class, while a VerbNet argument is comparable to a FrameNet frame element, but again it is not unique to a predicate. The semantic part of VerbNet frames is comparable to a composition of FrameNet frames. These differences are due to the fact that VerbNet focuses on verb classes rather than on conceptual structures.

¹³ <<http://www.ontologydesignpatterns.org/ont/lmm/ovn2lmm.owl>>.

On the other hand, based on OntoFrameNet and OntoVerbNet, it is easier to compare the two lexical knowledge bases on a formal basis, e.g. by restricting the matching to VerbNet arguments against FrameNet generic frame elements, or by finding recurrent arguments in VerbNet predicates, and trying to approximate core predicate structures.

Now, since FrameNet frames can be matched against VerbNet predicates, one can check the consistency between core frame elements and arguments shared across the predicates that hold for different verb classes. Moreover, FrameNet frames can be matched against VNFrames: we will check the consistency of the about 100 predicates from VerbNet as a top-level for FrameNet frames.

VerbNet arguments seem to match frame elements: in that case, argument categories can be matched to or used to populate FrameNet semantic types for frame elements when missing. Whatever matching pattern is taken, one will know what entities are involved in the matching, and what consequences will derive. For example, VerbNet arguments are not unique to predicates, while frame elements are unique to frames, therefore it is more appropriate to match OntoVerbNet arguments with OntoFrameNet generic frame elements.

Additional metamodels can be added in order to add new matching possibilities in lexical resource integration. WordNet is a first-class candidate because it is extensively used, its metamodel is already built, and an alignment is straightforward, e.g. the Synset class can be aligned as in axiom 1.73, and it can be used to feed argument and frame element with semantic types of a finer granularity. VerbNet categories can then be matched against synsets, and possibly proposed as an alternative top-level for synsets, comparable to WordNet *lexical names*, also known as ‘super-senses’. The following sample axioms make it viable to map VerbNet categories to super-senses, synsets to super-senses, and therefore synsets to categories:

$$\text{Synset} \subseteq (\text{Mea} \cap \text{C}) \quad (1.73)$$

$$\text{Super Sense} \subseteq (\text{Mea} \cap \text{C}) \quad (1.74)$$

$$\forall (x)(\text{Super Sense}(x) \rightarrow \exists (y)(\text{Synset}(y) \wedge \text{specializes}(x, y)) \quad (1.75)$$

$$\forall (x, y, z)((\text{mappable To}(x, y) \wedge \text{specializes}(z, x)) \rightarrow \text{specializes}(z, y)) \quad (1.76)$$

$$\begin{aligned} \text{mappableTo}(x, y, r_1, r_2) \rightarrow \text{Mea}(x) \wedge \text{Mea}(y) \wedge \\ \text{Resource}(r_1,) \wedge \text{Resource}(r_2) \wedge \\ \text{belongsTo}(x, r_1,) \wedge \text{belongsTo}(y, r_2,) \end{aligned} \quad (1.77)$$

5. Mapping and transformation patterns

Comparison between Semion-based elements can be formalized by defining appropriate relations, which are used as *mapping patterns* between elements from different lexical resources:

$$\begin{aligned}
 & \text{mappableConcept}(x, y, r_1, r_2) =_{df} \\
 & \text{mappableTo}(x, y, r_1, r_2) \wedge C(x) \wedge (C(y)) \quad (1.78) \\
 & \text{mappableConcept FN2V N}(x, y, r_1, r_2) =_{df} \\
 & \text{mappableConcept}(x, y, r_1, r_2) \wedge \\
 & (r_1 = \text{FrameNet}) \wedge (r_2 = \text{VerbNet}) \\
 & (\text{GenFE}(x) \wedge \text{Argument}(y)) \vee \\
 & (\text{SemanticType}(x) \wedge \text{Category}(y)) \quad (1.79)
 \end{aligned}$$

for example, based on the mapping pattern in 1.79, one can safely assert that a certain generic frame element abstracted from FrameNet, e.g. ofn:Agent, is mappable to a VerbNet argument, e.g. ovn:Agent:

$$\begin{aligned}
 & \text{mappableConceptFN2V N}(\text{Agent}, \text{Agent}, \text{FrameNet} \\
 & \text{VerbNet}) \quad (1.80)
 \end{aligned}$$

Finally, a sample *transformation pattern* shows how Semion constructs (e.g. a $(\text{Mea} \cap C)$) can be transformed to some formal semantic construct, e.g. a *Class*:

$$\begin{aligned}
 & \forall (x, y)(\text{conceptsTransformableToClasses}(x, y) \rightarrow \\
 & ((\text{Mea}(x) \wedge C(x)) \rightarrow \text{Class}(y)) \quad (1.81)
 \end{aligned}$$

When adopting axiom 1.81, we accept that any lexical resource element x that is an instance of a lexical resource element type aligned to $(\text{Mea} \cap C)$ can only be encoded as an ontology element y that has *Class* semantics, e.g. an owl:Class (in the Web Ontology Language). Therefore, by adopting 1.81, we also assume that all FrameNet, VerbNet, or WordNet elements that are instances of element types aligned to $(\text{Mea} \cap C)$ must be encoded as classes, so that formal operations on them will be founded on a(n adopted) shared semantics.

Appropriate recipes including transformation patterns can be used to manage large integration scenarios on heterogeneous lexical knowledge.

6. Conclusions

This article introduces a formal method to represent the intended semantics of lexical data, and to reconcile it with formal semantics, in a customizable way. The method described employs the flexibility and expressive power of c.DnS, a ontology of contexts and schematic entities. Based on c.DnS, a metamodel for lexical data, called Semion, has been described, and exemplified on well-known resources like WordNet, VerbNet and FrameNet. Mapping and transformation patterns have been exemplified as well.

Current work on Semion is focused on its extension to heterogeneous resources that contain lexical data, such as DBPedia and DMOZ, and on devising Semion plugins that describe the model underlying typical NLP tasks, such as taxonomy induction, frame detection, latent semantic indexing, word sense disambiguation, etc. Application work concentrates on large-scale data integration, where different datasets are aligned to Semion by means of appropriate metamodels.

ACKNOWLEDGEMENTS: I thank Alfio Gliozzo and Davide Picca, who have co-authored the LMM meta-model, an inspiration (and an applicative context) for Semion.

7. References

The Description Logic Handbook: Theory, Implementation, and Applications, ed. by F. Baader, D. Calvanese, D.L. McGuinness, D. Nardi, P.F. Patel-Schneider, New York, Cambridge University Press 2003.

C.F. BAKER, C.J. FILLMORE, J.B. LOWE, *The Berkeley FrameNet Project*. Proceedings of the 1998 International Conference on Computational Linguistics (COLING-ACL).

S. BECHHOFFER, F. HARMELEN, J. HENDLER, D. HORROCKS, I. MCGUINNESS, P. PATEL-SCHNEIDER, L.A. STEIN, *OWL Web Ontology Language Reference*. Technical report, W3C 2004.

T. BERNERS-LEE, W. HALL, J.A. HENDLER, N. SHADBOLT, D.J. WEITZNER, *Web Science*, in «Science», CCCXIII, 2006.

R.J. BRACHMAN, *A Structural Paradigm for Representing Knowledge*. Ph.d. thesis, Harvard University 1977.

P. BUITELAAR, K. CHOI, A. GANGEMI, C. HUANG, A. OLTRAMARI, Proc. of ISWC 2007 Workshop OntoLex2007, Busan , Korea <<http://olp.dfki.de/OntoLex07/>>.

P. CIMIANO, P. HAASE, M. HEROLD, M. MANTEL, P. BUITELAAR, *LexOnto: A Model for Ontology Lexicons for Ontology-based NLP*, 2007 <<http://olp.dfki.de/OntoLex07/>>.

D. DAVIDSON, *The Logical Form of Action Sentences*, in *The Logic of Decision and Action*, Pittsburgh, University of Pittsburgh Press 1967 (2nd ed.).

G. DE GIACOMO, M. LENZERINI, R. ROSATI, *Towards Higher-Order DL-Lite*. Proc. of the 21st International Workshop on Description Logics (DL2008), Dresden 2008.

E. DE LUCA, M. EUL, A. NÜRNBERGER, *Multilingual Query-Reformulation using an RDF-OWL EuroWordNet Representation*, 2007.

U. ECO, *Trattato di semiotica generale*, Milano, Bompiani 1975.

WordNet. An Electronic Lexical Database, ed. by C. Fellbaum, Cambridge, MIT Press 1998.

C. FILLMORE, *Frame semantics*, in «Linguistics in the Morning Calm», 1982, pp. 111-137.

C. FILLMORE, P. KAY, C. O'CONNOR, *Regularity and Idiomaticity in Grammatical Constructions: The Case of Let Alone*, in «Language», LXIV, 1988, pp. 501-538.

G. FRANCOPOULO, M. GEORGE, N. CALZOLARI, M. MONACHINI, N. BEL, M. PET, C. SORIA, *Lexical Markup Framework (LMF)*. Proc. of the International Conference on Language Resources and Evaluation (LREC), Genova 2006.

A. GALTON, *Time and Change for AI*, in «Handbook of Logic in Artificial Intelligence and Logic Programming», IV, 1995, pp. 175-240.

A. GANGEMI, *Ontology Design Patterns for Semantic Web Content*. Proceedings of the Fourth International Semantic Web Conference, Galway, Ireland, ed. by M. Musen *et al.*, Berlin, Springer 2005.

A. GANGEMI, *Norms and Plans as Unification Criteria for Social Collectives*, in «Journal of Autonomous Agents and Multi-Agent Systems», XVI, 3, 2008.

A. GANGEMI, *What's in a Schema?*, in *Ontology and the Lexicon*, ed. by C. Huang, N. Calzolari, A. Gangemi, A. Lenci, A. Oltramari, L. Prevot, Cambridge University Press 2010.

A. GANGEMI, N. GUARINO, C. MASOLO, A. OLTRAMARI, *Sweetening WordNet with DOLCE*, in «AI Mag.», XXIV, 3, 2003, pp. 13-24.

T. GRUBER, *Ontology of Folksonomy: A Mash-up of Apples and Oranges*, in «Int'l Journal on Semantic Web and Information Systems», III, 2, 2007.

N. GUARINO, *Formal ontology and information systems*, in *Formal Ontology in Information Systems*. Proceedings of FOIS'98, 1998, pp. 3-17.

Ontology and the Lexicon, ed. by C. Huang, N. Calzolari, A. Gangemi, A. Lenci, A. Oltramari, L. Prevot, Cambridge, Cambridge University Press 2010.

R. JAKOBSON, *Closing Statements: Linguistics and Poetics*, in *Style In Language*, ed. by T.A. Sebeok, Cambridge, MIT Press 1990, pp. 350-377.

M. JOHNSON, *The Body in the Mind: the Bodily Basis of Meaning, Imagination, and Reason*, Chicago-London, University of Chicago Press 1987.

K. Kipper, H.T. Dang, M. Palmer, *Class-Based Construction of a Verb Lexicon*. AAAI-2000 Seventeenth National Conference on Artificial Intelligence.

G. LAKOFF, *Women, Fire, and Dangerous Things: What Categories Reveal about the Mind*, Chicago, University of Chicago Press 1987.

R. LANGACKER, *Concept, Image, and Symbol: The Cognitive Basis of Grammar*, Berlin-New York, Mouton de Gruyter 1990.

C. MASOLO, L. VIEU, E. BOTTAZZI, C. CATENACCI, R. FERRARIO, A. GANGEMI, N. GUARINO, *Social Roles and their Descriptions*. Proc. of Knowledge Representation and Reasoning (KR), ed. by C. Welty, D. Dubois, AAAI Press 2004, pp. 267-277.

S. McCORMICK, C. LIESKE, A. CULUM, *A Flexible Language Data Standard*, OLIF II 2004 <<http://www.olif.net/>>.

M. MINSKY, *A Framework for Representing Knowledge*, in *The Psychology of Computer Vision*, ed. by P. Winston, McGraw-Hill 1975.

M. PAZIENZA, A. STELLATO, *Linguistic Enrichment of Ontologies: a Methodological Framework*, Workshop on Interfacing Ontologies and Lexical Resources for Semantic Web Technologies (OntoLex 2006), 2006.

C.S. PEIRCE, *On Signs and the Categories*, in *Collected Papers of C.S. Peirce*, ed. by C. Hartshorne, P. Weiss, A. Burks, Cambridge, Harvard University Press 1958.

W. PETERS, E. MONTIEL-PONSODA, G. AGUADO DE CEA, A. GÚMEZ-PÉREZ, *Localizing Ontologies in OWL*, 2007 <<http://olp.dfki.de/OntoLex07/>>.

D. PICCA, A. GANGEMI, A. GLIOZZO, *LMM: an OWL Metamodel to Represent Heterogeneous Lexical Knowledge*. Proc. of the International Conference on Language Resources and Evaluation (LREC), Marrakech, Morocco, 2008.

V. PRESUTTI, A. GANGEMI, S. DAVID, G. AGUADO DE CEA, M. SUAREZ-FIGUEROA, E. MONTIEL-PONSODA, M. POVEDA, *Library of design patterns for collaborative development of networked ontologies*, Deliverable D2.5.1, NeOn project 2008.

A. REINACH, *The Apriori Foundations of the Civil Law*, in «Aletheia», III, 1983, pp. 1-142.

J. RUPPENHOFER, M. ELLSWORTH, M.R.L. PETRUCK, C.R. JOHNSON, J. SCHEFFCZYK, *FrameNet II: Extended Theory and Practice*, 2006 <<http://framenet.icsi.berkeley.edu/book/book.html>>.

F.D. SAUSSURE, *Cours de linguistique générale*, Lausanne, Payot 1906.

J. SCHEFFCZYK, C.F. BAKER, S. NARAYANAN, *Reasoning over Natural Language Text by Means of FrameNet and Ontologies*, in *Ontology and the Lexicon*, ed. by C. Huang, N. Calzolari, A. Gangemi, A. Lenci, A. Oltramari, L. Prevot, Cambridge, Cambridge University Press 2008.

J.R. SEARLE, *Speech Acts: An Essay on the Philosophy of Language*, Cambridge, Cambridge University Press 1969.

J.R. SEARLE, *The Construction of Social Reality*, New York, Free Press 1995.

L. TALMY, *Toward a Cognitive Semantics*, Cambridge, MIT Press 2003.

W3C 2004, *OWL Web Ontology Language Family of Specifications* <<http://www.w3.org/2004/OWL>>.

C. WELTY, *Formal Ontology for Subject*, in «J. Knowledge and Data Engineering», XXXI, 1999, pp. 155-182.

ALDO GANGEMI

Systems biology challenges computer science

ABSTRACT: The definition of theory and methods to model, simulate, analyze and control biological systems poses new challenges to the computer science community because the complexity of such systems is many order of magnitude greater than actual information technology platforms.

We summarize the underpinning ideas of process algebras for mobility and we claim that they are the right abstraction (language) for systems biology. Our claim is supported by the analysis of important related issues: the interfaces which could be used to hide unwanted mathematical details; the computational environments that could be developed and exploited; the potential impact of the approach.

1. *Introduction*

The Human Genome Project (HGP) dates back to 1985 in its early formulation when Robert Sinsheimer arranged a meeting on the human genome sequencing at the University of California, Santa Cruz. The idea is to completely decode the human genome and to store it in a readable and accessible form. The human genome is made up of 22 pairs of chromosomes (autosomes) plus the X and Y chromosomes (for determining sex – a pair XX for female and a pair XY for male). The 23 pairs of chromosomes contain the whole information needed to build a human being from an egg. Chromosomes are made up of very long DNA molecules. Somewhere in the DNA molecules, genes are hidden. They are simply sections of a chemical message that runs along the spine of the DNA. Four molecular bases or nucleotides run the full length of the DNA molecule: Adenine (A), Guanine (G), Cytosine (C) and Thymine (T). The order of the nucleotides instructs the cells on how to build proteins that then regulate the behaviour of the cells themselves.

In language terms, the genome is a set of words built upon an alphabet of 4 characters denoting nucleotides. This is the abstraction that made feasible the HGP. Actually any breakthrough step in research is possible when simple abstractions and models are available for complex real systems.

The great amount of data to be analyzed and checked in the genome project made mandatory the use of computers and related software. The term *bioinformatics* was coined in 1987 at the Supercomputer Computations Research Institute by H.A. Lim, and it was later defined to be *the study of information content and information flow in biological systems*. With heavy use of computers and bioinformatics the HGP reached its first milestone in June 2000 when U.S. and U.K. jointly announced completion of the first draft of the human genome. Eventually in February 2001 Science published the data available from the private sector involved in the HGP and Nature published the data from the public sector.

In spite of the definition of bioinformatics given above in which information flow is mentioned, the HGP only deals with information content. A set of words is produced out of 4 characters. This is very related to the definition of the syntax of a finite language relying on an enumerative approach. Given an alphabet, we list all the possible words belonging to the language that we are defining: the *human genome language*. From computer science we know that the difficult and interesting part in the definition of a language is the assignment of meaning to words and the definition of rules to compose them: *the semantics of the language*. However we are very far from giving a semantics to the human genome language that reflects the actual behaviour of cells.

Cells, and more generally biological systems, are self-controlled and self-regulating dynamic complex systems consisting of components that interact in space and time. The relationships that prevail between structure, function, and regulation in cellular networks are still largely unknown. The great complexity of such systems as well as the costs of real experiments impose the need of principles and models of functioning of adaptive biological systems in order to provide formal methods of analysis. This approach should lead to a better understanding, modelling and simulation of complex bio-systems. Furthermore, organic systems are made up of a hierarchy of components at a different level of functioning that interact perfectly to form a whole: genes, cells, organs, organisms. This particular structure of organic systems should help to scale understanding from models of parts of complex adaptive

systems to a larger whole system thus giving hints on the compositional properties of these systems. Summing up biological systems are *self-controlled, self-regulating, self-organizing into hierarchical structures, autonomous, dynamic, scalable, communicating* complex systems with no centralized point of control and where any component has only a local, limited knowledge of the whole system.

The definition of theory and methods to model, simulate, analyze and control biological systems poses new challenges to the computer science community in the large because the complexity of such systems is many order of magnitude greater than actual Information Technology (IT) platforms even considering Internet. So the understanding of biological systems is the first step towards the attempt of devising new design, modelling and control paradigms for a new generation of IT complex systems via reverse engineering of organic systems.

The very same challenge is posed to biologists that start studying biological systems as a whole rather than sectioning them in components and looking at single sub-parts in isolation. This new research approach is imposing a paradigmatic shift from the classical reductionism to a hypothesis driven research as stated by the so-called systems biology.

In this paper we recall the underpinning ideas of process algebras for mobility and we claim that they are the right abstraction (language) for systems biology. Our claim is supported by the analysis of important related issues: the interfaces which could be used to hide the unwanted mathematical details from biologists; the computational environments that could be exploited; the potential impact of the approach.

2. Process Algebras for Modeling Bio-Systems

The characteristics of biological systems outlined above are the same as those of distributed and mobile systems. Many processes are active simultaneously over a set of physical resources for which they compete while cooperating to accomplish a common goal. Acquisition of a resource from a process or reception of a message upon which choices have to be taken can surely affect the future behavior of the whole system and even change the logical interconnection structure amongst processes. Trust barriers and administrative domains work as membranes that can be passed only by those processes that possess the right keys – hence the concept of localization of processes is an im-

portant one. This feature can be handled by mobile process algebras, the first of which (π -calculus)¹ was introduced by Milner, Parrow and Walker in the eighties.

Process algebras have been introduced to model rigorously concurrent systems, namely systems made up of a set of communicating and cooperating processes. Processes are syntactically described by means of few operators used to compose elementary actions (say α) over distributed channels (denoted hereafter by their names, given in small case letters). These operators are: sequentialization ($\alpha.P$), parallel composition ($P \mid Q$), name declaration ((νx)), and recursion ($\text{rec } x:P$).

The intuitive meaning of sequentialization is that the atomic action α is the first that the process $\alpha.P$ can execute. Atomic actions can be the output of a name b over a channel a ($a!b$), the reception of a datum on a channel a that will replace the placeholder x in the prefixed process ($a?x$), or an internal action of the system (τ). Reception $a?x$ binds the free occurrences of the variable x within P . The ν operator in $(\nu x)P$ declares x to be private to P . The parallel composition $P \mid Q$ allows the processes P and Q to be executed independently of one another and they can communicate if they share a communication channel. For instance $a!b:P \mid a?y:Q$ can perform a communication by sending b over the channel a from the left hand process to the right hand one yielding $P \mid Q \{b/y\}$, with $\{b/y\}$ being the substitution of b for the free occurrences of y in Q . Eventually, $\text{rec } x:P$ stands for the possible unfolding of process P as many times as needed. Sometimes recursion is represented through the *bang* operator ($!P$) that is interpreted as many copies of P as needed.

The formal semantics of process algebras is usually given in terms of the logics based Plotkin's structural operational semantics². The dynamic behavior of systems is then expressed in terms of transition systems, i.e. oriented and labeled graphs where the nodes are the states of the system and the arcs represent, via their labels, the actions that make the system change from one state to the other.

An important property of mobile process algebras is *compositionality*. The meaning or the behavior of a complex system is expressed in terms of the meaning of its components. This allows one to concentrate

¹ R. MILNER, J. PARROW, D. WALKER, *A Calculus of Mobile Processes (I and II)*, in «Information and Computation», C, 1992, pp. 1-77.

² G.D. PLOTKIN, *Structural Approach to Operational Semantics*, in «Technical Report DAIMI-FN-19», 1981.

on the basic operations that a system can perform and to obtain the whole behavior through composition of these basic building blocks. Compositionality is surely the key issue needed by systems biology to become effective. Indeed the huge amount of data available and the complexity of the interacting networks analyzed, make it impossible to define *formally* the behavior of biological systems when they are considered as a whole.

The paradigmatic interpretation of biological systems into process algebras for mobility can be described as follows. Processes are the biological components. Sharing of channels establishes the interconnection topology of the system and represents the interaction potentials of components together with their affinity. Scope of channels represents the boundaries within which interactions through such channels may occur. Since channel names can be sent as data along channels, the interconnection topology varies dynamically so modeling the impact of an interaction on the future behavior of the whole system. The above interpretation immediately provides a dynamic description of the temporal evolution of the system in hand: we only need to run the program.

Biological models are driven by a lot of quantitative information concerning for instance energy, time, affinity, distance, electrostatic charge, number of components. Therefore on the computer science side we must resort to stochastic variants of mobile process algebras like, e.g., the stochastic π -calculus³, and the biochemical stochastic π -calculus⁴. These extensions provide the formal tools to include numbers in system specifications.

We now have the computer science counterpart of systems biology: models are specified by using stochastic process algebras, experiments are carried out in silico (rather than wet biology) relying on analysis, verification and simulation techniques developed in the last decades in the field of concurrency theory and refined for the new applicative domain.

The neat result is that we can try hundreds of experiments without

³ C. PRIAMI, *Stochastic π -Calculus*, in «The Computer Journal», XXXVIII, 1995, pp. 578-589.

⁴ C. PRIAMI, A. REGEV, W. SILVERMAN, E. SHAPIRO, *Application of a Stochastic Name-Passing Calculus to Representation and Simulation of Molecular Processes*, in «Information Processing Letters», LXXX, 2001.

using reactants or animals – of course we must tune the techniques adopted by comparing the *in silico* results with real wet experiments. Once the framework is tuned we can use it in a predictive fashion. *In silico* experiments can help biologists selecting real experiment strategies among a plethora of possible ones. For instance it is possible to investigate how a new drug can break a signaling pathway leading to a disease by leaving unaffected as many functionalities of a complete signaling network as possible (reduction of side effects of drugs), or even how transduction mechanisms leading to DNA damage are activated.

The main added value for systems biology in joining this computer science approach is given by the abstraction mechanisms that computer scientists have been developing for concurrent systems over the last thirty years, so that a considerable speed up in life science research could be possible especially on the fields of predictive, preventive and personalized medicine, drug design and gene therapy, toxicological research, environmental research. There is a general understanding in the scientific community that computer science will be as indispensable for biology as mathematics has been for physics.

On the other hand it should be made clear that computer science cannot be only a service for biology otherwise the model cannot work due to the different expectations of the two research communities. The best way to proceed is to create a new interdisciplinary community in which all the components have the same weight. A new definition of bioinformatics by NCBI and available at the url <http://www.ncbi.nlm.gov/Education> catches exactly this point: *bioinformatics is the science in which biology and computer science join together to ease new biological discovery and to define new computational paradigms inspired by living systems.*

The added value for computer science in joining systems biology is well-expressed in the last part of the definition above. The identification of abstraction mechanisms to model living systems may hopefully provide us with new computational and information processing paradigms that could successfully be applied in programming the net at various level of coordination: from grid computing to global computing. Another obvious application of reusing biological information in the computer science domain concerns security: once we have a complete model of the immune systems we can recreate the same conditions in artificial systems like the net (also this strategy begins to be investiga-

ted, see e.g. a recent work by Chao and Forrest)⁵.

A first step towards the definition of new computational models is the understanding of the living systems that can be accomplished if we provide biologists with integrated environments in which they can model and simulate systems. This poses two technological challenges. First, we should provide them with environments that hide from the user as many technical details as possible. Second, we should be able to fully integrate the new tools with the ones that are already in use, like data bases and data/literature mining frameworks. These two technological aspects are separately expanded below.

3. *Integrated Computational Environments*

Public data bases like KEGG or BioCarta are becoming an indispensable tool for biologists because they contain almost all the information available on genomes and regulatory or signaling networks. Unfortunately they lack dynamic information on the behavior of biological systems. At their best, they can include multiple references to published papers where the relevant behavioral aspects are separately described. Biologists rely on literature mining tools to recover the papers that are most interesting for their cases. Most of these data bases support exportation of their models through XML. However, there exist no standard representation of the information stored in different data bases, and this makes it difficult to compare data.

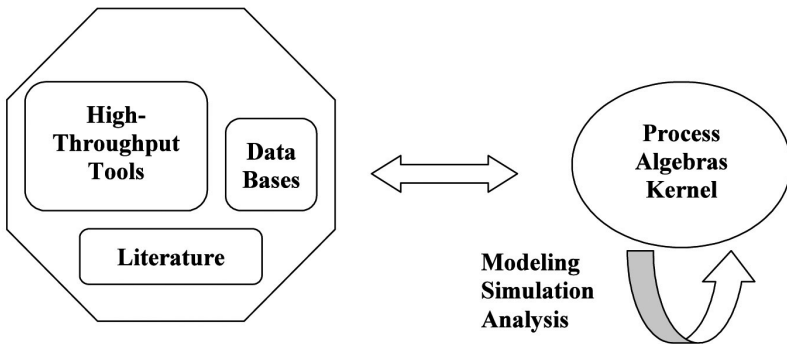
Integration of this framework with the environments we proposed can happen in two ways. The information stored into the data bases can be used to (semi-)automatically extract/infer process algebra models of the systems in hand. On the other way, the resulting model could be stored in appropriate data bases to have a collection of behaviors that can be visualized or analyzed through an interface equipped with the run-time support of the process algebra adopted.

Most appreciated results can be achieved by exploiting meta modeling tools to hide the formal details from biologists. We suggest relying on an interface between the tools usually adopted by biologists and

⁵ D.L. CHAO, S. FORREST, *Information Immune Systems*, in «Proceedings of the First International Conference on Artificial Immune Systems (ICARIS)», 2002, pp. 132-140.

the formal kernel of the environment. The graphical language should be formal enough to allow the extraction of process algebra specifications and to visualize back the results of the simulation/analysis carried out on the process terms. Possible languages to be used here are the Kohn's graphical formalism⁶ with its Kitano's extension⁷, the diagrammatic cell language⁸, BioUML (see <<http://www.biouml.org/>>) or SB-UML⁹.

Finally, also high-throughput tools are nowadays largely available to produce in short time a huge amount of data. The problem that biologists must face is then the interpretation and selection of the produced data. An important aspect of any *in silico* lab is surely the connection between the high-throughput tools and those for either data mining in the classical sense, or for inferring behavioral models of the phenomenon at hand. Summing up, we devise an architecture like the one in Figure 1 to completely describe the biological *in silico* lab that should be the arena where looking for new biology and new



computational paradigms.

⁶ K.W. KOHN, *Molecular Interaction Map of the Mammalian Cell Cycle Control and DNA Repair Systems*, in «Molecular Biology of the Cell», X, 1999, pp. 2703-2734.

⁷ H. KITANO, *A Graphical Notation for Biological Networks*, in «BioSilico», I, 2003, pp. 169-176.

⁸ R. MAIMON, S. BROWNING, *Diagrammatic Notation and Computational Grammar for Gene Networks*. Proc. of the Second International Conference on Systems Biology 2001, pp. 311-317.

⁹ M. ROUX-ROUQUIE, N. CARITEY, L. GAUBERT, C. ROSENTHAL-SABROUX, *Using the Unified Modeling Language (UML) to Guide Systemic Description of Biological Processes and Systems*, in «BioSystems», 2004.

Fig. 1. Environment schema

4. *Potential Impact and Perspectives*

Understanding cellular signalling etc. is the logical next step for biological science after the massive effort which is currently being directed towards proteomics. The wealth of information generated on protein-protein interactions is useless unless we can start to understand (and quantify) the pathways they are involved in (i.e. the big picture).

In this respect, we believe that the formal stochastic modeling of the dynamics of biological systems will have a sensible impact over medicine, biology and pharmacology. In the near future it could improve the target discovery for pharmacology industries thus enhancing the efficiency of medicine and reducing their collateral effects. Also, new biological knowledge could be provided by predicting the behaviour of biological systems and therefore driving the selection of real experiments among a plethora of possible ones via simulation and analysis of models. In the long-term, the formal modeling of bio-systems might have benefits in a number of health and societal problems.

- *Predictive, preventive and personalized medicine.* Predictive medicine aims at analyze DNA sequences from small amount of blood to predict a probabilistic health history of individuals (expert says that it will be available within 10-15 years). Preventive medicine will use a system approach (in the next 15-25 years) to place defective genes within a lab biological context and to simulate their behaviour (possibly in silico) to circumvent their limitations. Finally, personalized medicine will examine the gene set of any individual and will prevent diseases on actual data rather than statistical ones.
- *Drug design and gene therapy.* The understanding of the genetic regulatory network in full details will help pharmacology to develop drugs that acts exclusively on the signalling mechanisms that cause the disease thus limiting collateral and undesired effects. Furthermore, the full understanding of gene functioning can help design genetic therapy that removes anomalies from affected individuals. Gene based drugs for common conditions such as diabetes and high blood pressure should be available in 15-20 years.
- *Toxicological research.* This research aims at discovering correlations between responses to toxic agents and gene profile of individuals. The possibility of studying and simulating in silico such interactions

could be a breakthrough in this field.

- *Environmental research.* Due to the scale understanding objective of the project based on the complex self-organizing hierarchy of living organisms, we can push a step further the hierarchy and study population of organisms. This should help understanding and simulating how eco-systems are affected when introducing a new organism or removing an organism. This could be of interest also to medicine when new viruses emerge or when viruses cross species barriers.
- *Building organs and tissues.* The complete understanding of the gene functioning and the complete control of development of organisms could lead to the production of organs and tissues from cells of the individual to which they must be transplanted thus avoiding the rejection problem. Alternatively, the organs of the donor could be genetically modified to prevent rejection when transplanted. Of course the understanding of the dynamics of evolution from genes to organs through cells is a starting point for this and our project will help in this by providing simulation and analysis tool operating on large amount of data that cannot be handled manually.
- *Promote the cross-fertilization of different disciplines.* The tuning of our environment through in vitro experiments can bring benefits to computer science in assessing well-established theories or in changing them according to real needs. Concomitantly, bio-molecular research can receive important inputs and predictions from studies of specific systems, as well as general insights borrowed and adapted from computer science. Furthermore, the cross-fertilization between biology and computer science establishes firm grounds for the further growth of bioinformatics.

We conclude the presentation commenting on the other side of the coin: the perspectives in computer science. Computer and bio-molecular systems both start from a small set of elementary components from which, layer by layer, more complex entities are constructed with evermore sophisticated functions. Computers are networked to perform larger and larger computations; cells form multi-cellular organisms. All existing computers have an essentially similar core design and basic functions, but address a wide range of tasks. Analogously, all cells have a similar core design, yet can survive in radically different environments or fulfill widely differing functions. Of course, bio-molecular systems exist independently of our awareness or understanding of them, whereas computer systems exist because we understand, desi-

gn and build them. Nevertheless, the abstractions, tools and methods used to study complex systems and biological systems can illuminate a paradigm shift in developing computing machines. In fact our accumulated knowledge about bio-molecular systems could allow us to use them as computational devices¹⁰.

ACKNOWLEDGMENTS: The authors have been partially supported by the FIRB project '*Modelli formali per Sistemi Biochimici*'.

CORRADO PRIAMI, PAOLA QUAGLIA

¹⁰ A. REGEV, E. SHAPIRO, *Cells as Computations*, in «Nature», CDXIX, 2002, p. 343.

Un approccio geometrico all'analisi dei testi letterari

1. Introduzione

Modelli di cooperazione testuale sono stati proposti fin dai primi anni Sessanta da Umberto Eco¹ nei quali il testo viene visto come aperto, ossia sottoposto a un'interazione con la comunità dei lettori: il testo diventava in questo modo una sorta di potenzialità di fruizione, che il lettore sfrutta colmando gli interstizi che il testo gli offre. Il lettore diventava un soggetto attivo che stabiliva un legame tra il testo e il senso dei significati possibili nella cultura e nella società in cui agiva.

Il nostro tentativo è quello di calare queste intuizioni all'interno di una prospettiva computazionale: in particolare, di proporre un modello dinamico di significato lessicale del testo utilizzando strumenti non supervisionanti dell'apprendimento automatico del linguaggio (*unsupervised machine learning*). In particolare, il modello che qui proponiamo, verte su un'interpretazione della semantica lessicale dei testi letterari, da un punto di vista geometrico e non logico, ovvero a partire da analisi meramente distribuzionali delle parole nel testo, e non da un'analisi logica delle posizioni e del ruolo che rivestono le parole in una frase rispetto a costrutti predefiniti.

Il nostro intento è quello di presentare un modello capace di far emergere, non solo il flusso narrativo degli eventi e dell'intreccio presenti nel testo, ma di vederli nella prospettiva paradigmatica di un significato comune a una comunità ideale di lettori. Se assumiamo come esempio l'*incipit* del primo romanzo di Alberto Moravia, *Gli Indifferenti*, che sottoporremo in seguito a un'analisi dettagliata con il nostro modello («entrò Carla; aveva indossato un vestitino di lanetta

¹ U. Eco, *Opera Aperta, Forma e indeterminazione delle poetiche contemporanee*, Milano, Bompiani 1962.

marrone con la gonna così corta, che bastò quel movimento di chiudere l'uscio per fargliela salire di un buon palmo sopra le pieghe lente che le facevano le calze intorno alle gambe»)², il lettore è in grado di assumere una correlazione sintagmatica delle parole *gonna*, *corta*, *salire* capaci di evocare l'evento di un atteggiamento disinibito della ragazza che si presenta agli astanti, e una paradigmatica, ponendo in correlazione il termine *gonna* con una serie di termini assenti nel testo, ma presenti nella sua conoscenza del mondo femminile (abbigliamento femminile, attillata, sotto le ginocchia, di jeans, fuori moda, ...).

Il concetto di Campo semantico, definito nel paragrafo successivo, ovvero un insieme di parole strettamente correlate tra loro in rapporto preferibilmente distintivo, contribuisce in maniera significativa a cogliere proprio questa conoscenza, esemplificata nel Campo di *gonna*. Laddove, i contributi alla generazione del Campo semantico possono essere brani di un testo specifico, oppure quelli di molti testi atti a rappresentare una conoscenza sociale diffusa. Nel primo caso si avrà un'accezione del significato di *gonna* molto calato nel sapere del testo campione (supponiamo il romanzo di Moravia), con possibili interpretazioni come quella riportata in precedenza. Nel secondo invece si otterrà un'estensione in cui gli 'interstizi del testo' verranno coperti dalla conoscenza distribuita in molti testi, come modello ideale della conoscenza del lettore. Non solo: in alcuni casi verranno colmati dalla competenza del lettore stesso come soggetto attivo, in particolare *soggetto critico*, che sfrutterà le evidenze del Campo semantico generato dal testo stesso per formulare nuove interpretazioni di quest'ultimo (come abbiamo cercato di mettere in luce nell'Appendice finale con un intervento del critico letterario Roberto Gigliucci).

2. *Il Campo semantico*

La teoria dei Campi semantici considera una rappresentazione semantica come un concetto emergente dalla relazione tra il significato di termini all'interno di una qualche relazione. In altre parole, il Campo semantico è un insieme di parole strettamente correlate tra loro, e ogni parola del Campo viene espressa in termini di differenza rispetto alle altre.

² A. MORAVIA, *Gli Indifferenti*, in ID., *Opere*, 1, Milano, Bompiani 2000.

Introdotti dal linguista olandese Jost Trier negli anni Trenta, i Campi semantici, secondo l'autore dovevano far emergere l'idea che non esisteva nessuna parola isolata nella conoscenza del parlante, ma ognuna faceva parte di una struttura, molto prima di far intervenire una grammatica o comunque qualche logica aggregativa basata sul discorso.

«I Campi sono le realtà linguistiche vive fra le parole singole e il lessico nel suo insieme; come sottoinsiemi condividono con la parola la caratteristica di articolarsi in unità superiori, e con il lessico la caratteristica di articolarsi in unità inferiori»³.

La teoria dei Campi non necessita quindi di una definizione formale delle relazioni tra parole, ma si accontenta di individuare un principio di contrasto, in grado di rispondere ai requisiti predetti. Per esempio il Campo semantico dei colori può essere rappresentato grazie alle proprietà di brillantezza e percentuale di miscelazione di tre colori primari: Rosso, Blu, Verde. Inoltre il Campo semantico permette di disambiguare termini polisemici, grazie alla capacità di definirne il contesto d'uso: per esempio: il *calcio*, in ambito chimico è un elemento, mentre in ambito sportivo un *gioco*. Le proprietà detenute dalla parola calcio, in chimica, sono diverse da quelle del gioco del calcio.

Il Campo semantico di un termine abbraccia sia le definizioni di quel termine sia le sue caratteristiche concettuali (*conceptual feature*)⁴. Nelle scienze cognitive sono stati proposti differenti approcci che tendono a dare una valenza concreta a queste *feature*, in modo da calarle nelle azioni e nelle percezioni di un ambito comunicazionale, sviluppato in un determinato contesto.

2.1. Un'interpretazione geometrica del Campo semantico

Scopo della nostra ricerca è quello di cercare un'interpretazione geometrica del Campo semantico alla luce delle idee esposte, basandoci su metodologie di riduzione di Spazi vettoriali (LSA: Latent Semantic Analysis), rispetto a una componente dello spazio capace di mettere

³ J. TRIER, *Das sprachliche Feld. Eine Auseinandersetzung*. Neue Fachbücher für Wissenschaft und Jugendbildung, 10, ***, *** ***. (trad. it. *Il Campo semantico. Una discussione*. Conferenza abbreviata in [Cuppari, 2000], appendice I.

⁴ D.P. VINSON, G. VIGLIOCCO, S. CAPPA, S. SIRI, *The Breakdown of Semantic Knowledge: Insights from a Statistical Model of Meaning Representation*, London, UCL Department of Psychology ed., May 2003.

in evidenza la correlazione tra i termini dei testi originali. A partire da LSA, si è quindi costruita una definizione di Campo semantico statico e una di Campo semantico dinamico (al variare dell'evoluzione di un testo narrativo).

Riprendendo un nostro precedente intervento: «Il concetto di Campo semantico nasce dall'osservazione che i Campi semantici e i giochi linguistici sono naturalmente riflessi nel testo, e come tali indagabili in una prospettiva empirico-geometrica. In altre parole, sulla base di un processo – secondo molti – versi simile a quello descritto con LSA, è possibile inferire l'esistenza di domini semantici servendosi delle analisi statistiche delle distribuzioni dei termini in uno spazio geometrico, all'interno del quale le co-occorrenze dei termini nei testi sono riflesse. Mediante considerazioni di tipo geometrico si è potuto osservare come nel testo si manifesti un'aggregazione del lessico in domini semantici, proprietà altresì nota come *coerenza-lessicale*, e come tale proprietà evidenzii delle relazioni latenti difficilmente esplicabili con i metodi delle logiche formali»⁵.

Ad esempio, nel romanzo *Gli Indifferenti*, la nozione di *noia* è centrale nell'analisi dei personaggi e della condizione psicologica che manifestano nei discorsi diretti e nel commentario del narratore. Nonostante la parola *noia* occorra soltanto 18 volte nel romanzo, e occupi soltanto la seicentesima posizione nella classifica delle frequenze dei termini, è tra quelle, insieme a *indifferenza*, che più caratterizzano le tendenze espressive dell'autore, e le sue intuizioni rispetto alla caratterizzazione della società borghese che descrive. Se analizziamo il Campo semantico di *noia*⁶, questo è composto da termini, quali *esistenza*, *avventura*, *falsità* difficilmente catturabili mediante un'analisi collocazionale (*collocational analysis*) e utilizzando rappresentazioni

⁵ R. BASILI, A.M. GHIOZZO, P. MAROCCO, *Teorie geometriche del significato*, <<http://www.rescogitans.it/main.php?articleid=228>>.

⁶ Definire il Campo semantico di un fenomeno di carattere psicologico, come quello di *noia* all'interno di un contesto narrativo è una sorta di variante rispetto alla nozione classica di Campo semantico, in particolare perché *noia* può essere dotata di diverse connotazioni in funzione dei contesti, quindi è un fenomeno abbastanza instabile rispetto agli esempi precedenti, come il Campo semantico dei colori o del gioco del calcio. L'utilizzo di uno strumento geometrico come LSA, come vedremo, contribuisce comunque a dare una definizione di carattere quantitativo, capace di generalizzare il Campo semantico anche a fenomeni 'più sfumati' come in questo caso.

strutturate 'dall'esterno', capaci di mettere in rilievo quei concetti che un supervisore ha tautologicamente già inserito nel modello.

3. LSA (*Latent Semantic Analysis*)

La Latent Semantic Analysis (LSA) è stata proposta verso gli inizi degli anni Novanta come un modello geometrico per la rappresentazione della semantica lessicale dei testi. L'assunto della teoria è che il significato possa essere catturato a partire dalla distribuzione delle parole nei testi, e dalle mutue correlazioni delle parole stesse in testi, o porzioni di testo, diversi, facenti parte del medesimo *corpus* (es. nella frase: «entrò Carla; aveva indossato un vestitino di lanetta marrone con la gonna così corta», se il termine *gonna* in un'altra porzione del testo è associato a *provocare*, significa che la terna: *gonna*, *corta* e *provocare* costituirà una marca evidente di un concetto, e si potranno considerare i termini *corta* e *provocare* come associabili a uno stesso concetto semantico, sebbene non siano esplicitati in nessuna frase). Quest'idea è sostenuta formalmente dall'utilizzo di una nozione matematica (Principal Component Analysis), applicata a un sistema di riduzione dimensionale dello spazio originale, grazie al quale, nello spazio ridotto, le rilevanze semantiche sono più facilmente rintracciabili.

3.1. Breve digressione su LSA

Nei modelli tradizionali di Information Retrieval (VSM: Vector Space Model), uno spazio geometrico è definito in maniera tale che ogni dimensione dello spazio sia una parola occorrente in una data collezione di documenti (*corpus*). Ogni documento viene rappresentato come un vettore nello spazio d'origine, le cui componenti rappresentano l'insieme dei termini del *corpus*, con valore 0_{ij} se i termine *i* non è presente nel documento *j*, e 1_{ij} se invece è presente (al posto di 1 è possibile inserire un peso *w*_{ij} che stabilisce l'importanza di quel termine in quel documento).

Il maggior limite di questo approccio distribuzionale è quello di volere che i termini occupino dimensioni indipendenti e ortogonali dello spazio dei documenti. Molte relazioni significative tra termini, come quella dell'esempio descritto in precedenza, sono ignorate. Inoltre il VSM è caratterizzato da un grande numero di dimensioni, grande quanto sono i termini nel *corpus*.

L'approccio LSA invece cattura le relazioni statistiche termine-termine e tende a raggruppare insieme termini che esprimono informazioni simili, al contempo lo spazio originale è sostituito con uno molto più piccolo, grazie a un processo di decomposizione, chiamato SVD (Singular Value Decomposition). Data una matrice originale *termine-per-documenti* M che descrive lo spazio originale, questa viene trasformata in tre nuove matrici T , S , D in maniera tale che valga l'uguaglianza $M = TSD^T$. Il prodotto sulla destra cattura la stessa informazione statistica di M in uno spazio k -dimensionale (con K molto più piccola di N : dim. originale dello spazio), dove ogni dimensione rappresenta una delle *feature* LSA (un concetto puramente geometrico ed implicito nello spazio, non corrispondente direttamente a un'evidenza linguistica). D è la matrice dei documenti, e ogni sua colonna rappresenta uno dei K -concetti derivati, T è la matrice dei termini le cui colonne sono i K -concetti derivati e infine S è la matrice dei valori singolari, che esprime in maniera decrescente il significato delle componenti principali di LSA. La dimensione di S stabilisce il numero di componenti significative nelle quali operare.

In definitiva, i benefici di LSA nell'ambito dell'NLP sono i seguenti:

- La somiglianza tra vettori (sia quelli delle parole che quelli dei documenti), nello spazio ridotto, risponde meglio a criteri semantici rispetto a quella dello spazio originale, senza far uso di strumenti complessi di misura, ma rimanendo nelle scelte tradizionali di metriche.
- Il principio di dualità che vale nello spazio ridotto permette di fare un confronto tra una parola e un documento. Ovvero, le proprietà testuali e lessicali possono essere rintracciate in maniera congiunta: è possibile misurare la pertinenza, nel senso della sua semantica latente, di una parola rispetto a un testo, o a una porzione di testo.
- LSA non richiede particolari vincoli sul tipo di processing richiesto, se non quello dell'algoritmo di base utilizzato (SVD). Questo significa che LSA scopre le correlazioni semantiche in una maniera completamente automatica.

3.2. Il paradigma semantico di LSA

LSA rappresenta un paradigma alternativo agli approcci presenti in teorie del significato di tipo logico e metalinguistico (es.: strutture predicative in linguistica generativa o formalismi logici per rappresentazioni ontologiche). LSA spinge verso un approccio alla semantica dei

testi in una prospettiva geometrica e analitica sfruttando le metodologie dei *Vector Space Model*, fortemente utilizzati in Information Retrieval. I concetti emergono dai testi, come conseguenze di metriche di similarità, che quantificano relazioni immerse nel testo stesso, e sono adottabili sia per il confronto tra parole, tra testi, e tra parole e testi.

La possibilità di avere a disposizione una misura di relazione tra parole, crea un significativo ponte tra l'Information Retrieval e certi studi della linguistica del Novecento, come quello sui Campi semantici. Poiché il principio fondatore di un Campo semantico è in pratica quello di vedere ogni parola del Campo in termini di differenza/somiglianza semantica rispetto alle altre, all'interno di un criterio di mantenimento della coerenza lessicale, una misura di relazione tra parole quantifica proprio questa esigenza e la proietta naturalmente in una disciplina formale, come l'analisi vettoriale.

L'idea che qui proponiamo è quindi di esplorare tecniche di acquisizione dei Campi semantici, basandoci su strumenti consolidati dell'Information Retrieval, come supporto all'interpretazione critica dei testi narrativi. Poiché le tecniche di LSA possono essere adottate su *corpus* differenti, con la possibilità di farli interagire per mettere in rilievo maggiori evidenze del testo, l'idea è quella di agire su diversi livelli:

- Un *livello enciclopedico* per catturare le conoscenze testuali sul senso comune dei termini a livello generale di una comunità culturale.
- Un *livello interno all'Opera letteraria*, come rappresentazione del mero significato emergente dall'opera, senza ausilio di conoscenze esterne.

L'integrazione dei due spazi, quello del sapere sociale e quello del testo, permette un confronto tra il significato del Campo semantico centrato su una parola (es. *indifferenza* nel romanzo *Gli Indifferenti*) all'interno della competenza del testo e a livello della competenza sociale. Questa organizzazione del dominio in oggetto ci permette anche di rileggere i Campi semantici rispetto al carattere sintagmatico e paradigmatico del linguaggio: il primo livello viene raggiunto dalle sole parole interne al testo dell'opera, che contribuiscono alla formazione dei diversi Campi semantici; il secondo livello invece è costituito dalle parole che non sono presenti nel testo dell'opera, ma sono potenzialmente sostituibili a quelle presenti. Quest'ultimo potrebbe essere visto come una sorta di meta-campo rispetto al precedente, un

involucro in cui i concetti del Campo semantico dell'opera assumono un significato se contrapposti a quelli del livello superiore.

Alla luce di questo approccio verranno creati due *corpus*, uno interno all'opera e uno esteso (a contributi critici esterni), con due spazi LSA correlati LSA_C e LSA_E . Nel primo emergeranno i fenomeni relativi al testo, nel secondo quelli frammisti anche a un sapere più diffuso, atto a esplorare intrinsecamente e automaticamente le possibilità critiche.

Ad esempio il lessico del Campo semantico del termine indifferenza. $L_{\text{indifferenza}} = (\text{vita, noia, scena, prova, vero, volontà, esistenza, proposito, ambiente, mancanza, incapacità, vanità})$ non contiene due termini che nel *corpus* esteso a testi critici sul romanzo invece emergono: *borghesia* e *fascismo*. Questi ultimi diventano particolarmente significativi se ricondotti ai termini appartenenti al lessico primario $L_{\text{indifferenza}}$ e possono restituire il contesto storico e il sostrato familiare in cui la proprietà dell'indifferenza è presente nel romanzo.

3.2.1. *Catturare una conoscenza sociale come Campo semantico (basato su LSA)*

In uno spazio LSA la proprietà di prossimità semantica tra termini può essere vista come una stima quantificabile dell'associatività tra concetti. In altre parole, il Campo semantico diventa quell'insieme di termini che, nella loro rappresentazione vettoriale, sono 'vicini' rispetto a una misurazione dello spazio.

Prendiamo un esempio canonico nella descrizione dei Campi semantici: l'insieme dei colori⁷. Se consideriamo un testo saggistico sulla teoria dei colori, il Campo semantico che ne emerge sarà quello formato dai nomi dei colori stessi: 'giallo, blu, arancione, ...'. Questi ultimi sono semanticamente omogenei rispetto al significato di colore, e in contrapposizione tra loro rispetto al concetto di colore (più sono i colori che emergono nel Campo semantico del testo, meglio emergerà che cosa rappresentano i colori per la popolazione descritta nel testo). Un sistema di misurazione nello spazio LSA restituisce immediatamente la prima di queste proprietà: l'omogeneità semantica, e la possibilità di collezionare termini che insieme possano essere esaustivi rispetto alla modellizzazione di un concetto. Per la seconda proprietà, il discorso è un po' più articolato.

⁷ U. Eco, *Trattato di semiotica generale*, Milano, Bompiani 1975, p. 124.

Quest'ultima, ossia la caratteristica di mutua indipendenza dei termini del Campo semantico rispetto al concetto che devono evocare, è quella che consente di denotare il concetto generale di colore (questo avviene perché, se utilizzassimo in prevalenza sinonimi di un colore particolare, ad esempio il bianco, o piccole gradazioni cromatiche dello stesso, non riusciremmo mai a descrivere l'insieme dei colori). LSA, utilizzata nella maniera standard, non esclude i sinonimi, anzi li raggruppa: per cui, nel nostro esempio dei colori, troveremmo anche il bianco di zinco, il bianco di titanio, ecc. che differiscono di pochissimo rispetto al concetto di bianco. Quel che ci aspettiamo però è di trovare il bianco di zinco e di titanio particolarmente vicini agli altri tipi di bianco, e in generale a espressioni lessicali che evocano il bianco, in maniera tale che tutti questi termini possano rientrare in un unico cluster, quello del bianco. Mentre nel caso del rosso, quest'ultimo dovrebbe trovarsi a una distanza maggiore, poiché, in un testo saggistico ideale sui colori, ci aspettiamo di trovare un insieme di correlazioni tra quest'ultimi che privilegerà quelle tra gli elementi dotati di proprietà comuni. E, nel nostro testo saggistico ideale di teoria dei colori, i documenti che parleranno di rosso e bianco, o comunque che veicoleranno le correlazioni tra rosso e bianco, saranno probabilmente di meno rispetto ad altri che metteranno in rilievo fenomeni percettivi e pittorici più rilevanti⁸.

In definitiva LSA, da un punto di vista ideale (quindi osservata su testi omogenei rispetto al senso ideale racchiuso in un particolare Campo semantico), dovrebbe restituire un insieme omogeneo di fenomeni, tali che quelli di carattere sinonimico possano essere ricondotti a un rappresentante unico – a questo proposito si veda anche la discussione nel paragrafo successivo sul rapporto tra nome e definizione – mentre quelli in contrasto tra loro, ma necessari per veicolare il concetto generale del Campo semantico, rimangano indipendenti (sorta di primitive) del Campo semantico stesso⁹.

⁸ Ad esempio molti documenti metteranno in evidenza le relazioni tra i colori fondamentali: rosso, verde e blu, altri documenti tratteranno dei colori che maggiormente vengono utilizzati in una sintesi additiva o attrattiva, i colori complementari ai fondamentali, e così via.

⁹ Questo principio può emergere particolarmente se vengono adottate tecniche di riduzione geometrica più avanzate rispetto a LSA, le quali guardano non solo alla distribuzione dei termini nel testo, ma anche a una semantica associata a un grafo

3.2.2. I termini singolari

La geometrizzazione del significato rappresenta un valido strumento per un approccio alla semantica dei nomi propri¹⁰. Come osserva Piero Perconti «oggi il campo delle tesi in competizione relativamente alla semantica dei nomi propri è conteso tra le teorie descrittiviste, secondo le quali il significato dei nomi propri consisterebbe in una o più descrizioni che possono essere sostituite con il nome corrispondente senza mutare il valore di verità delle frasi in cui esse compaiono, e le teorie causali del riferimento, secondo cui i nomi propri non avrebbero alcun senso (fregeano) e si riferirebbero direttamente ai loro individui tramite una opportuna catena causale iniziata con un battesimo»¹¹. Innanzitutto un nome proprio inserito in Campo semantico ricopre una funzione equivalente a quella di un qualunque altro termine contenuto nel Campo: da questo punto di vista, una definizione del nome proprio dovrebbe essere geometricamente ‘vicina’ al nome stesso (in termini più propri i due vettori, quello del nome proprio e quello della sua definizione, dovrebbero avere una misura simile), ammesso che questa definizione sia effettuata utilizzando i termini del dizionario generato dal *corpus* in esame. Più strettamente correlate sono il nome e la sua definizione, più quest’ultima ha una maggiore ‘vicinanza’ nello spazio geometrico in esame. Le coppie *lemma-definizione* rimandano a un aspetto enciclopedico che mette nella stessa posizione i termini generali da quelli singolari, nel loro carattere lemmatico, ammesso che il *corpus* sia sufficientemente grande e ben distribuito, in modo tale da poter fornire definizioni appropriate dei lemmi. Da un punto di vista della quantificazione del significato, quel che ci si aspetta da metodologie geometriche come quelle esposte, è di avere cluster *lemma-termini della definizione* particolarmente individuabili all’inter-

esterno, come LPP (*Locality preserving projection*). Con queste tecniche, i fenomeni caratterizzati da relazioni di sinonimia possono più facilmente essere raggruppati in cluster, in maniera tale da rimanere ancora più fedeli all’accezione di Campo semantico, secondo la tradizione linguistica storica, che ne connota i rappresentanti come indipendenti.

¹⁰ Più in generale i termini singolari, ossia i nomi propri, i deittici e le descrizioni definite: tutto ciò che viene usato per fare riferimento a un unico oggetto o gruppo di oggetti.

¹¹ P. PERCONTI, *X-phil. Una ‘nuova’ tendenza in filosofia?*, <<http://www.rescogitans.it>>.

no del Campo semantico del concetto a cui è ricondotto il lemma¹².

In definitiva, la competenza semantica rimane 'chiusa' all'interno del *corpus* utilizzato nell'esperimento, arricchibile certo con altri *corpus* affini o meta-semantici rispetto al *corpus* di riferimento (per esempio se il *corpus* è l'opera di un autore, un *corpus* significativo può essere quello dell'insieme dei testi critici su quell'autore), ma senza la necessità di ricondurla a una semantica esterna.

3.2.3. *Il Campo semantico dinamico*

La nozione di Campo semantico in LSA che abbiamo appena messo in rilievo, è in un certo senso, staticamente applicata a un *corpus* precostituito, che può essere quello della conoscenza sociale (idealmente: un insieme di testi ed enciclopedie in grado di riprodurre una cultura), oppure quello di un testo specifico (il saggio sui colori, un romanzo di Moravia).

Questa nozione può però essere studiata dinamicamente per riprodurre le variazioni di un Campo semantico, attraverso un insieme di parti di un *corpus*, caratterizzate da una sequenza temporale o comunque da un'evoluzione narrativa. Ognuna di queste parti può essere misurata rispetto alle parole che meglio le evidenziano.

Si noti come ogni unità testuale di un'opera possa rappresentare una sorta di pseudo-documento descritto come un vettore (un punto) nello spazio LSA, e come questo vettore possa essere confrontato con parole e concetti, all'interno dello stesso spazio, grazie alle proprietà duali precedentemente citate. In particolare un'unità testuale (ad esempio un capitolo o un'altra porzione di testo di un romanzo) può essere confrontata con il Campo semantico che meglio lo rappresenta. In questo modo è possibile osservare come, in certe parti di un'opera narrativa, vengano privilegiati certi concetti rispetto ad altre parti, e come questi concetti si evolvano nell'opera. In altri termini, come il Campo semantico associato a un determinato contesto espressivo possa variare, al variare delle vicende descritte. L'unione di questi contesti espressivi potrà restituire il Campo semantico globale di quei concetti (ritornando all'esempio precedente di *noia* nel romanzo *Gli Indifferenti*, il Campo semantico globale sarà costituito da tutti i carat-

¹² Ogni Campo semantico sufficientemente esteso contiene al proprio interno sottocampi semantici, per cui l'insieme *lemma-termini della definizione* potrebbe essere visto a sua volta come un Campo semantico.

teri distintivi che via via nei diversi capitoli evocheranno il senso di noia nel romanzo).

Una visione dinamica dei fenomeni permette altresì una rappresentazione grafico-funzionale più interessante, in quanto la rilevanza tra capitoli e concetti può essere vista come un punto che si evolve nel tempo, ossia una curva. Il Campo semantico, in questa accezione dinamica, è quindi la combinazione delle curve dei concetti che lo compongono, che può essere a sua volta visto come un'unica curva risultante. Quest'ultima ha ovviamente un carattere sintetico, ma è in grado di fornire una rappresentazione visuale semplice ed efficace dell'evoluzione del fenomeno.

4. Metodologie applicative

Il campo d'indagine è stato quello del primo romanzo di Alberto Moravia, *Gli Indifferenti* (1929). Il testo è costituito da 16 capitoli, 91059 parole (token) e 3273 lemmi indipendenti.

L'unità testuale presa in esame è stata il paragrafo (dal capoverso al punto a capo): sono stati esaminati 1854 paragrafi, con una media di 116 paragrafi per capitolo.

La fase sperimentale è stata condotta sia su un *corpus* formato dal testo originale LSA_C , sia da una versione arricchita con un'opera critica LSA_E (13041 token, 1920 parole differenti) che ha permesso di estendere i Campi semantici a termini del testo narrativo insieme ad altri del testo critico.

Sono state utilizzate diverse pesature, per l'assegnamento dei valori iniziali della matrice termini-documenti, a cui è stata applicata LSA. La dimensione utilizzata nello spazio finale è stata 100 (componenti principali). Infine sono stati adottati anche filtri grammaticali per distinguere Verbi, Nomi, Aggettivi e Avverbi.

4.1. Estrazione del significato dai Campi semantici

La generazione di Campo semantico è stata ottenuta attraverso una misura della distanza negli spazi LSA_C e LSA_E (si veda formula 1 in Appendice II.6) Nelle diverse prove è stata adottata una soglia variabile di appartenenza al Campo (da 0.4 a 0.6 unità, misurate con la metodologia dea Cosine Similarità, si veda formula 2 in Appendice II.6). Presentiamo di seguito alcuni esempi nel *corpus* LSA_E :

- $c=noia$, $L_{noia} = (esistenza, atto, avventura, fatalità, falsità, familiare, ragazza, abitudini, ...)$
- $c=indifferenza$, $L_{indifferenza} = (vita, noia, scena, prova, vero, volontà, esistenza, proposito, ambiente, mancanza, incapacità, vanità, borghesia, fascismo)$
- $c=Carla$, $L_{Carla} = (osservare, baciaronno, torpore)$
- $c=Leo$, $L_{Leo} = (suonare, ingiunse, stirò, cammina, fastidio, signor)$

È interessante verificare come un considerevole numero di concetti semantici sia catturato in una modalità completamente non supervisionata, e senza utilizzare nessuna analisi sintattica o riduzione delle parole alla loro radice (lemma) (per cui le parole sono evidenziate nella forma in cui appaiono nel discorso). Innanzitutto nei diversi Campi emergono i caratteri distintivi (es.: *noia/volontà, falsità/familiare*, e altri). In secondo luogo, le correlazioni al livello di plot, come la coppia *noia/falsità*, costituiscono delle proprietà di un fenomeno che Moravia descrive, quello di una famiglia borghese in fase decadente, che conserva abitudini del passato, le quali appaiono oltremodo false e desuete nei nuovi contesti ai quali cerca di adattarsi.

Mentre il lessico dei primi due Campi semantici L_{noia} e $L_{indifferenza}$ restituisce una significativa esperienza del lettore, quello associato ai due personaggi protagonisti, Carla e Leo non ne cattura propriamente il significato. Questo accade per diversi motivi: i nomi propri, in un'opera narrativa (a differenza di un articolo di giornale) sono spesso affetti da referenze anaforiche, e nel nostro semplice modello non abbiamo adottato nessuna analisi anaforica per riconoscere ad esempio i pronomi correlati ai soggetti in gioco. Inoltre i personaggi del racconto sono molto pochi, per cui Carla e Leo si trovano ad agire pressoché in tutte le situazioni descritte nel romanzo, introducendo così molto rumore e molta variabilità di significati.

4.2. Sviluppo dinamico di Campi semantici

In questo paragrafo viene evidenziato l'output del tool utilizzato dalla nostra analisi, che implementa il modello matematico delle funzioni di rilevanza dei concetti per unità testuale e per capitolo, espresse nei precedenti paragrafi.

In generale osserviamo che il comportamento delle diverse parole di un Campo semantico tende a essere omogeneo nel testo: ossia, in termini geometrici, le curve correlate alle parole hanno un andamento sincrono nell'evoluzione del testo (crescono o decrescono insie-

me), come risulta evidenziato dall'esempio in figura 1, dove vengono mostrate nel grafico in alto le curve delle singole parole del Campo, e in quello in basso, la sintesi geometrica del Campo stesso (si veda formula 4 in Appendice II.6).

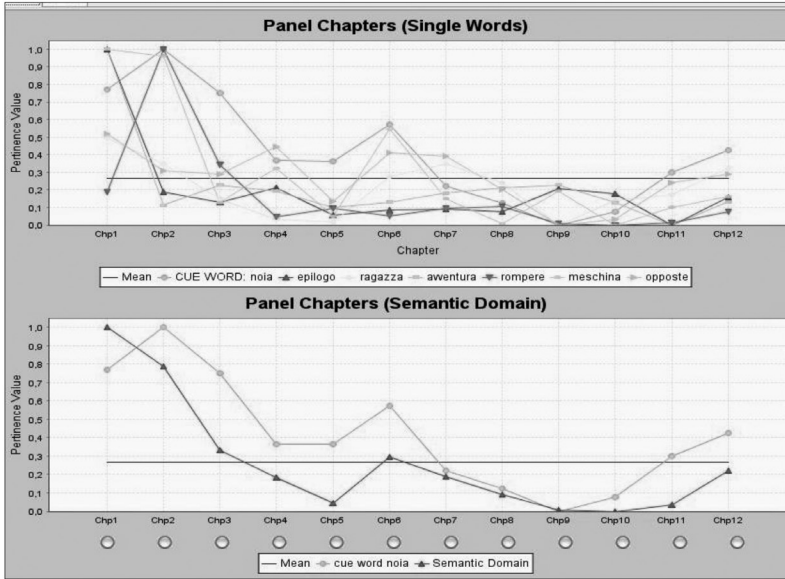


Fig. 1. Evoluzione del Campo semantico di *Noia* in *Gli Indifferenti* (primi 12 capp.)

In secondo luogo, alcuni capitoli mettono in particolare evidenza l'emergere della parola *noia*, con un Campo semantico particolarmente valorizzato dalla ricorrenza delle parole del Campo, e dai concetti correlati. Analisi più accurate delle curve, come quella della deviazione standard dei paragrafi in un capitolo, possono mettere in evidenza gli eventi e le frasi in cui il fenomeno è particolarmente rilevante, e le correlazioni rispetto ai brani in cui è inerte (vedi figura 2 che rappresenta la distribuzione dei concetti predetti nelle unità testuali di un capitolo scelto. Infine la figura 3 visualizza il testo del romanzo relativo a un'unità testuale, per dettagli si veda formula 3 in Appendice II.6).

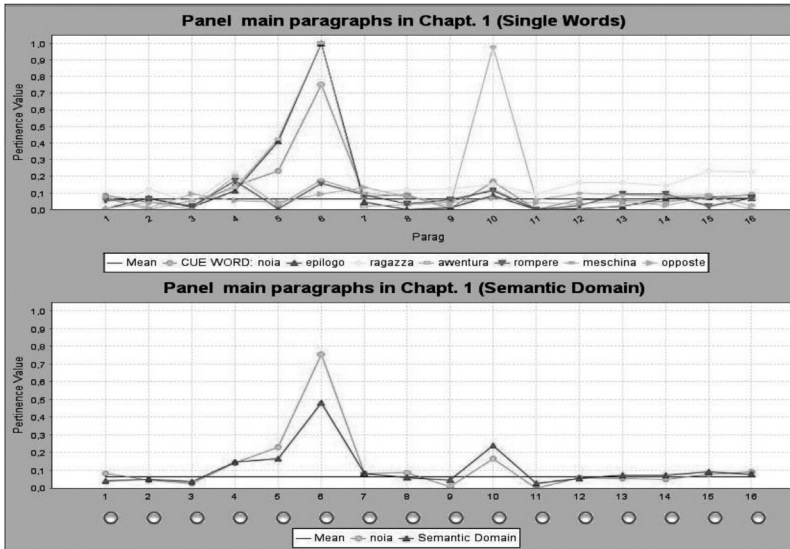


Fig. 2. Evoluzione del Campo semantico di *Noia* in *Gli Indifferenti*
(paragrafi del capitolo 1)

Indirizzo C:\UnDataSperimentali\programmi\TOOL LSA GRAPHICS 4\saved pages\Paragraph 6 (Chapter 1).html

Google

Cerca

Segnalibri

961 blocchi

Controllo

Invia a

Impostazioni

Text Paragraph 6 (Chapter 1)

length Paragraph (number of chrs): 1534
Number of words of the original SD : 3
Number of words of the final (user chosen) SD : 3

Ella fece di nuovo il vano gesto di respingerlo, ma ancor piu' fiaccamente di prima, che' ora la vinceva una specie di volonta' rassegnata; perche' rifiutare Leo? Questa virtu' l'avrebbe rigettata in braccio alla noia e al meschino disgusto delle abitudini; e le pareva inoltre, per un gusto fatalistico di simmetrie morali, che questa avventura quasi familiare fosse il solo epilogo che la sua vita meritasse; dopo, tutto sarebbe stato nuovo; la vita e lei stessa; guardava quella faccia dell'uomo, la', tesa verso la sua: "Finirla", pensava "rovinare tutto " e le girava la testa come a chi si prepara a gettarsi a capofitto nel vuoto, ma invece supplico': "Lasciami", e tento' di nuovo di svincolarsi; pensava vagamente prima di respingere Leo e poi di cederli, non sapeva perche', forse per avere il tempo di considerare tutto il rischio che affrontava, forse per un resto di civetteria; si dibatte' invano; la sua voce sommessa, ansiosa e sfiduciata ripeteva in fretta la preghiera inutile: "Restiamo buoni amici Leo, vuoi? Buoni amici come prima" ma la veste tirata le scopriva le gambe, e c'era in tutto il suo atteggiamento renitente e in quei gesti che faceva per coprirsi e per difendersi, e in quelle voci che le strappavano le strette libertine dell'uomo, una vergogna, un rossore, un disonore che nessuna liberazione avrebbe potuto piu' abolire.

Fig. 3. L'unità testuale emergente dal picco del par. 6 (Semantic Domain = 0.5)

5. Appendice I. Il punto di vista di un critico letterario

L'applicazione di LSA a un testo letterario prospetta linee di indagine quanto mai galvanizzanti per il ricercatore umanista. Mentre l'applicazione a un testo filosofico (per quanto ospitante sonetti e formalmente molto creativo) come gli *Eroici furori* di Giordano Bruno si è rivelata notevolmente proficua sul piano del contributo allo studio del lessico filosofico dell'autore e degli intrecci concettuali¹³, la stessa applicazione a un romanzo come *Gli Indifferenti* di Moravia è un primo passo verso l'allestimento di una strumentazione straordinariamente dinamica per il critico letterario.

Si tratta di compiere un salto epistemologico nel campo della collaborazione fra scienze informatiche e scienze letterarie. Ovvero passare da un tipo di prodotti sostanzialmente assimilabili a un *corpus* testuale interrogabile da un motore di ricerca piuttosto evoluto, come il DBT (Data Base Testuale dell'enciclopedia LIZ)¹⁴, a un modello di strumento che ricorra a tecniche computazionali per estrarre dai testi letterari non solo informazioni linguistiche ma anche concettuali¹⁵.

L'utilizzo di LSA comporta un varco epistemologico, che gli umanisti potrebbero aver difficoltà a superare. Infatti la natura geometrica, topologica dell'inquisizione di LSA può far sorgere obiezioni di inintelligenza assoluta del testo da parte della macchina e quindi di prevaricazione del superficiale rispetto al profondo. Su questa proble-

¹³ S. BASSI, F. DELL'ORLETTA, D. ESPOSITO, A. LENCI, *Computational Linguistics Meets Philology: a Latent Semantic Analysis of Giordano Bruno's texts*, in *LREC 2006: 5th International Conference of Language Resources and Evaluation*, Genova 2006.

¹⁴ Letteratura Italiana Zanichelli (LIZ), a cura di P. Stoppelli ed E. Picchi, Bologna, Zanichelli 2004.

¹⁵ Si sono fatti passi importanti in questa direzione, ad esempio progettando un tool che analizzi, con l'apprendimento progressivo, l'entità dello scarto di una traduzione poetica d'autore rispetto al testo in lingua originale, per raggiungere poi in un secondo momento la possibilità di definire su questa base lo *stile* dell'autore che traduce (Gigliucci e Marocco in *Cyberletteratura: tra mondi testuali e mondi virtuali*, a cura di A. Di Stefano, Roma 2006); si è proceduto nella analisi delle descrizioni negli *Indifferenti* da parte della macchina istruita per l'identificazione di questa particolare ontologia degli elementi narrativi (Basili-Moschitti-Pennacchiotti e Di Stefano, in *Cyberletteratura. Atti del Convegno [Roma 9-10 giugno 2005]*, Roma, Centro Stampa Nuova Cultura 2006).

matica, che coinvolge il confronto anche speculativo fra modelli logici e modelli geometrici, mi permetto di dire, da umanista letterato, che la difficoltà sul piano metodologico per noi non esiste: il livello delle occorrenze e delle co-occorrenze di lemmi in un corpo testuale è dato rilevante per un analista letterario comunque, perché non si tratta soltanto di escutere dati consapevoli, ordinati dall'autore con volontà razionale, ma anche dati inconsapevoli o appena subliminari, o comunque di coerenza interna all'opera d'arte nella sua autonomia. E questo discorso non vale certo soltanto per una critica psicanalitica o strutturalista, ma per ogni critica letteraria. Inoltre, dato secondario ma non indifferente, LSA permette di superare il problema di alterità linguistiche dei testi (testi antichi o in lingua latina o straniera), non avendo bisogno di un effettivo *thesaurus* e scommettendo solo sul principio che il particolare lavoro topico si traduce in significativo.

Ecco che allora LSA permette una intersezione formidabile con lo studio dei testi creativi dal punto di vista delle loro dinamiche interne non sempre ordinate in coscienza retorica o espressiva ma produttive di semantica appunto latente. A prescindere, del resto, dal fatto che sia l'autore consapevolmente a dispiegarle o che siano disseminate da una forza interna di tenuta microcosmica dell'opera letteraria. In ogni caso, se pensiamo alla critica tematica francese storica (Richard, Weber, Mauron, Bachelard ecc.) con l'intuizione di studiare non solo reti motiviche consce (*topoi*, stereotipi ecc.) ma anche inconsapevoli, e agli esiti tematologici fino ai nostri giorni, riconosciamo che la critica letteraria può aver bisogno di uno strumento come LSA, integrandolo ai propri processi ermeneutici.

Io credo che un buon uso di LSA da parte dell'operatore umanistico sia quello di accedervi per verificare un'ipotesi o più ipotesi già formulate. Cioè indirizzare la ricerca, che altrimenti sarebbe randomizzata, seguendo una intuizione personale, conseguente alla lettura e analisi del testo. La partenza può essere anche soltanto l'escussione di un'immagine o di una parola, che si ritiene riassuntiva di un discorso complesso e centrale, portante, su cui si possa fondare l'intero edificio interpretativo. Questo è uno dei modi possibili di iniziare, non l'unico.

Ritornando al testo *Gli Indifferenti*, l'idea di partenza è che la parola *gonna* abbia una capacità di riverberazione semantico-immaginale e una entità sintetica nucleare, posta peraltro ad apertura (celebre) del romanzo: «Entrò Carla; aveva indossato un vestitino di lanetta marrone con la gonna così corta, che bastò quel movimento di chiudere l'uscio per fargliela salire di un buon palmo sopra le pieghe lente che

le facevano le calze intorno alle gambe»¹⁶. Il punto è allora partire da *gonna* per costruire poi nuove famiglie e reti concettuali-immaginali possibili. Daremo soltanto pochi cenni di un *work in progress* in tal senso.

Con la parola *gonna* (centroide del Campo semantico relativo) non si sono avuti risultati in prima istanza molto apprezzabili, ma incrociando famiglie di lemmi siamo giunti ad es. alla famiglia generata da *divano*. Nel Campo di *divano*, oltre a *boudoir* troviamo *cosce* e *gonna* in un episodio importante, l'incontro di Michele e Lisa nel capitolo 5. Qui, a un certo punto dell'incontro ripugnante e ridicolo fra i due, capita che la gonna salga sul ventre di Lisa, e questo è senz'altro da ricollegare all'inizio del romanzo, sempre con *gonna* come parola-nucleo, che unisce quindi Lisa a Carla.

Un fruscio; il braccio di Michele scivolò dietro la schiena della donna e le circondò la vita. "No" ella disse con voce chiara, senza muoversi o voltarsi, come se avesse risposto ad una domanda interiore; Michele sorrise di malavoglia, ch  sentiva un certo turbamento invaderlo, e l'attir  pi  strettamente; [...] agitava le braccia, le gambe; infine cadde fuori del divano e in un movimento convulso di tutto il corpo la veste le risal  sul ventre e le grosse coscie di una bianchezza adombrata di muscoli apparvero; allora Michele lasci  la stretta; Lisa risedette e riabbass  la gonna sopra le ginocchia¹⁷.

La *gonna* nella sua centralit , quindi, nella sua rappresentativit  del femminile plurale ma coerente, e soprattutto la *gonna* che risale su per le cosce e il ventre. La *gonna* come ineluttabilit  del femminile ma anche come copertura-scopertura grottesca tanto quanto sensuale. Eccetera, non   qui la sede per una analisi del romanzo.

Sempre la ricerca sul Campo semantico generato da *divano* individua una zona del capitolo 5 dove si descrive il boudoir di Lisa. Il sistema di screpolature, sdruciture, strappi del tessuto del divano e della tappezzeria, il sistema della corruzione, insomma, appunto sempre ridicola e ripugnante (p. 46): siamo di fronte a un equivalente oggettuale delle scoperture di vestiti femminili che disvelano cosce e quant'altro.

A una articolazione interpretativa siffatta (appena abbozzata) si arriverebbe comunque anche solo leggendo attentamente il testo?

¹⁶ MORAVIA, *Gli Indifferenti* cit., p. 5.

¹⁷ *Ibid.*, pp. 51 sgg.

Ovvero, un critico non ha bisogno di LSA per formulare una ipotesi ermeneutica? Ovvio che sì, ma il punto è avere ausili alla ricerca, da sempre: questa è l'evoluzione scientifica del rapporto fra informatica (in particolare l'informatica computazionale) e umanistica. La frammentazione del testo imposta da LSA fa effettivamente riflettere l'utente su particolari che nell'insieme possono sfuggire. Il sistema delle co-occorrenze permette di fondare geometricamente relazioni concettuali e immaginali cui magari non avremmo dato rilevanza nella lettura corrente. E così via.

Naturalmente la ricerca su domini semantici generati da certe parole può avere risultati più immediati ma meno singolari. Ad es. se partiamo dal Campo semantico di *strada* abbiamo una coerente famiglia di lemmi che individuano gli esterni di *Gli Indifferenti* con una tenuta descrittivo-simbolica piuttosto prevedibile da parte del critico letterario. Così per *lampada*, rispetto agli *interni*, o per *luce* ecc. Meglio, secondo noi, orientare la ricerca su una ipotesi di lavoro specifica prefigurata e da verificare, possibilmente non ancorata alla coerenza obbligatoria dei procedimenti retorici ma fondata su un'idea anche arida e criticamente creativa, pur rischiando un borderline fra interpretazione e sovrainterpretazione, salvo poi riaggiustare razionalmente il tutto durante e dopo l'inchiesta.

Si tratta insomma di credere nello strumento informatico come ausilio attivo non solo nell'analisi dei processi linguistici e retorici, ma anche nella costruzione di discorsi interpretativi profondi. D'altronde, per il critico letterario, la semantica è sempre latente.

6. *Appendice II. Un modello matematico per i Campi semantici*

Da un punto di vista formale, per formulare un modello di Campo semantico statico e dinamico, è stata prima definita una misura di distanza semantica negli Spazi vettoriali dove il Campo emerge ('lo spazio LSA'), quindi è stata derivata una formula per definire i Campi semantici e la loro evoluzione temporale.

La funzione di distanza scelta è la Cosine Similarity ($|\cdot|_{\text{LSA}}$) nello spazio vettoriale ridotto dopo aver eseguito il processo LSA: data w generica parola nel *corpus* C , c concetto centrale del Campo semantico (dove c è anche un termine in C), τ soglia empirica (numero positivo $\in [0,1]$) di valutazione della consistenza del Campo semantico, dati $\vec{c}$ e $\vec{w} \in \Gamma_{\text{LSA}}$ (Lessico del *corpus* nello spazio vettoriale ridotto

dalla funzione LSA) definiamo il lessico di un Campo semantico L_C come $L_C = \{ w \in C \text{ tale che } |\vec{c} - \vec{w}|_{LSA} < \tau \}$ (**formula 1: lessico di un Campo semantico**).

In questo contesto, data un'unità testuale $t \in C$, la similarità semantica tra un concetto c e l'unità t viene definita come

$\text{sim}(c, t) = |\vec{c} - \vec{t}| = (\sum_i c_i * t_i) / |\vec{c}|_{LSA} |\vec{t}|_{LSA}$ (**formula 2: similarità semantica**)

dove $\vec{c} = \sum_{w \in L_C} \vec{w}$, $\vec{t} = \sum_{w \in t} \vec{w}$, $c_i * t_i$ è il prodotto tra le rispettive componenti i -esime dei due vettori $\vec{c}$ e $\vec{w}$ nello spazio LSA, e L_C è il lessico del Campo semantico definito in precedenza.

In questo modo viene definita la possibilità di stimare il 'valore' di un concetto all'interno di un'unità testuale (nei nostri esperimenti è stato scelto il paragrafo come unità).

Una quantificazione dello sviluppo narrativo del concetto c nel testo viene ottenuto da una funzione discreta $f: \tilde{N} \times T \rightarrow [0, 1]$, dove T è l'insieme delle unità testuali t_i e $\tilde{N}$ l'insieme dei concetti. Data μ rappresenta la media e σ la deviazione standard della distribuzione $\text{sim}(c, t_i)$ la funzione di sviluppo narrativo è definita come $f(t_i, c) = (\text{sim}(c, t_i) - \mu) / \sigma$ (**formula 3: sviluppo narrativo di un'unità testuale**).

Quindi lo sviluppo del concetto narrativo può anche essere visto in altre segmentazioni dell'opera, per esempio quelle canoniche dei capitoli, costituite ciascuna da un insieme di unità testuali. Più formalmente:

dato un Capitolo $C_h \subset C$, la semantica globale del capitolo nello sviluppo dell'opera può essere vista come una funzione che mappa i diversi contributi locali di f in un valore aggregativo del capitolo. Ossia $f(C_h, c) = \psi(f(t_i, c)) = 1/N \sum_{t_i \in C_h} f(t_i, c)$ (**formula 4: sviluppo narrativo di un capitolo**).

dove N è il numero di unità t_i nel capitolo. Questa equazione rappresenta la media standard dei contributi delle diverse funzioni delle unità testuali di cui è composto il capitolo.

Ulteriori varianti della funzione ψ possono essere definite inserendo altri contributi oltre la media, per esempio la deviazione standard moltiplicata per un coefficiente: quest'ultimo contributo metterebbe in rilievo distribuzioni dotate di picchi (quindi con frasi dove il concetto è reso particolarmente evidente) rispetto ad altre con distribuzioni uniformi (ossia dove il concetto è più uniformemente distribuito nel testo del capitolo).

Ubiquitous computing Challenges for the software engineer

ABSTRACT: Software engineers are facing revolutionary changes. The boundaries of the closed world in which they used to live, the context in which they operated and their products were embedded are now changing continuously, autonomously, and unpredictably. It is necessary to adapt to changes quickly, flexibly, cheaply, and reliably. Software products evolved from monolithic and closed architectures to distributed, open, and decentralized architectures. Networks replaced single machines, components and services replaced ad-hoc developments from scratch. In many cases, computers are disappearing as recognizable units: computing capabilities are spread in the environment as a myriad of ad-hoc interacting autonomous devices. This article traces through the evolution of these transformations from the viewpoint of the software engineer, who is responsible for the design of new applications. It tries to identify the challenges and possible directions for solutions.

1. *Introduction*

Computational devices are everywhere and software provides them with functionalities that makes them indispensable in our society and even in our everyday life. A typical example is the ability to access the Internet and its services almost everywhere and anytime, enabled by a variety of devices, such as personal computers, PDAs, or cellular phones. Cameras incorporate functionalities to classify and sort photos and can get positional information via GPS to attach as descriptors to outdoor photos. Cellular phones incorporate computing facilities, such as spreadsheets or word processors. Home appliances, such as a refrigerator, can advise on items to buy.

Computational intelligence, however, is not simply added to existing devices. A silent revolution is happening today, which is spreading computational intelligence in the environment in which we live and operate. The ‘traditional’ computer, as we all know it – i.e.,

a clearly recognizable computational device, with its functional units and wiring equipments – is disappearing, being replaced by a myriad of autonomous intelligent devices that are diffused in the environment. The environment, in fact, is increasingly populated with sensors and actuators, which can coordinate their behaviors to provide some useful functionalities for the humans or for other devices belonging to some shared context. This allows, for example, to develop home environment that can recognize the presence of people and adapt their functionalities to their abilities and preferences.

This emerging world has been called *ubiquitous* or *pervasive* computing. A vision of ubiquitous computing can be traced back to the pioneering work of Mark Weiser at Xerox PARC. Let us quote here some of the visionary statements by Mark Weiser that are now becoming real:

The most profound technologies are those that disappear. They weave themselves into the fabric of everyday life until they are indistinguishable from it.

Mark Weiser also spoke of *calm technology* as:

technology that recedes into the background of our lives.

A typical property of ubiquitous computing settings is the continuous presence, in time and space, of computing power, and the availability of a wide and heterogeneous set of networking technologies, devices, and infrastructures, such as wired networks, wireless low-range, wireless wide-range, with infrastructure (access points, routing nodes), without infrastructure (ad-hoc). Recently, the focus has been directed to the creation of systems that can compose and re-compose their functionalities in a dynamic manner. The terms *self-organizing* and *autonomic* systems are often used to denote systems that can behave in a self-managed manner, especially to react to changes in the context in which systems operate. Context changes may for example be due to changes in the physical context due to people mobility.

Several examples of ubiquitous computing scenarios are already in place or are being developed. A typical case, is the so called *ambient intelligence* scenario [1]: from home automation to guided visits in museums, from interacting functionalities within cars to functionalities that connect cars with the external environment and with other cars, from medical systems within hospitals to continuous remote patient monitoring and assistance in their homes. Other examples may be

found at the enterprise level, where individual information systems are increasingly interconnected in a dynamic fashion to support networked enterprises that can federate in a flexible and changeable manner.

This article tries to identify the challenges that continuously evolving and dynamic ubiquitous computing applications are posing to software engineers. We start in Section 2 with a retrospective analysis that tries to capture how software engineering methods were progressively evolving up to the current situation. In Section 3 we then turn our attention to the consequences of ubiquitous computing on software engineering and to the research challenges. Section 4 draws some final conclusions.

2. A Brief Retrospective View of Software Engineering Evolution

The goals of software engineering are to develop methods and tools for the development and evolution of software applications that satisfy some predetermined quality attributes under predetermined schedule and budget constraints, and to define and manage suitable development and evolution processes. Hence software engineering has two main concerns: the software *product* and the software *process*. As we will show in the sequel, both have been evolving towards increasing levels of dynamism, flexibility, and decentralization¹.

The origin of this evolution can be traced back to 1968, when software engineering was recognized as a crucial scientific topic [3] and the *waterfall development process* was developed [4] as a reference process model for software engineers. The waterfall process was an attempt to bring discipline and predictability to software development. It specified a sequential, phased workflow where a requirements elicitation phase is followed by a design phase, followed by a coding phase, and finally by a validation phase. The proposed process model was fixed and static: the decomposition into phases and the link between one phase and the next were precisely defined. The conditions for exiting a phase and entering the next were also precisely formulated, with the goal of making any rework on previously completed phases unnecessary, and even forbidden. The underlying motivation was

¹ For a more detailed history of software engineering, the reader can refer to [2].

that rework – i.e., later changes in the software – was perceived as detrimental to quality, responsible for high development costs, and for late time to market. This led to investing efforts in the elicitation, specification, and analysis of the needed requirements, which had to be frozen before development of the entire application could start.

This approach made very strong implicit and tacit assumptions about the world in which the software was going to be embedded. This world was assumed to be static, i.e. the system's requirements were assumed not to change. The problem was just to elicit the requirements right, so that they could be fully specified prior to development. Completely and precisely specified requirements would ensure that that no need for software changes would arise during development and after delivery.

Another assumption was that the organizations in which the software had to be used had a monolithic structure. Indeed, at that time companies were isolated entities with centralized operations and management. Communication and coordination among different companies were limited and not crucial for their business. Centralized solutions were therefore natural in this setting.

Based on these assumptions and on the software technology available at the time, engineers developed monolithic software systems to be run on mainframes. Even though the size of the application required several people to share the development effort, little attention was put on the modular structure of the application. Separately developed functions were in any case bound together at translation time and no features were available to support change in the running application. To make a change, the system had to be put in offline mode, the source code had to be modified, each module had to be recompiled, rebound, and redeployed. Because little or no attention was put on the modular structure, changes were difficult to make and their effect was often unpredictable.

The need to accommodate change and software evolution asked for more flexible approaches. It became clear that, in most practical circumstances, system requirements cannot be fully gathered upfront, because sometimes customers do not know in advance what they actually require. If they do, requirements are then likely to change, maybe before an application that satisfies their initial specification has been developed completely. This of course questions the tacit assumptions on which the waterfall model is based. Moreover, software architectures turned out to be difficult to modify, because the tight coupling among the various parts made it impossible to isolate the

effect of changes to restricted portions of the application. Seemingly minor changes could have a global effect that disrupted the integrity of the whole system.

Finally, the evolution of business organizations also asked for more flexibility. The structure of monolithic and centralized organizations evolved towards more dynamic aggregates of autonomous and decentralized units. Software development processes also changed, because off-the-shelf components and frameworks became progressively available. Software development became distributed and decentralized over different organizations: component developers and system integrators.

A first conceptual contribution toward supporting evolution of software systems was the principle of *design for change* and the definition of techniques supporting it. Parnas [5-7] suggested that software designers should pay much attention in the requirements phase to understand the future changes a system is likely to undergo. If changes can be anticipated, design should try to achieve a structure where the effect of change is isolation within restricted parts of the system.

Parnas elaborated a design technique, called *information hiding*, through which an application is decomposed into modules, where the major sources of anticipated requirements changes are encapsulated and hidden as module secrets and inaccessible to other modules. A clear separation between *module interface* and *module implementation* decouples the stable design decisions concerning the use of a module by other modules from the changeable parts that are hidden inside it. Module clients are unaffected by changes as long as changes do not affect the interface. This great design principle enables separation of concerns, a key principle that supports multi-person design, and software evolution.

The concepts of information hiding and interface vs. implementation were incorporated in programming languages and supported separate development and separate compilation of units in large system [8]. As an example, let us consider the Ada programming language, which was designed in the late 1970s.

Further advances were provided by object-oriented design and object-oriented programming languages, which provided improved support to modularization, incremental development, and dynamic change. According to object-oriented design, a software system is decomposed into *classes*, which encapsulate data and operations. Furthermore, the binding between an operation requested by a client module and the implementation provided by a server module can change dynamically at run time. This is a fundamental feature to support software evolu-

tion in a flexible manner. Moreover, most object-oriented languages provide it in a way that retains the safety of static type checking [9].

Off-the-shelf components [10] marked the next major evolution step, which asked for changes in the methods used to design applications. The traditional dominant design approaches based on top-down decomposition were at least partly replaced by bottom-up integration. With the advent of off-the-shelf components, process development became decentralized, since different organizations at different times are in charge of different parts of a system. This has clear advantages in terms of efficiency of the development process, which can proceed at a higher speed, since large portions of the solution are supported by reusable components. Decentralized responsibilities also imply less control over the whole system, because no single organization is in charge of it. For example, the evolution of components to incorporate new features or replace existing features is not under control of the system integrator.

Components are often part of a system that is distributed over different machines. Indeed, distribution has been another major step in the evolution of software composition. In a distributed system, the binding among components is performed by the *middleware* [11].

Continuous evolution, adaptation, and decentralization are nowadays reaching unprecedented levels. The main reason is that in ubiquitous computing the boundaries between the software system and the external environment are subject to continuous change. For example, in ambient intelligence settings, often computational nodes are mobile and the physical topology of the system changes dynamically. The logical structure (i.e., software composition) is also required to change to support location-aware services.

Similar problems arise in the case of new enterprise-level application. Emerging enterprise models are based on the notion of networked organizations that dynamically federate their behaviors to achieve their business goals. Goal-oriented federated organizations may be supported by dynamic software compositions at the information system level. In this case, the composition of the distributed software is made out of elements owned and run by different organizations (Web services). Each individual organization exposes fragments of its information system that offer services of possible use for others and, in turn, it may exploit the services offered by others. The binding between a requested and a provided service may change dynamically.

We have used the term *open-world software* to describe this emerging situation [12]. In the early days of software engineering, a closed-world assumption was made. It assumed that the requirements speci-

fying the border between the environment and the system to develop were fairly stable. It assumed that one organization was in charge of developing, operating, and evolving the software. This turned out to be contradicted by reality. We observe today that instability is the norm: software applications are in a kind of *perpetual beta version*. Software increasingly lives in an open world, where the boundaries with the real world change continuously, even as the software is executing. Moreover, the parts that compose a running application are decentralized components, developed, administered, run, and evolved by independent parties. Such components are often called *services* and the resulting architecture is called *service-oriented architecture* – SOA.

SOAs structure applications as a composition of parts (services) that are developed and run as autonomous application units by independent developers. In principle, many potential service providers and service consumers exist. Service providers contribute to a service marketplace. Service consumers, in turn, may compose existing services provided by others to develop new added-value services, thus playing also the role of service providers. A typical way to integrate a number of services to form a new, composite service is by means of *orchestration*. In this case, a workflow process coordinates the execution of external services by invoking them according to its internal logic.

SOAs [13] can be characterized by the fact that (1) there is a global registry where, once a new service is delivered, the description of the function it performs is stored, and (2) a discovery function is provided, through which one can search for a service yielding some desired function. Once a service has been discovered, it can be used remotely through direct binding. In principle, registration, discovery, and binding can be performed dynamically at run time.

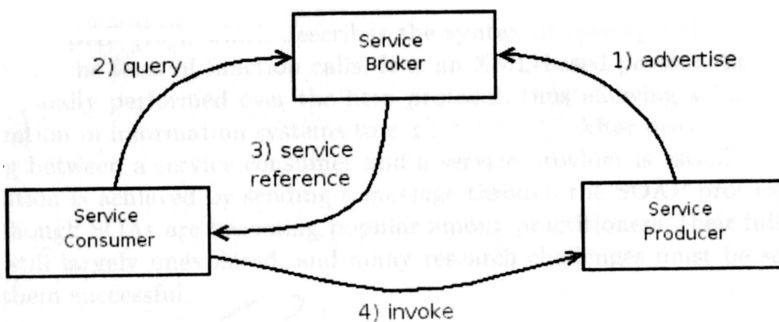


Fig. 1. Service oriented architecture

The figure clearly illustrates the roles of *service consumer*, *service provider* and *service broker*. The service provider provides a service. The service is described through a specification (which we may assume to consist roughly of its interface description). The provider advertises the presence of the service through the broker, which maintains a link to the service and stores its interface. The consumer who needs a service must submit a request to a broker, providing a specification of the requested service. By querying the broker for a service, it gets back a link to a provider. From this point on the consumer uses the operations offered by the provider by communicating with it directly.

Many middleware solutions exist which implement support for SOAs; for example, JINI network technology [14] or OSGI [15]). At the enterprise-level, *Web services* [16, 17], and the rich set of standard solutions that accompany them, are emerging as a very promising approach among practitioners.

IBM [18] defines a Web Service as:

[...] a new breed of Web application. They are self-contained, self-describing, modular applications that can be published, located, and invoked across the Web. Web services perform functions, which can be anything from simple requests to complicated business processes. Once a Web service is deployed, other applications (and other Web services) can discover and invoke the deployed service [...].

Web service (WS) technology offers a big advantage over other service-oriented technologies in that it aims at making interoperation between different platforms easy. This is particularly useful in an open-world environment, where it is impossible to force all the actors to adopt a specific technology to develop their systems. WS technology consists of a set of standards that specify the communication protocols and the interfaces exported by services, without imposing constraints on the internal implementation of the services. Conversely, other solutions, like Jini, are based on Java both for the communication and for the implementation of services. Interoperability allows a seamless integration of the web service solutions developed by different enterprises.

Communication among Web services is based on the Simple Object Access Protocol (SOAP [19]), which describes the syntax of messages that can be exchanged in the form of function calls. It is an XML-based protocol and interaction is usually performed over the http protocol, thus allowing a less intrusive intervention in infor-

mation systems to export services. After discovery, once the binding between a service consumer and a service provider is established, communication is achieved by sending a message through the SOAP protocol.

Although SOAs are becoming popular among practitioners, their full potential is still largely unexplored, and many research challenges must be solved to make them successful.

3. Challenges

SOAs provide what can be called a *discovery-based binding* among components. Potentially, as we observed, the choice of possible targets for a binding and the establishment of the binding can be performed dynamically at run time. In practice, however, the current practice is based on fairly static binding schemes, established at deployment time. Thus, for example, a workflow that orchestrates external services is bound to a predetermined set of services. In ubiquitous computing settings, however, dynamic compositions are needed to support context-aware behaviors. This would require the workflow to bind to external services in a dynamic way, by making the discovery function a run-time functionality.

To push this idea to its extreme, one could design a workflow that composes external services in a way that only the specification of the required external services must be known at design time. The specification should describe not only the signature of the service, but also its required behavior (in terms of functional and nonfunctional properties). At run time, any published service that complies with the specification would then be a candidate for use.

This vision is currently investigated by several research efforts. We are currently distilling a design methodology along these lines from a number of projects in which we have been involved, including the EU SeCSE and PLASTIC projects, and the ArtDeco FIRB Project. The main open research challenges on which we focus fall into three areas: service specification, verification, and self-management.

Service specification is needed to support discovery-based binding. Although software specification has been an active research area since the late 1960's, no consolidated results are available and used in practice. Furthermore, in the case of services the problem has its own specific slant. We need to be able to specify both the syntactic and the semantic aspects of services. And we need to specify not only func-

nal properties, but also nonfunctional ones. How this can be done is an active research area.

Service verification is another crucial aspect, and it is well known that dynamic evolution may conflict with correctness. Traditionally, verification is performed off-line, during software development, prior to run time. The underlying assumption is that once the software is running it does not change. But here the situation is totally different. Software indeed changes at run time. In particular, in a SOA, service providers who own a service may change the service implementation without any notice, and the change might violate (maybe inadvertently) the contract established in the specification. These problems ask for continuous verification, that extends to run time. Several research efforts are currently addressing the problem this problem.

Service self-management is a required functionality to support continuous operation even in the presence of anomalous or unpredicted situations. For example, run-time verification might discover that an external service invoked by a composite service does not satisfy the specification. The composite service might exhibit a self-managed behavior through which it may try to overcome the fault. It might, for example, retry the invocation – hoping for a transient failure – or look for a replacement by activating a discovery phase. It might even try to re-plan its workflow to try to match its required functionality with the offered services advertised by registry. Self-managing systems are also called *autonomic* [20].

4. Conclusions

Ubiquitous computing is leading to profound changes in software. The demand for dynamic change and adaptation on one side, and the availability of decentralized services on the other, are leading towards software architectures where both components and their interactions may change dynamically at run time. Presently, this is done in an ad-hoc manner. Moreover, flexibility and dynamism conflict with dependability. Since ubiquitous applications may support critical functions, the need to reconcile them is becoming a central problem in the software engineering research agenda. There is a need for languages, methods, and tools that may support more systematic and predictable designs, which may possibly lead to the desired levels of dependability.

ACKNOWLEDGEMENTS: Partially supported by the PRIN Project Art Deco.

5. References

- [1] K. DUCATEL, M. BOGDANOWICZ, F. SCAPOLO, J. LEIJTEN, J. BURGELMA, *Scenarios for Ambient Intelligence in 2010* (ISTAG 2001 Final Report). IPTS, Seville 2000.
- [2] C. GHEZZI, M. JAZAYERI, D. MANDRIOLI, *Fundamentals of Software Engineering*, Prentice Hall 2003.
- [3] P. NAUR, B. RANDELL, J. BUXTON, *Software Engineering: Concepts and Techniques: Proceedings of the NATO Conferences*. Petrocelli/Charter 1976.
- [4] W. ROYCE, *Managing the Development of Large Software Systems*. Proceedings of IEEE WESCON 1970, pp. 1-9.
- [5] D. PARNAS, *On the Criteria to be Used in Decomposing Systems into Modules*, in «Communications of the ACM» XV, 12, 1972, pp. 1053-1058.
- [6] D. PARNAS, *On the Design and Development of Program Families*, in «IEEE Transactions on Software Engineering», II, 1, 1976, pp. 1-9.
- [7] D. PARNAS, *Designing Software for Ease of Extension and Contraction*. Proceedings of the 3rd international conference on Software engineering 1978, pp. 264-277.
- [8] B. RYDER, M. SOFFA, M. BURNETT, *The Impact of Software Engineering Research on Modern Programming Languages*, in «ACM Transactions on Software Engineering and Methodology (TOSEM)», XIV, 4, 2005, pp. 431-477.
- [9] C. GHEZZI, M. JAZAYERI, *Programming Language Concepts*, John Wiley & Sons, Inc. 1997.
- [10] C. SZYPERSKI, *Component Oriented Programming*, Springer 1998.
- [11] W. EMMERICH, *Engineering Distributed Objects*, John Wiley & Sons 2000.
- [12] L. BARESI, E.D. NITTO, C. GHEZZI, *Toward Open-World Software: Issue and Challenges*, in «IEEE Computer», XXXIX, 10, 2006, pp. 36-43.
- [13] M. PAPAZOGLU, D. GEORGAKOPOULOS, *Service-Oriented Computing: Introduction*, in «Communications of ACM», XLVI, 10, 2003.
- [14] J. WALDO, *Jini Architecture for Network-Centric Computing*, in «Communications of the ACM», XLII, 7, 1999, pp. 76-82.
- [15] O. ALLIANCE, *OSGI Service Platform. Release 3, March 2003*, IOS Press 2003.
- [16] G. ALONSO, H. KUNO, F. CASATI, V. MACHIRAJU, *Web Services: Concepts, Architectures and Applications*, Springer 2004.

- [17] M. BICHLER, K. LIN, *Service-Oriented Computing*, in «Computer», XXXIX, 3, 2006, pp. 99-101.
- [18] U. WAHLI, M. TOMLINSON, O. ZIMMERMANN, W. DERUYCK, D. HENDRIKS, *Web Services Wizardry with WebSphere Studio Application Developer*, IBM Redbook 2002.
- [19] D. BOX, D. EHNEBUSKE, G. KAKIVAYA, A. LAYMAN, N. MENDELSON, H. NIELSEN, S. THATTE, D. WINER, *Simple Object Access Protocol (SOAP) 1.1*.
- [20] J. KEPHART, *Research Challenges of Autonomic Computing*. 27th International Conference on Software Engineering (ICSE'05) 2005.

CARLO GHEZZI

Progettare interfacce web

1. *Impostare bene la navigazione aiuta l'utente a trovare quello che cerca e anche quello che non sa di volere*

Per 'interfaccia' si intende un insieme di strumenti che consente l'interazione con un ambiente, nel nostro caso il Web.

Le interfacce in genere (e il Web non fa eccezione) devono obbedire ad alcune regole, determinate perlopiù dal funzionamento della mente umana. Per chi progetta interfacce è quindi utile avere almeno dei rudimenti di psicologia cognitiva e percettiva, per creare interfacce ergonomiche, intuitive da usare.

La definizione di questa professionalità è *interaction designer*, e nel mondo ci sono diversi specialisti con *skill* notevoli in questo campo, ma chiunque progetti un sito web, anche solo per proprio piacere personale, che lo sappia o no, è un *interaction designer*.

2. *Due tipi di interfacce web*

Fondamentalmente su Web ci troviamo a progettare due tipologie di interfacce, per la navigazione e per l'interazione.

La prima ha la sua espressione principalmente nei menu (oltre che nei link), mentre per la seconda si usano essenzialmente le form.

Poiché la motivazione dell'utente è diversa, e diversa è la funzionalità, queste due tipologie di interfaccia sono sostanzialmente differenti anche da un punto di vista visivo.

Delle due, quella attualmente di gran lunga più importante e più studiata è l'interfaccia di navigazione, se non altro perché non esiste sito che non ne abbia bisogno, mentre esistono molti siti che non presentano form né altri tipi di interazione col server remoto.

* Estratto da S. POSTAI, *Web design in pratica*, Milano, Tecniche Nuove 2006, su gentile concessione dell'editore.

3. *I principi dell'interfaccia (interaction architecture)*

Quello che un'interfaccia (di navigazione o interazione) deve dare all'utente è in buon sostanza il controllo della situazione. È un quadro comandi che l'utente deve sentire di padroneggiare, tanto da non accorgersi nemmeno di cosa sta usando. Quando camminate pensate a come spostare il peso da una gamba all'altra mentre alzate alternativamente i piedi per portarli, uno dopo l'altro, davanti a voi... spostare il peso nuovamente e così via? Certamente no. È per questo che camminare costa poca fatica intellettuale (richiede poche risorse attentive direbbe uno psicologo cognitivo). Le risorse attentive (come i soldi) non bastano mai e quando cominciano a scarseggiare la persona entra in ansia. Potrebbe anche abbandonare quello che sta facendo, se gli pare di non farcela.

Steve Krug ha chiamato «Don't make me think» il suo fortunato libro sull'usabilità, proprio per mettere l'accento su questo fatto. «Non farmi pensare» vuol dire lasciare libero l'utente di dedicarsi al vero scopo della sua interazione, senza farlo concentrare sul modo per raggiungerlo. La padronanza che l'utente vuole avere dell'ambiente in cui si muove non è una questione di 'potere' ma di comprensione.

4. *Predittività*

Significa la possibilità, per l'utente, di capire in anticipo cosa succederà quando lui compirà una determinata azione. Dove andrà seguendo un link o cosa succederà dopo che ha cliccato il pulsante di submit su una form.

Questo vuol dire semplicemente che è necessario dare all'utente tutte le informazioni e gli strumenti di cui ha bisogno per raggiungere i suoi obiettivi.

5. *Semplificazione*

Non date all'utente più scelte di quante siano necessarie, solo perché il vostro sistema consente di farlo. Alcuni siti consentono all'utente di essere personalizzati (grafica, colori, disposizione degli elementi) ma raramente l'utente lo fa. Le maschere di ricerca avanzata dei motori di ricerca vengono usate da un'esigua percentuale di utenti. Le opzioni

più complesse devono essere date come secondarie rispetto alle opzioni più semplici.

Il navigatore non vuole scegliere troppo: dobbiamo prenderci le nostre responsabilità: il progettista siamo noi, non lui.

Inoltre, quando è sensato, prevedete nelle forme alcune scelte di default: per esempio, se immaginate che un'opzione sia molto più frequente di un'altra (per esempio 'cittadinanza italiana') fatela già pre-selezionata.

6. *Informazioni sullo stato del processo*

Non abbandonate l'utente a se stesso, incerto se il processo si sia avviato, se il server sia eccezionalmente lento, se l'applicazione web proprio non funziona o se abbia sbagliato qualcosa. Se l'applicazione è importante ed è possibile che il server non sia istantaneo nella risposta, prevedete di mandare un qualche avvertimento (corrisponde alla musicchetta che sentite al telefono in attesa che vi passino l'interno desiderato).

Lo rassicura che tutto sta funzionando: deve solo avere un momento di pazienza. Può essere il classico orologio con le lancette che girano o, ancora meglio, l'avanzamento di una barra che progressivamente si completa.

7. *Coerenza*

Influenza grandemente la predittività ed è una dote che ogni interfaccia dovrebbe avere.

Sostanzialmente si può riassumere nel fatto che oggetti che hanno la stessa apparenza devono avere lo stesso funzionamento e oggetti che hanno funzionamento diverso devono essere diversi anche nell'aspetto.

Un tipico esempio di incoerenza sono quei menu apparentemente ordinati di voci visualizzate tutte in modo simile: peccato che alcuni siano link ad aree del sito, altri link a siti esterni, altri aprano dei pdf, altri aprano finestre pop up, altri il programma di posta elettronica.

Coerenza non vuol dire fare i bottoni tutti uguali, ma rendere prevedibile il loro comportamento.

8. *Risparmiate il tempo dell'utente, non quello del computer (o il vostro)*

Molte interazioni sono difficili da eseguire perché chi le ha progettate ha trovato una soluzione semplice (per lui) che costa solo poche righe di codice e viene risolta lato client senza impegnare risorse di elaborazione del server.

Tipicamente questo succede quando si effettuano i controlli di congruenza dei campi delle form in Javascript, il che produce raramente una buona gestione dell'errore.

In una form on line (composta da oltre 50 campi!) se si prova ad effettuare un submit senza compilarne nemmeno uno si ottiene un alert riferito al primo dei campi non compilati. Lo script adoperato (un freeware di pubblico dominio) infatti gestisce un errore alla volta. In questo modo l'utente che compia più di un errore se li vedrà segnalati uno alla volta, senza contare che, per la natura degli alert javascript, dovrà chiuderli (rischiando di dimenticarsi il messaggio di errore) prima di riuscire a tornare alla pagina da compilare.

Il modo giusto per gestire l'errore è invece lato server, offrendo all'utente una pagina che gli illustri gli errori fatti e segnali i campi da correggere.

Se nella pagina riproposta per la correzione c'era una password che viene cancellata per motivi di sicurezza, anche questo verrà segnalato. 'Riscrivere la password, in quanto è stata cancellata per motivi di sicurezza'.

9. *Prevedete un 'annulla' ad ogni passo*

L'utente deve sempre poter annullare quello che ha appena fatto.

Questo non significa affatto che vicino ad ogni bottone 'submit' deve sempre esserci un bottone 'reset' che cancella quanto appena scritto.

Significa che, dopo aver cliccato un bottone submit deve esserci sempre la possibilità di tornare indietro di un passo.

A volte questo non è possibile restando all'interno del sistema, in quando il bottone era quello che concludeva il processo, dopo le necessarie richieste di conferma. In questo caso ci sarà un'e-mail o il telefono del customer care.

Nessun pulsante deve essere una condanna.

Se cliccando il back del browser l'utente riceverà un messaggio del tipo 'data missing', 'impossibile ricaricare la pagina perché la sessione

è conclusa' o altre irritanti tecnicaglie del genere si sentirà come Neo quando vede Matrix per la prima volta (per gli amanti del genere).

Prevediamo che possa farlo e prevediamo un messaggio sensato che risponda alle sue motivazioni (evidentemente cambiare o annullare quanto appena concluso).

10. *Aiutate l'utente a concentrarsi*

Nelle pagine di interazione complessa, eliminate quanto più possibile le distrazioni: lasciate solo l'intestazione del sito e il menu principale. Via i box promozionali, i banner ed ogni altro tipo di stimolo non necessario: lasciatelo lavorare in pace.

11. *La 'prima volta' e alla fidelizzazione*

Prima di rilasciare un progetto chiedetevi se la navigazione/interazione è chiara per chi non l'ha mai fatta e se ci sono aiuti sufficienti.

Per gli utenti abituali vanno previste scorciatoie sufficienti e sufficientemente chiare. Quando un utente acquista familiarità con un sito, se ne fa un modello mentale ed è in grado di utilizzare la navigazione in modo diverso, senza effettuare tutti i passaggi che in una prima visita sono necessari per capire i contenuti, i servizi offerti, il modo per fruirne, e così via.

Il nuovo utilizzatore ha bisogno di orientarsi, l'utilizzatore abituale, di arrivare velocemente allo scopo.

Qualsiasi sito ha entrambi gli utenti e si deve tenerne conto.

12. *Il focus dell'attenzione è uno solo*

La mente (umana o animale che sia) riesce a prestare attenzione a una cosa sola alla volta: quando il focus dell'attenzione è concentrato in un punto, l'utente non si accorge di cambiamenti in altre aree della pagina. Quindi se va dato un feedback a pagina ferma (per esempio mediante javascript) il messaggio va visualizzato nell'area in cui ha appena agito l'utente (tipicamente vicino al bottone che ha cliccato).

Per lo stesso motivo le spiegazioni o le *legenda* vanno posizionate nei pressi dell'elemento da chiarire.

Si sente spesso dire che in Rete le persone non leggono le istruzio-

ni, perché l'help desk di diversi servizi viene chiamato per ottenere informazioni che erano presenti nella pagina che l'utente sta consultando: spesso la colpa è della posizione sbagliata. L'utente cerca di capire come utilizzare una determinata funzione e guarda solo nei pressi immediati dell'elemento che gli crea dei problemi.

13. *Rendete l'informazione visivamente disponibile*

Nella memoria a breve termine le persone possono mantenere mediamente 7 elementi (massimo 9). Poiché in un sito web è piuttosto raro che si possano esaurire le possibilità di navigazione con 7/9 voci, è importante che almeno le principali opzioni siano tutte visibili nella stessa schermata, in modo che l'utente le abbia contemporaneamente sott'occhio.

Questa esigenza può essere in contrasto con quelle dell'accessibilità: un cieco non riuscirà a memorizzare tutte le opzioni (oltre le 7/9) e un ipovedente, ingrandendo i caratteri, non avrà comunque nella stessa schermata la panoramica delle possibilità.

Dovendo progettare un sito ad alta accessibilità, può essere necessario prevedere anche un percorso di navigazione più lungo, in modo che la scelta avvenga tra minori opzioni.

14. *Rendete il funzionamento chiaramente comprensibile*

Soprattutto in interazioni complesse, cioè quando l'utente utilizza a distanza delle vere e proprie applicazioni affacciate su Web, è fondamentale che capisca il concetto che sta alla base dell'interazione.

Non si tratta solo di progettare form usabili ed ergonomiche, ma di concepire un funzionamento che possa essere comprensibile nel suo complesso e non solo relativamente al singolo 'atomo' di interazione.

Qualsiasi persona è facilitata nel suo compito se capisce l'ambiente nel quale si muove: agire alla cieca, semplicemente compiendo un passo dopo l'altro seguendo delle istruzioni, è dispersivo, frustrante e fonte di errori.

Per esempio: in quasi tutti i sistemi di home banking, prima di scegliere un servizio (estratto conto, bonifico, o altro) è necessario selezionare il proprio conto corrente da una lista. Se l'utente ha due conti, il significato di questa selezione gli è immediatamente chiaro, mentre

invece se ne ha uno solo, proverà ad effettuare direttamente l'operazione, mentre il sistema gli chiede di selezionare il 'tipo di rapporto'.

Basterebbero poche righe di programmazione per rendere più amichevole questa interazione: se il conto corrente è uno solo, il sistema si seleziona già su quello e non chiede nessuna ulteriore selezione.

L'utente non ha difficoltà a selezionare una singola voce da una tendina, ma poiché non ne capisce il senso, fatica a capire cosa deve fare.

In questo (come in molti) casi basta una leggera modifica al funzionamento per rendere tutto più comprensibile.

15. *Eliminate il 'rumore'*

Il 'rumore cognitivo' è costituito da tutte le presenze non necessarie che distraggono l'utente dal compito.

Una buona interfaccia rende evidenti le informazioni fondamentali per portare a termine il compito e raggiungere l'obiettivo, mentre lascia in secondo piano altre informazioni di carattere più generale, o di infrequente utilizzo.

Questa 'pulizia' è doverosa nel campo dell'interazione mentre è meno necessaria (e più difficile da realizzare) nel campo dell'informazione: se la presenza non necessaria è costituita da una notizia che all'utente Mario non interessa, ma all'utente Giovanni potrebbe interessare moltissimo, questo fa parte dei compromessi inevitabili nella vita (per non parlare del Web).

SOFIA POSTAI

Biblioteca digitale: esperienze e prospettive

1. *Introduzione*

I contenuti digitali, il Web, gli sviluppi di servizi, come l'e-commerce e l'e-learning, si combinano attualmente per creare una moltitudine di opportunità per l'accesso in linea alle risorse informative ed anche per il possibile ri-uso di queste. Tutte queste opportunità devono essere colte soprattutto dalle istituzioni culturali pubbliche, cioè biblioteche, archivi e musei, che hanno come dovere istituzionale quello di facilitare l'apprendimento e l'accesso alla conoscenza della società. La spinta alla costruzione di biblioteche digitali che, come isole organizzate nel mare caotico dell'informazione del Web, assolvano questo compito è quindi quanto mai attuale; tuttavia non è di facile realizzazione. Una serie di problematiche da affrontare e da cercare di risolvere si presentano ad esempio per la sostenibilità delle collezioni digitali ed anche per il coordinamento delle iniziative nella digitalizzazione, necessario per una reale interoperabilità.

Il quadro di riferimento che abbiamo in Italia per l'evoluzione dei servizi digitali delle istituzioni culturali è necessariamente inserito nel programma più ampio delle iniziative che l'Unione Europea ha preso per la realizzazione di biblioteche digitali, di cui la più recente è quella dell'European Digital Library. I risultati della consultazione 'i2010 European Digital Library', realizzata nel 2006, hanno indicato precise priorità su cui concentrare gli interventi. Tra queste priorità è sicuramente da indicare l'interoperabilità e la conservazione preventiva (o preservazione) degli oggetti digitali che, insieme alla comune adesione a standard tecnici, richiede nuove sinergie e coordinamento tra le istituzioni culturali. L'Unione Europea è stata inoltre lo stimolo per una serie di progetti, che sono stati finanziati per affrontare determinate tematiche per l'accesso alle collezioni digitali, progetti a cui l'Italia partecipa o di cui è la coordinatrice. Lo scenario delle biblioteche digitali in Italia, a cui facciamo riferimento, è quindi uno scenario pienamente integrato nel contesto europeo. Lo scopo di questo articolo

è quello di fornire una prospettiva aggiornata sulle realizzazioni attuali e sulle prospettive di evoluzione che possono essere evidenziate. I dati su cui basiamo l'evidenza del quadro attuale delle biblioteche digitali in Italia rappresentano il risultato dell'indagine compiuta dal Progetto *Digital Libraries Applications*, promosso e finanziato dalla Fondazione Rinascimento Digitale¹ e conclusosi nel 2006.

2. Metodologia

Il Progetto *Digital Libraries Applications* ha cercato di predisporre una metodologia per la valutazione della soddisfazione dell'utenza e di stimolare un brainstorming con le diverse istituzioni culturali, per capire come queste stiano diversificando le tradizionali strategie di servizio. Il Progetto si è posto tre obiettivi principali:

- realizzare un'indagine dell'utenza delle biblioteche digitali,
- approfondire le tematiche di organizzazione e di coordinamento delle biblioteche digitali,
- indagare le problematiche tecnologiche, definendo i problemi di accessibilità, i diversi modelli di architettura dei servizi, i risultati della ricerca in corso.

Per realizzare l'indagine, sono stati riuniti in un Gruppo di studio alcuni esperti di istituzioni culturali, rappresentativi delle diverse tipologie di biblioteca digitale e con diverse esperienze e conoscenze².

¹ La Fondazione 'Rinascimento Digitale – Nuove Tecnologie per i Beni Culturali', presieduta dal Prof. Galluzzi, è nata nel 2005 al fine di promuovere l'applicazione, secondo standard di elevata qualità, delle nuove tecnologie dell'informazione e della comunicazione per la valorizzazione dei beni culturali, anche in collaborazione con altri enti, iniziative di ricerca, consulenza, documentazione, promozione, formazione e divulgazione.

² Gli esperti del Gruppo di studio sono stati numerosi e purtroppo non mi è possibile citarli tutti. La collaborazione è stata preziosa ed il presente lavoro non sarebbe stato possibile senza il contributo che ognuno di loro ha dato. L'elenco completo degli esperti che hanno collaborato all'indagine è consultabile in: A.M. TAMMARO, S. CASATI, D. LUZZI, *Biblioteche digitali in Italia*: <http://www.rinascimento-digitale.it/DLA/RapportoSintesi_DLA.pdf>.

La metodologia dell'indagine attuata dal Gruppo di studio del Progetto *Digital Libraries Applications* si è basata su diversi strumenti di raccolta dei dati:

- una rassegna delle applicazioni di biblioteca digitale esistenti, basata sulla documentazione e sull'esperienza personale degli esperti del Gruppo di studio;
- un'indagine rivolta allo staff delle diverse istituzioni culturali, per capire il modello di servizio;
- interviste ad esperti (usando il metodo Delphi) per concentrarsi sulle problematiche evidenziate;
- una comparazione delle funzionalità ed accessibilità dei siti web delle biblioteche digitali;
- un'indagine dell'utenza per testare una possibile metodologia di valutazione della biblioteca digitale.

Il modello di valutazione usato, concentrato sulle istituzioni culturali e sugli utenti dei loro servizi, è sintetizzato nella Tabella 1.

Tab. 1 Modello di valutazione

Istituzioni Culturali	Utenti	Output (Risultati)	Outcomes (Impatto)
Approcci e strategie di servizio delle biblioteche digitali	Bisogni, priorità e percezione dei servizi	Soddisfazione dell'utente Misurazione:	Raggiungimento della <i>mission</i> dell'Istituzione culturale.
- Contributi degli esperti alla discussione	Analisi demografica degli utenti	Gap tra aspettative e percezione dei servizi	Misurazione: Quanto la biblioteca digitale supporta le attività abituali dell'utente?
- Indagine dello stato dell'arte delle biblioteche digitali in Italia	Fattori Socio-economici che hanno un impatto sugli usi delle applicazioni delle biblioteche digitali	Frequenza d'uso delle risorse digitali e dei servizi	Cosa non sarebbe stato possibile senza la biblioteca digitale?
- Contenuti digitali e funzionalità siti web	(come e-learning, e-government, e-commerce)		

L'indagine è stata svolta in tre fasi. La prima fase è stata dedicata a raccogliere informazioni sulle attuali biblioteche digitali, insieme alla discussione con gli esperti per la definizione del contesto teorico e socio-economico di riferimento delle biblioteche digitali. In questa fase sono stati distribuiti i questionari allo staff di diverse tipologie di biblioteche digitali, coordinati dai diversi esperti facenti parte del Gruppo di studio, che hanno poi analizzato e sintetizzato i dati. Sono state inoltre svolte delle interviste agli esperti per evidenziare il contesto di riferimento delle biblioteche digitali in Italia. I risultati ottenuti sono stati utilizzati nella seconda fase, impiegata per definire gli strumenti di misurazione per l'indagine dell'utente, proposti in tre aree:

1. risorse digitali, servizi e loro uso;
2. soddisfazione dell'utente per le risorse digitali ed i servizi;
3. misurazione dell'impatto.

Infine, nella terza fase, il Gruppo di studio ha realizzato l'indagine ed analizzato i dati. L'indagine ha sicuramente dei limiti, a cominciare dalla scelta di realizzare una molteplicità di studi di caso, per un confronto finale dei risultati, invece che decidere per un'indagine più estesa. Questa scelta, che è stata necessaria per mancanza di risorse, tuttavia è stata anche guidata dalla consapevolezza della diversità dell'utenza di riferimento delle istituzioni culturali e quindi dalla difficoltà di capire le problematiche dell'utenza in un quadro ampio e generale. La ripetizione di una molteplicità di studi di caso, potrà portare successivamente ad un confronto dei dati, per evidenziare linee comuni di interesse e di intervento. Inoltre l'indagine non ha considerato gli utenti remoti e non ha incluso i non utenti della biblioteca digitale.

Lo scopo del modello di valutazione è stato quello di facilitare una prima rilevazione di dati qualitativi, senza abbandonare la raccolta di dati statistici quantitativi. Da una rassegna della letteratura e della documentazione, il Gruppo di studio ha evidenziato che la valutazione della biblioteca digitale si è finora concentrata su misurazioni quantitative di uso delle risorse e dei servizi, tuttavia pochi studi hanno cercato di definire la prospettiva dell'utente. Il Gruppo di studio ha preso come riferimento alcuni degli studi più importanti, tra cui il Progetto E-measures dello SCONUL³, il Progetto di ARL, chiamato

³ Lo scopo era di produrre un insieme di statistiche per misurare l'uso dei servizi

E-Metrics⁴ ed il Progetto COUNTER⁵. Il Progetto eVALUED⁶ infine, si è rivelato di particolare interesse per gli scopi del Gruppo di studio, in quanto ha sviluppato un pacchetto per facilitare la valutazione delle biblioteche digitali⁷. La soddisfazione dell'utente è stata identificata come misura del gap tra le priorità che l'utente ha indicato

digitali nelle biblioteche universitarie in UK. Il Progetto si è basato sulla rilevazione periodica realizzata nella totalità delle biblioteche universitarie inglesi. In conclusione dopo due anni, il Progetto ha deciso di sospendere la raccolta dei dati, perché ritenuti non veramente significativi per capire il reale rendimento delle biblioteche digitali. The Society of College, National and University Libraries (SCONUL) produce regolarmente l'Annual Library Statistics, dove sono riportate anche tutte le misure sviluppate dal progetto durante la sua esistenza. Cfr. S. Town, *E-Measures: A Comprehensive Waste of Time*, in «VINE», XXXIV, 4, 2004, pp. 190-195.

⁴ R. MILLER, S. SCHMIDT, *E-metrics: Measures for Electronic Resources*, Keynote Delivered at the 4th Northumbria International Conference on Performance Measurement in Library and Information Services, Pittsburgh, 12-16 august, 2001, ARL <<http://www.arl.org/>>, <<http://www.arl.org/stats/newmeas/emetrics/miller-schmidt.pdf>>.

⁵ Un Gruppo di lavoro in collaborazione con alcuni fornitori (Working Group on Database Vendor Statistics) ha indagato come raccogliere i dati dalle banche dati dei fornitori; dai risultati di questo Gruppo ha preso poi avvio il Progetto COUNTER. Cfr. J.C. BLIXRUD, *Measures for Electronic Use: The Arl E-Metrics Project*. IFLA Satellite Conference Statistics in practice - Measuring & Managing, Loughborough, UK, 13-25 august, 2002, Loughborough University <<http://www.lboro.ac.uk/>>, <<http://www.lboro.ac.uk/departments/dils/lisu/downloads/statsinpractice-pdfs/blixrud.pdf>>; EAD., *Assessing Library Performance: New Measures, Methods, and Models*. IATUL Proceedings Libraries and Education in the Networked Information Environment, Ankara, Turkey, 2-5 june, 2003, IATUL <<http://www.iatul.org/>>, <http://www.iatul.org/conference/proceedings/vol13/papers/BLIXRUD_fulltext.pdf>. Il Progetto COUNTER ha definito le misure standard che i fornitori di servizi di accesso a risorse digitali rendono attualmente disponibili.

⁶ Cominciato nel 2001 e completato nel 2004, ha cercato di andare oltre la misurazione degli indicatori di rendimento per affrontare la valutazione dei risultati in rapporto alla fornitura di servizi informativi elettronici. S. THEBRIDGE, *Evaluating Electronic Information Services: A Toolkit for Practitioners*, in «Library and Information Research», XXVII, 87, 2003, pp. 38-46, E-LIS <<http://eprints.rclis.org/>>, <<http://eprints.rclis.org/archive/00003421/>>.

⁷ Descritto da: S. McNICOL, *The Evalued Toolkit: A Framework for the Qualitative Evaluation of Electronic Information Services*, in «VINE», XXXIV, 4, 2004, pp. 172-175.

per le risorse digitali e per i servizi e la percezione che lo stesso ha dei servizi e delle risorse digitali disponibili. A questa misura del gap, è stata aggiunta la misura della frequenza d'uso delle risorse e dei servizi e contemporaneamente dell'essere utente interno e remoto di diverse biblioteche digitali. Per una misurazione dell'impatto, cioè una misura di come i servizi della biblioteca digitale contribuiscano al successo dell'utente, la scelta del Gruppo di studio è stata quella di ritenere che questo sia strettamente legato al successo dell'istituzione di appartenenza della biblioteca digitale, come espressa nella *mission* o in altri documenti programmatici dell'istituzione culturale. Si è quindi individuata la necessità di trovare una misura che potesse identificare ciò che è critico per la missione della singola biblioteca digitale, come anche per le attività dei suoi utenti. Per una indicazione di impatto e per una sua misurazione, ci si è limitati ad assicurarsi che i servizi della specifica biblioteca digitale analizzata nello studio di caso, fossero forniti in modo da essere di supporto (o se si vuole utili) alle attività dei loro utenti ed alle loro abitudini attuali di ricerca ed uso dell'informazione. La misurazione dell'impatto così inteso è stata indagata attraverso i commenti liberi nel questionario e le specifiche domande richieste nelle interviste. In particolare si è tenuto particolarmente conto di commenti qualitativi negativi e neutri (cioè né negativi né positivi).

Di solito la valutazione di soddisfazione dell'utente, le misurazioni dell'uso e dell'impatto dei servizi della biblioteca digitale sono fatte separatamente. Tuttavia il Gruppo di studio ha ritenuto che i tre processi di misurazione e valutazione dovevano essere complementari, e sicuramente i risultati, messi a confronto, non avrebbero dovuto essere in conflitto. Sono stati quindi elaborati alcuni criteri metodologici, che sono tra i risultati probabilmente più interessanti dell'indagine dell'utenza ed una parte importante del Progetto *Digital Libraries Applications*.

3. Definizione funzionale di biblioteca digitale

La discussione sulla definizione di biblioteca digitale è stato uno dei primi temi di discussione del Gruppo *Digital Libraries Applications*. Il tentativo di formulare una definizione comune di biblioteca digitale è stato originato da un'esigenza concreta, non da un problema astratto o accademico. Era necessario, infatti, stabilire un linguaggio minimo

comune, per garantire la comprensione e lo scambio di opinioni tra esperti provenienti da archivi, biblioteche e musei, insieme a studiosi ed informatici con un diverso background. Fin dai primi incontri è stato evidente che sarebbe stato difficile arrivare ad un accordo, senza condividere prima concetti come quello di biblioteca, archivio, deposito, collezioni, che rivelano il diverso approccio disciplinare delle istituzioni culturali. Sarebbe stato importante poter usare questa occasione per una discussione approfondita sui fondamenti delle discipline del patrimonio culturale, ma necessariamente si è dovuto limitare il confronto interdisciplinare al nuovo ambito digitale, concentrandosi sulle problematiche che sono comuni ad archivi, biblioteche e musei.

La discussione è stata molto animata perché, al di là della forma e della sintassi, gli esperti partecipanti al Gruppo di studio erano coscienti che la definizione servisse non solo a fare una fotografia di sintesi delle attuali applicazioni di biblioteca digitale ma avesse il compito importante di fornire un riferimento sintetico dell'ambito in cui ci si muove, una specie di *vision statement*, dove dovrebbe avere rilievo un concetto ideale di biblioteca digitale, magari desunto dalle buone pratiche emergenti a livello internazionale. La biblioteca digitale è quindi un concetto in evoluzione e bisogna prendere atto che ci sono diverse interpretazioni del concetto di biblioteca digitale, qui definite come un concetto tradizionale ed uno più innovativo.

Il concetto tradizionale, su cui ci si è accordati per una definizione funzionale, è focalizzato sulle funzioni di organizzazione della collezione della biblioteca digitale:

La biblioteca digitale è un insieme organizzato di risorse, strumenti, servizi e personale specializzato a cui una o più organizzazioni offrono prontamente e economicamente l'accesso, interpretano e distribuiscono le informazioni e assicurano la persistenza nel tempo delle collezioni digitali.

Il concetto più innovativo vede la biblioteca digitale come uno strumento al centro di un'attività intellettuale, che non ha confini logici, concettuali, fisici o temporali o barriere all'informazione. Piuttosto che concentrarsi sui contenuti e sulla loro organizzazione, il concetto innovativo di biblioteca digitale è centrato sulle persone, con lo scopo di fornire loro esperienze interessanti, nuove, personalizzate. Questa visione innovativa di biblioteca digitale sembra adattarsi bene al concetto di Spazio Informativo ed è stata espressa nel Reference Model

della biblioteca digitale realizzato dal Progetto DELOS, che ha partecipato al Gruppo di studio⁸.

I due concetti qui evidenziati non sono in opposizione ma sono complementari l'uno all'altro: uno si concentra sull'accesso e sulla organizzazione di risorse digitali, mentre l'altro si concentra su servizi innovativi di visualizzazione e condivisione di conoscenza. I due concetti hanno tuttavia implicazioni diverse per le motivazioni per cui si realizza una biblioteca digitale, per le risorse che fanno parte della collezione e per i servizi da rendere disponibili. Mentre è condiviso da tutti che la biblioteca digitale debba avere un valore aggiunto rispetto alle biblioteche tradizionali, i due concetti di biblioteca digitale definiscono questo valore in maniera diversa. Il concetto tradizionale si concentra sul recupero dell'informazione, che viene migliorato attraverso l'accesso tramite il Web, consentendo una valorizzazione del patrimonio culturale nazionale e utilizzando i motori di ricerca per una ricerca veloce ed efficace. Il concetto innovativo di biblioteca digitale focalizza il supporto alla creazione di conoscenza come il valore aggiunto che la biblioteca digitale può dare. In questo caso le funzionalità che vengono evidenziate come essenziali sono quelle di supporto all'apprendimento e quelle che consentono la creazione di un ambiente partecipativo e collaborativo per la condivisione della conoscenza. Nel concetto più innovativo inoltre, si focalizza l'aggregazione di collezioni di documenti multimediali distribuiti, dati sensibili, informazione mobile, servizi di elaborazione pervasivi piuttosto che testi digitalizzati come immagine e spesso localizzati centralmente, che caratterizza il concetto più tradizionale. Infine per i servizi, il concetto tradizionale indica come principale servizio l'accesso alle collezioni, con la possibilità di scaricare immagini, documenti ed oggetti digitali posseduti da una o più istituzioni

⁸ Gli *Information Spaces* sono un concetto elaborato nel campo del *Computer Supported Cooperative Work (CSCW)* da Snowdon, Churchill, e Freon, che hanno sviluppato visioni future come 'Connected Communities' e 'Inhabited Information Spaces', dove il secondo è strettamente correlato alla visione di biblioteche digitali. In particolare, questi sono: «spaces and places where people and digital data can meet in fruitful exchange, i.e., they are effective social workspaces where digital information can be created, explored, manipulated and exchanged». Cfr.: D.N. SNOWDON, E.F. CHURCHILL, E. FRÉCON, *Inhabited Information Spaces – Living with your Data*, London, Springer-Verlag 2005.

culturali, insieme a nuovi servizi come i servizi di reference e l'integrazione con servizi di e-learning; il concetto più innovativo prevede possibilità di 'letture' e integrazioni a diversi livelli delle collezioni digitali accessibili in Rete.

L'elemento su cui è rimasto un sostanziale disaccordo tra i due concetti riguarda la necessità di un intermediario (il personale specializzato), che la maggioranza degli esperti del Gruppo (provenienti soprattutto dall'ambito delle biblioteche) ha voluto introdurre nella definizione funzionale di biblioteca digitale che è stata poi utilizzata per l'indagine sui servizi. Il concetto tradizionale di biblioteca digitale prevede infatti la necessaria presenza di un'istituzione che garantisca la gestione e di personale specializzato che è responsabile dei servizi di accesso. Gli argomenti contro l'intermediazione, proposto dai sostenitori del concetto più innovativo, indicavano che il termine personale specializzato sbilanciava la definizione funzionale adottata, in quanto risorse, strumenti e servizi non vengono qualificati, ma ciascuno potrebbe avere diritto ad analoga qualificazione di specializzato. Alla realizzazione di una biblioteca digitale concorrono competenze eterogenee e quindi diverse figure di indirizzo culturale, di progettazione e professionali, come tecnici informatici, bibliotecari, archivisti, editori, esperti nella gestione dell'informazione, specialisti di determinate tematiche o argomenti, inclusi gli utenti. Tra queste, la figura del bibliotecario digitale pare conservare il ruolo tradizionale di intermediario, con ad esempio la funzione di tracciare le strategie di fruibilità delle risorse, la politica dei servizi, la scelta dei contenuti digitali. A questo ruolo tradizionale, nel concetto innovativo di biblioteca digitale, potrebbe aggiungersi un ruolo che potremmo definire di collaboratore, che comprende la funzione di collaborare alla realizzazione ed alla cura scientifica del prodotto digitale e dei metadati di corredo.

Biblioteca digitale o biblioteche digitali? Un altro punto caldo di discussione. La scelta del singolare o del plurale porta ad evidenziare prima di tutto la confusione che sembra esistere attualmente tra i due termini, spesso usati come sinonimo, di collezione digitale e biblioteca digitale. Questa confusione nasce soprattutto in ambito bibliotecario, dove di fatto si tende a chiamare biblioteca digitale quello che più propriamente si dovrebbe chiamare biblioteca ibrida o multimediale. In questo caso, la discussione nell'ambito del Gruppo ha chiarito che il Web e la disponibilità di risorse digitali non organizzate non possono essere confuse con una biblioteca digitale: ad esempio il Progetto Google Book non è una biblioteca digitale. Chiaramente è opportuno

considerare una molteplicità di modelli di biblioteca digitale, prendendo in considerazione diverse modalità di accesso ed organizzazione delle risorse digitali per diverse comunità di utenti. Le applicazioni di biblioteche digitali che il Progetto della Fondazione Rinascimento Digitale ha analizzato sono quelle realizzate dalle istituzioni culturali e queste, secondo la definizione più tradizionale del concetto che è stata adottata per l'indagine, sono raccolte organizzate per il servizio, sia con accesso locale che remoto. Questo modello può contenere materiale digitalizzato, quali surrogati digitali di libri, giornali, fotografie, disegni, mappe, oggetti museali e altro materiale posseduto dalle istituzioni culturali, oppure materiale prodotto direttamente in formato digitale, come avviene sempre più di frequente nel settore della comunicazione scientifica. Altro modello può essere rappresentato da una biblioteca digitale intesa come la biblioteca digitale di un Paese che, come indicato dalla Commissione Europea, nasce da un programma nazionale di digitalizzazione dei contenuti culturali e scientifici presenti in tutte le istituzioni culturali, al fine di sostenere e promuovere l'identità culturale del Paese e produrre un impatto positivo sull'istruzione, il turismo e l'industria. Per la costruzione di questa biblioteca digitale nazionale si devono definire, a livello centrale ministeriale, gli aspetti organizzativi, giuridici, tecnici e scientifici. La biblioteca digitale di una nazione, come in Italia la Biblioteca Digitale Italiana, si integra ed include le biblioteche nazionali e le biblioteche delle altre istituzioni culturali.

Un altro aspetto importante che andava analizzato è stato la cooperazione tra biblioteche digitali. La cooperazione in ambito digitale non è da considerarsi opzionale, come nelle biblioteche tradizionali, ma una scelta strategica per la gestione di una o più funzionalità della biblioteca digitale stessa. La cooperazione dovrebbe essere legata alla biblioteca digitale fin dalla definizione del concetto: cioè non possiamo definire come biblioteca digitale un servizio che non si basi su un'ampia base cooperativa. Anche su questo aspetto i pareri sono stati contrastanti. Alcuni hanno voluto puntualizzare la diversità e specificità delle singole biblioteche digitali, che rendono necessario un punto di accesso locale per la soddisfazione prioritaria di un'utenza strettamente locale. Su questo aspetto della cooperazione il Gruppo *Digital Libraries Application* ha dovuto focalizzare la sua attività di analisi per chiarire i diversi punti di vista che andranno necessariamente armonizzati al fine di garantire l'interoperabilità. Il criterio alla base di ogni scelta cooperativa dovrà essere la convenienza

per l'utente. Anche una semplice federazione potrebbe far ottenere dei vantaggi, come aggregare delle collezioni digitali distribuite con un'unica interfaccia di ricerca. La cooperazione è maggiormente necessaria tra istituzioni culturali affini, per costruire una massa critica di contenuti e per estendere il numero degli utenti⁹. Questi sono infatti gli obiettivi che l'Unione Europea persegue in ambito digitale quando stimola la convergenza di archivi, biblioteche e musei (in sigla ALM).

Infine gli esperti del Gruppo di studio si sono domandati: Quale impatto la biblioteca digitale potrà avere sulle biblioteche, archivi e musei come spazi fisici? La risposta a questa domanda non è stata possibile. È tutto da esplorare il rapporto tra istituzioni culturali reali e la loro rappresentazione digitale. Quello che è chiaro è che sono diverse, ma non è ancora chiaro l'impatto reciproco. La convivenza nella stessa struttura, ad esempio della biblioteca tradizionale e della sua rappresentazione digitale, crea spesso competizione per le risorse disponibili, incluso lo spazio fisico, di cui per altro la biblioteca digitale continua ad avere bisogno. C'è da evidenziare tuttavia anche un possibile impatto positivo: la biblioteca digitale può promuovere l'attività delle istituzioni culturali come luoghi fisici ed incentivare l'utenza alla conoscenza e all'uso del patrimonio tradizionale. Nell'ambito della biblioteca digitale, si ha la possibilità, in alternativa al modello fino ad oggi proposto di fruizione diretta da parte del pubblico del 'bene culturale', cioè della visita in biblioteca, dell'utilizzazione alternativa del modello proposto invece dall'editoria, cioè di una fruizione indiretta dell'opera, tramite la sua rappresentazione digitale attraverso la mediazione del 'prodotto culturale' da parte di un servizio di biblioteca digitale, che, pur essendo sempre un'istituzione culturale, non ha più le caratteristiche di unicità e originalità possedute dal 'bene culturale', poiché è un oggetto appositamente pensato e realizzato industrialmente per il mercato della cultura. Le istituzioni culturali hanno una grande opportunità di valorizzare i propri contenuti culturali digitali attraverso la diffusione in Rete. Possiamo notare infatti i diversi effetti avuti sul bene culturale originale da parte dei due diversi modelli di fruizione:

⁹ Cfr. B. USHERWOOD, K. WILSON, J. BRYSON, *Relevant Repositories of Public Knowledge? Libraries, Museums and Archives in the Information Age*, in «Journal of Librarianship and Information Science», XXXVI, 2, 2005, pp. 89-98.

- l'uno, quello tradizionale basato sul modello 'turistico' di visita, concentrandosi solo sullo sfruttamento del 'bene culturale' senza creare alcunché, non avrà nel migliore dei casi, altro risultato che il lasciare inalterato il bene sfruttato, ma potrebbe anche, in casi più sfortunati, danneggiarlo sia materialmente che nel valore intrinseco perché troppo esposto in pubblico;
- l'altro, quello basato sul modello editoriale, prevedendo in ogni caso una rielaborazione intellettuale del bene per ottenere il 'prodotto culturale', non avrà nel peggiore dei casi altro risultato che lasciare inalterato il bene potendo aggiungere al giacimento, nel caso in cui il prodotto si riveli un'opera di qualità, un nuovo bene che finirà per aumentarne il valore.

Secondo Granelli, una quota sempre più crescente di scambi economici nella loro forma più innovativa sarà riferibile alla commercializzazione di esperienze culturali, più che di beni e servizi prodotti industrialmente¹⁰. Una parte importante del potere economico sarà nelle mani degli intermediari culturali, veri controllori di un patrimonio immateriale:

Saranno coloro che daranno senso a prodotti ed esperienze – sempre più bisognose di essere riempite di senso per avere valore. L'eccesso di informazione tipico dell'era digitale aumenterà il rumore di fondo e la conseguente esigenza di avere filtri autorevoli che guidino la nostra attenzione verso i temi rilevanti.

4. *Organizzazione del flusso di lavoro*

Nella discussione del Gruppo di studio, un altro argomento che è stato a lungo discusso è stato quello dell'infrastruttura necessaria per la biblioteca digitale e di chi debba realizzarla. Il Gruppo di studio ha adottato la definizione di infrastruttura di Borgman¹¹, intesa come

¹⁰ Cfr. A. GRANELLI, *Tecnologie e patrimonio culturale: riflessioni per una via italiana all'innovazione*, in «Sociologia. Rivista quadrimestrale di Scienze storiche e sociali», XXXIX, 2, 2005, pp. 149-155.

¹¹ Cfr. C.L. BORGMAN, *From Gutenberg to the Global Information Infrastructure (GII): Access to Information in the Networked World*, Cambridge MA, MIT Press 2000.

combinazione di persone (o organizzazioni), tecnologie, reti, contenuti digitali, ed ha evidenziato che l'infrastruttura necessaria debba essere il risultato della sinergia di più livelli di attività, da quelle di indirizzo, che coinvolgono i politici e gli amministratori delle istituzioni culturali, a quelle più tecniche che coinvolgono gli informatici ed i professionisti coinvolti nelle biblioteche digitali.

Di fronte alle grandi opportunità ed alla promessa di nuovi servizi di accesso che prospettano alcune tecnologie disponibili, come quelle della facilità di memorizzazione e di recupero veloce di grandi quantità di risorse digitali attraverso Internet, si oppongono alcuni ostacoli legati alla 'governance' del sistema delle biblioteche digitali ed altri connessi alla riorganizzazione del flusso di lavoro: tutti questi aspetti sono strettamente connessi. Su questi aspetti non c'è attualmente chiarezza, anche perché i vari interessati (o 'stakeholders') hanno divergenze di orientamento, in alcuni casi, ad esempio tra editori commerciali ed utenti, hanno interessi a volte opposti. In particolare, la situazione delle biblioteche digitali rivela uno scenario frammentato in cui prevalgono progetti isolati e non coordinati.

Il quadro di riferimento a cui ci si dovrebbe rapportare per la costruzione della biblioteca digitale è invece quello profondamente innovativo del contesto della Rete, o 'networking'. Letteralmente 'networking' significa lavorare in Rete, con un'enfasi che non è nella Rete, come tecnologia e canale di comunicazione, ma nel lavorare o meglio collaborare insieme. La grande opportunità delle nuove tecnologie dell'informazione e della comunicazione va individuata soprattutto nel facilitare la cooperazione tra istituzioni e, all'interno della stessa istituzione, l'integrazione di uffici o settori ora separati. Si può anche dire che la mera applicazione delle nuove tecnologie senza che si attui una diversa organizzazione e con una prospettiva di automazione avanzata ma isolata, vanifica i vantaggi che l'utente potrebbe avere dalla biblioteca digitale. Tralasciamo in questo articolo gli aspetti di 'governance' del sistema, per descrivere l'organizzazione interna delle biblioteche digitali in Italia, come è risultata dall'indagine svolta.

4.1. *Selezione della collezione digitale*

Lo scopo e gli obiettivi che si pongono le diverse istituzioni culturali per costruire una collezione digitale sono strettamente legate alla *mission* dell'istituzione ed all'utente di riferimento per il servizio. Sono quindi da considerare diversi approcci, che dipendono dalle diverse finalità istituzionali. Come conseguenza del cambiamento organizzativo, peraltro ancora in corso e non basato su una precisa strategia,

l'indagine dello staff ha rivelato che le collezioni digitali sono spesso realizzate nell'ambito di un progetto scritto, che ha la funzione di documento di pianificazione e di gestione del budget, in cui vengono decisi i tempi di realizzazione con i diversi stati di avanzamento. Nel caso di collezioni di risorse digitali 'born digital' è stato più difficile evidenziare un documento strategico che definisca la politica di sviluppo della collezione; tuttavia la negoziazione delle licenze è stata accentrata in un apposito ufficio.

Lo scopo che generalmente tutte le istituzioni culturali perseguono nella creazione di una collezione digitale è quello di migliorare l'accesso alle risorse, sia che queste facciano parte della propria collezione analogica (nel caso di progetti di conversione da analogico a digitale o di surrogati) o che vengano selezionate ed acquisite, anche attraverso licenze di accesso (nel caso di risorse originariamente digitali o born digital). Nell'ambito del generico scopo di migliorare l'accesso, sono stati elencati diversi obiettivi, che possono essere combinati o uno prevalere sugli altri, come:

- una migliore conservazione degli originali, in particolare quelli rari e fragili o il loro restauro virtuale;
- la creazione di particolari percorsi tematici, predisposti per utenti specifici o aggregando tipologie diverse di risorse, ad esempio per mostre virtuali permanenti;
- la realizzazione di un progetto editoriale, come ad esempio la riproduzione digitale del fondo locale, di una collezione di libri rari, o la creazione di pubblicazioni originariamente digitali, come nel caso di alcune università;
- l'estensione al digitale di alcuni servizi tradizionali, come ad esempio il collegamento tra il catalogo ed il testo pieno o tra il catalogo e le immagini degli indici dei volumi e dei periodici.

Un ulteriore scopo, da aggiungere a quelli elencati, ma spesso taciuto nei risultati dell'indagine, è la possibilità di ri-uso della collezione digitale. Bisogna infatti essere consapevoli della potenzialità della collezione digitale di costituire un elemento di un possibile servizio per utenti diversi da quelli per cui è stata pensata. Questo implica delle problematiche di ri-usabilità, persistenza, interoperabilità, certificazione delle risorse. Un particolare ri-uso è quello da parte dei motori di ricerca del Web, con cui bisognerebbe favorire l'integrazione funzionale.

I criteri di selezione della collezione digitale vengono decisi da vari attori:

- il Comitato Guida della Biblioteca Digitale Italiana;
- gli stessi utenti, nel caso ad esempio della digitalizzazione a domanda o nel caso dell'utenza universitaria, anche se l'avvio dei Consorzi per la negoziazione di licenze ha dato più potere ai tecnici nei Sistemi bibliotecari di Ateneo;
- i responsabili dello sviluppo della collezione digitale, o della biblioteca digitale, previa autorizzazione di Comitati di riferimento o dell'organo di controllo gestionale, se necessario;
- gli enti finanziatori di appositi programmi per la digitalizzazione, a livello locale, nazionale o internazionale.

I criteri di selezione sono quindi diversi, adeguati alle finalità ed agli interessi dei diversi attori coinvolti nella decisione. Possiamo tuttavia riconoscere che c'è una spinta a digitalizzare che è comune a tutte le istituzioni culturali che sono state analizzate ed è quella della disponibilità del finanziamento (pur se spesso limitato alla sola acquisizione della risorsa e non alla sua gestione e preservazione). Un altro criterio di selezione riguarda il monitoraggio dell'effettivo utilizzo, ad esempio diffuso nel caso di risorse digitali in licenza di accesso, o nel caso di volumi a stampa molto richiesti che si decide di digitalizzare. Infine, un criterio ovvio a molte delle iniziative di digitalizzazione è rappresentato dall'evitare i vincoli posti dal copyright, selezionando in preferenza le collezioni di cui si possiedono tutti i diritti (o si ha autorizzazione da chi li possiede) oppure le così dette opere orfane (cioè quelle senza un detentore di diritti di proprietà intellettuale).

Altri criteri non sono stati elencati nei questionari ricevuti, probabilmente perché considerati come scontati. Ad esempio nessuno ha affermato di controllare prima di selezionare un documento da digitalizzare che questo non sia già stato convertito in digitale. Con il diffondersi delle attività di digitalizzazione è importante evitare delle duplicazioni. Il criterio di non duplicare è spesso implicito nelle esperienze in corso in Italia, ma è evidente che quando si seleziona del materiale si dovrebbe essere certi che non ci siano progetti simili o complementari già realizzati o in corso. È da notare che per un'efficace attività di selezione delle risorse digitali, mancano degli strumenti di riferimento a cominciare da un registro autorevole che contenga tutte le collezioni digitali di qualità, costantemente mantenuto. Finora questo strumento mancava; attualmente, con l'avvio

del Progetto MICHAEL¹², si potrà contare su questo indispensabile supporto.

L'attività di acquisizione e conversione della collezione digitale è quella che più spesso viene affidata all'esterno, completamente o in parte, a Centri di competenza tecnica per l'OCR. Poche sono le biblioteche digitali che hanno creato al loro interno un Centro di digitalizzazione o che usufruiscono di servizi comuni di conversione da analogico a digitale. Un buon esempio è stato, a partire degli anni '50, l'istituzione del Centro di Fotoriproduzione, Legatoria e Restauro degli Archivi di Stato¹³: questo ha rappresentato un punto di partenza molto utile per la digitalizzazione degli Archivi di Stato, anche perché la finalità originaria di creare un microfilm è analoga a quella di creare un oggetto digitale, qualcosa che può essere usato al posto di un originale. Si consideri che nel contesto dei musei il fine della riproduzione è sempre stato viceversa semplicemente quello di identificare l'originale (dipinti, sculture, edifici), arricchendo la scheda di catalogo di una immagine fotografica. È per la digitalizzazione di questi materiali fotografici che l'Istituto Centrale per il Catalogo e la Documentazione (ICCD) ha adottato degli standard¹⁴, ma vale la pena di sottolineare il fatto che questi riguardano la digitalizzazione di una ripresa fotografica dell'originale, non dell'oggetto originale in sé. Bisogna inoltre considerare esigenze diverse di risoluzione degli originali: da rappresentazioni digitali che trattano originali di piccolo formato, scansionate a risoluzioni medie, all'accesso remoto a documenti di maggiore formato e a risoluzione più alta che richiedono sforzi maggiori e tecnologie avanzate, che comprendono anche sistemi georeferenziati (GIS).

Le tipologie di materiale che vengono selezionate per la collezione digitale comprendono ogni possibile oggetto digitale, con una prevalenza attualmente di testi, sia libri che periodici, che hanno generalmente una copia a stampa corrispondente. Degna di nota una tendenza evidenziata nell'indagine che sta portando le singole istituzioni culturali ad allontanarsi da questo modello che potremmo chiamare

¹² MICHAEL <http://www.michael-culture.org/copyright_i.html>.

¹³ Centro di Fotoriproduzione, Legatoria e Restauro <<http://www.cflr.beniculturali.it/>>.

¹⁴ ICCD, 1998. *Normativa per l'acquisizione digitale delle immagini fotografiche*, ICCD <<http://www.iccd.beniculturali.it/>>, <<http://www.iccd.beniculturali.it/download/fo-todig.pdf>>.

di estensione del servizio tradizionale, verso un vero e proprio sistema informativo che integra diverse banche dati ed oggetti digitali multi-mediali e non testuali, e quindi rende possibile utilizzi diversi ed innovativi delle risorse digitali. Si prospetta un cambiamento importante ed una rivitalizzazione del ruolo delle istituzioni culturali nel fare cultura. In questo caso, come evidenzia Vagaggini¹⁵ occorre un cambio paradigmatico nella comunicazione svolta dalla biblioteca digitale:

[...] nel momento in cui il testo digitalizzato si confronta con la compresenza di altri testi, immagini, suoni, links di approfondimento, questo perde la sua caratteristica principale, la sequenza sintagmatica dei segni – le parole - a favore di una lettura paradigmatica. Ecco quindi che altre regole linguistiche e narrative si impongono, avvicinandosi al linguaggio proprio dell'immagine, da sempre dominio del settore mediatecale.

4.2. *Gestione della collezione digitale*

Quello che finora per molte istituzioni culturali come ad esempio le biblioteche, era l'iter del libro o del documento, è stato completamente trasformato dal nuovo ambiente digitale, in un flusso che non è più come una catena di montaggio, in cui ogni fase precede la successiva. La caratteristica essenziale del nuovo flusso di lavoro è che esso ha un tipico andamento elicoidale, in cui in ogni fase si deve tener conto di tutte le altre: nella fase di creazione dei contenuti, si deve sapere come l'oggetto culturale sarà presentato ed accessibile all'utente, ed inoltre si dovrà immediatamente decidere quale tipo di preservazione verrà assicurata all'oggetto. Da questa prima sommaria descrizione di una delle caratteristiche della realizzazione della biblioteca digitale, ne consegue un primo importante corollario: l'organizzazione all'interno di un'istituzione non potrà più essere quella pensata per l'iter del libro, in cui questo passa da un ufficio ad un altro, con poche o nessuna interazione tra gli addetti. In un'ottica di Rete inoltre, l'organizzazione coinvolgerà altre istituzioni, in quanto non possono esistere biblioteche digitali isolate. Un progetto di costruzione della biblioteca digitale ha il compito primario di definire gli aspetti di innovazione organizzativa, allargata alla cooperazione trasversale tra istituzioni appartenenti a settori affini o diversi, per garantire un servizio integrato all'utente. Anche in questa apertura al networking, elementi essen-

¹⁵ Cft. E. VAGAGGINI, *Mediateche*: <<http://www.rinascimento-digitale.it>>.

zialmente tecnici dovranno essere combinati con accordi organizzativi e collaborativi, insieme ad un governo complessivo del sistema a livello locale o nazionale. Una nuova organizzazione è legata quindi all'architettura dell'informazione in Rete, con una corretta progettazione di servizi condivisi e di servizi distribuiti.

La gestione della collezione digitale include un insieme di attività tecniche previste nel flusso di lavoro della collezione digitale, integrate dalla disponibilità di un'infrastruttura tecnologica complessa, che è necessaria per portare a compimento le attività stesse. Sono permeate dalle tecnologie, in vari aspetti, le fasi previste per la creazione dei contenuti, per la memorizzazione e gestione delle collezioni e per l'accesso a queste. Queste fasi riguardano la predisposizione di contenuti con appositi formati standard e l'accesso a questi contenuti culturali digitali attraverso i protocolli della Rete e del Web. Le fasi di progettazione e valutazione di un programma di biblioteca digitale insieme a quella dell'inserimento dei metadata, tutt'altro che marginali nel flusso di lavoro, rappresentano l'infrastruttura organizzativa e tecnica portante ed hanno tra gli obiettivi previsti anche quello di razionalizzare le risorse e migliorare l'organizzazione interna dell'istituzione, eliminando barriere ed ostacoli dovuti ad organizzazioni stratificate e non più funzionali.

Le istituzioni culturali hanno risposto alle nuove esigenze della gestione della collezione digitale con modalità organizzative diverse, che tuttavia hanno la caratteristica comune di evidenziare una tendenza alla cooperazione ed all'esternalizzazione. Le forme organizzative scelte sono:

- un'organizzazione separata all'interno della stessa istituzione;
- una maggiore delega di particolari attività e servizi ai privati;
- una spinta a forme di cooperazione e collaborazione all'interno ed all'esterno dell'istituzione.

Lo staff responsabile delle biblioteche digitali è di solito lo staff dell'istituzione, che riceve un incarico suppletivo, senza una particolare formazione. Nell'ambito di questa organizzazione, separata o integrata dal flusso normale delle attività, lo sviluppo della collezione digitale è sempre un lavoro aggiuntivo, per il quale ci si rivolge spesso ad esterni, organizzazioni private o istituzioni pubbliche, nel caso che non ci siano adeguate competenze, come ad esempio quelle che riguardano le tecnologie, o per attività e servizi che non sia possibile realizzare internamente, come ad esempio la conversione da formato

analogico a digitale. Anche l'attività di memorizzazione e di preservazione di medio e di lungo periodo non è gestita dalle biblioteche digitali che sono state analizzate, tranne eccezioni, ma affidata all'esterno a ditte o gestita in modo cooperativo nei Sistemi, nelle Reti e nei Consorzi. Raramente la problematica è considerata con la giusta consapevolezza: si può citare come buona pratica il Gestionale della Biblioteca digitale dell'Istituto e Museo di Storia della Scienza, oppure la realizzazione dei metadati MAG per gestire la preservazione degli oggetti digitali nell'ambito del Progetto di Biblioteca Digitale Italiana. La preservazione delle risorse digitali 'born digital' viene spesso affidata alla responsabilità dello stesso produttore della risorsa, come nel caso degli editori commerciali. In alcuni casi viene gestita all'interno della stessa istituzione che ha creato la risorsa digitale, ad esempio con la realizzazione dei depositi istituzionali. Sicuramente possiamo evidenziare la necessità di interventi per l'armonizzazione (volontaria o regolata) di una politica nazionale per la preservazione, anche con la realizzazione di *trusted repository* e con una maggiore collaborazione tra istituzioni per l'applicazione della Legge sul deposito legale. In particolare sembrano necessari degli interventi per la preservazione di metadata gestionali e formati aperti, la gestione degli IPR o diritti di proprietà intellettuale, oltre che di tecnologie per l'accesso. Il problema della preservazione è centrale per la sostenibilità della biblioteca digitale.

Si può dire che la sostenibilità della biblioteca digitale nel lungo periodo di solito è una problematica messa in fondo alla lista delle priorità delle biblioteche digitali italiane. Una problematica sottovalutata è che, per produrre un servizio, oppure anche per realizzare dei contenuti digitali, occorrono investimenti aggiuntivi e si corrono certi rischi. Tuttavia, anche se le risorse digitali ed i servizi basati su di esse sono di qualità, non potranno sopravvivere a lungo quando i finanziamenti del progetto saranno finiti. Costruire una biblioteca digitale è costoso. Inoltre, quando si parte con un nuovo programma di digitalizzazione, occorrono spese di investimento iniziale, che poi non si sa se potranno essere riusate o adattate per programmi successivi.

4.2.1. Cooperazione

Un'attenzione particolare è stata dedicata dal Gruppo di studio alle problematiche di coordinamento e del necessario equilibrio in ambito digitale tra centralizzazione e decentramento. La cooperazione è necessaria per perseguire obiettivi come l'interoperabilità, la condivisione dei costi, la preservazione, l'attenzione a non duplicare. Cosa

conviene allora centralizzare, perché è più efficiente per l'amministrazione e contemporaneamente più efficace per l'utente? Ad esempio, il Manifesto sulle Biblioteche Digitali dell'AIB¹⁶ indica le seguenti attività collaborative: dalla più semplice che è rappresentata dalla condivisione di linee guida e standard comuni, ad una progressione di attività sempre più complesse da organizzare, come l'adozione comune di software *open source* e standard aperti, e servizi come il Portale nazionale delle biblioteche digitali e la preservazione (in base alla recente legge sul deposito legale).

Quali sono le prime attività svolte in modo cooperativo dalle biblioteche digitali in Italia? A livello nazionale, possiamo indicare le attività del Progetto della Biblioteca Digitale Italiana ed in particolare Internet culturale che è il portale che si propone come punto di accesso per i vari settori dei beni culturali, compreso musica e cinema. Il portale è un punto di accesso non solo per le istituzioni culturali ma anche per altri istituti (come l'Istituto Luce ed il Touring Club) e consente servizi aggiuntivi, come l'e-commerce. Per la condivisione di linee guida e l'adesione a standard, il Progetto europeo MINERVA ha offerto la disponibilità di una serie di linee guida e di buone pratiche che, oltre ad essere primi esempi di strumenti cooperativi, sono particolarmente utili per l'orientamento del personale. Il Progetto Minerva ha inoltre realizzato una struttura di coordinamento sugli standard tecnici, chiamata OTEBAC¹⁷.

Nell'attuale stato dell'arte delle biblioteche digitali in Italia, si pos-

¹⁶ Manifesto AIB Manifesto per le Biblioteche digitali, presentato a Ravenna nel 2006 e consultabile a: <<http://www.aib.it/aib/cg/gbdigd05a.htm3>>. Il documento è diviso in tre parti: Principi, Modelli e Funzioni. Nella sezione dei Principi è si delinea un modello di biblioteca digitale articolato in strutture decentrate e autonome, costituite da servizi di mediazione per l'accesso alla conoscenza, con un modello di governance di tipo locale, basato sul coordinamento istituzionale e la condivisione di obiettivi e metodi con la comunità del territorio. A questa visione fa da sponda un atteggiamento di diffidenza verso il centralismo e soluzioni calate dall'alto. Altri principi forse più ovvi, sono che le biblioteche digitali sono in Rete e sono accessibili a tutti – il riferimento è al digital divide – e utilizzano standard riconosciuti e condivisi.

¹⁷ OTEBAC, l'Osservatorio tecnologico per i beni e le attività culturali, è un centro del Ministero per i Beni e le Attività culturali, per monitorare e sostenere gli istituti culturali nella realizzazione dei siti web e per una corretta digitalizzazione e accessibilità dei contenuti: <<http://www.otebac.it/>>.

sono indicare altri progetti cooperativi nazionali, come i progetti dei Consorzi in ambito universitario. Questi nati per la negoziazione di licenze ora estendono la cooperazione ai servizi di accesso, come ad esempio la costruzione cooperativa della Base di conoscenza di MetaLib, per l'aggregazione anche semantica di diverse banche dati e periodici elettronici. A livello delle biblioteche pubbliche, i Sistemi bibliotecari concentrano le loro risorse per servizi comuni. Ad esempio possiamo citare le esperienze cooperative delle biblioteche di ente locale, elencate da Bertini¹⁸. Un intero capitolo è dedicato alla cooperazione nell'ultima indagine sulle biblioteche italiane¹⁹ in particolare per le biblioteche pubbliche collegate in Rete, con esempi e studi di caso.

Malgrado gli indubbi vantaggi, tuttavia non è semplice realizzare dei programmi cooperativi, per un insieme di ragioni di ambito socio-economico. Appare evidente che la motivazione a cooperare ha bisogno di particolari incentivi e stimoli. Dobbiamo distinguere inoltre i due casi che appaiono evidentemente diversi, cioè della cooperazione tra istituzioni culturali pubbliche (in particolare archivi, biblioteche e musei) e della cooperazione tra istituzioni culturali pubbliche con organizzazioni o aziende private.

Nella cooperazione tra istituzioni culturali, il fenomeno della convergenza (che in sigla viene chiamato ALM Archivi, Biblioteche, Musei) è particolarmente vantaggioso per l'utente, che si trova facilitato ad esempio da ricerche integrate di più collezioni digitali. L'attività di catalogazione è quella in cui si possono evidenziare i risultati migliori ottenuti per la cooperazione tra archivi e biblioteche, con il coordinamento dell'ICCU e nell'ambito del progetto Biblioteca Digitale Italiana. Il fenomeno della convergenza sembra tuttavia voluto dai politici e dai responsabili amministrativi, ma ancora non apprezzato dai professionisti dei diversi settori. Gli esempi di effettiva convergenza che il Gruppo di studio ha rilevato in altre nazioni hanno avuto un discreto successo in caso di cooperazione diretta da uno spe-

¹⁸ Il Sistema bibliotecario degli enti locali organizza servizi comuni per il prestito interbibliotecario, il reference, il virtual reference desk, le linee guida, i servizi di document delivery. Cfr. V. BERTINI, *Biblioteche pubbliche*: <<http://www.rinascimento-digitale.it>>.

¹⁹ Cfr. V. PONZANI, G. SOLIMINE, *Rapporto sulle biblioteche italiane 2005-2006*, 2006, AIB <<http://www.aib.it/>>, <<http://www.aib.it/aib/editoria/2006/pub144.htm>>.

cifico finanziamento o in caso di cooperazione forzata, con la riunione sotto lo stesso ufficio ministeriale di archivi, biblioteche e musei. È inoltre apparso evidente che la cooperazione in ALM è una forma politica di comunicazione, in cui è necessaria sia una spinta dall'alto (ad esempio con opportuni finanziamenti) sia una condivisione dei valori della cooperazione dal basso. La comunicazione si rivela quindi come uno strumento strategico per la migliore cooperazione, facilitando lo scambio tra esigenze generali ed esigenze specifiche. Il finanziamento in ogni caso sarà necessario, perché uno degli ostacoli che è stato rilevato nell'indagine è rappresentato dalla scarsità di risorse delle istituzioni culturali che, in mancanza di finanziamenti aggiuntivi, ritengono di non poter distogliere risorse dalle attività istituzionali ritenute prioritarie.

Recentemente, per una migliore cooperazione si è proposto come organizzazione cooperativa distribuita il distretto culturale digitale. Il 'distretto culturale digitale' è un territorio in cui le istituzioni culturali possono essere al centro di una rete fatta dalle istituzioni locali e dalle imprese tecnologiche, dai servizi di accoglienza turistica e dal dipartimento di ricerca dell'università, dagli artigiani, dai progettisti di software e dai creatori di contenuti. Per Granelli²⁰ lo Stato dovrebbe guidare questi processi, dovrebbe usare la domanda pubblica per modulare lo sviluppo, dovrebbe creare aggregazione sui progetti, far nascere un tessuto che li tiene insieme secondo un disegno che sarà tanto più lungimirante quanto meno rigido e centralizzato.

4.2.2. *Sostenibilità e rapporto pubblico/privato*

Si intende qui per sostenibilità il prevedere e pianificare tutte le attività necessarie per mantenere il contesto istituzionale per la creazione e gestione della risorsa digitale, incluso la sua preservazione. Il recupero dei costi può essere un modo per ottenere la sostenibilità della biblioteca digitale, ma occorre costruire dei modelli economici²¹. Questo implica tuttavia un cambiamento culturale delle istituzioni pubbliche, che trova qualche resistenza ed è stato uno dei temi su cui si è concentrata la discussione del Gruppo di studio per decidere

²⁰ Cfr. GRANELLI, *Tecnologie e patrimonio culturale* cit.

²¹ CLIR, *Building and Sustaining Digital Collections: Models for Libraries and Museums*, 2001, CLIR <<http://www.clir.org>>, <<http://www.clir.org/pubs/reports/pub100/pub100.pdf>>.

i servizi garantiti a tutti ed i servizi invece a pagamento. Il servizio di base e gratuito è stato indicato nel servizio di accesso e ricerca a testo pieno. I servizi aggiuntivi a pagamento, secondo il Gruppo di studio, potrebbero comprendere servizi aggiuntivi, come ad esempio la stampa e il download nel proprio PC di documenti, immagini o parti della risorsa digitale.

Sia il pubblico che vende servizi o prodotti, sia il privato che intenda fornire servizi basati sui contenuti culturali digitali devono fare un investimento iniziale, che potrà essere necessario per i primi 3-4 anni prima che il servizio diventi sostenibile²². Nel caso di un investimento pubblico di sostegno all'avvio di progetti di biblioteca digitale, ad esempio per la costruzione di oggetti digitali e di metadati adeguati, i finanziatori pubblici non si dovrebbero aspettare un ritorno dell'investimento diretto, invece dovrebbero preoccuparsi di creare sinergie in alcuni specifici settori (come ad esempio l'e-learning o il turismo culturale), per essere sicuri di migliorare l'impatto del loro investimento in specifici scenari.

È possibile quindi una sinergia tra pubblico e privato? Bergamin, nella risposta all'intervista su questo tema, ha fatto notare che questo problema è ben esemplificato da Google Book e dalla posizione del governo francese: da una parte c'è l'approccio del liberismo e della libera competizione e collaborazione e dall'altra l'organizzazione centralizzata dello Stato, che persegue dei fini di utilità pubblica anche cercando di correggere le iniziative del privato.

Allo stato delle cose, nei pochi casi in cui il privato contribuisce alla realizzazione di una biblioteca digitale, questo interviene esclusivamente come finanziatore, o come fornitore di tecnologie informatiche. Esistono esempi in cui i privati si sono fatti carico di investire per la digitalizzazione di contenuti, come il Progetto DIGITAmi, in cui la Telecom ha intrapreso la digitalizzazione di un fondo della Biblioteca Soriani. Le collaborazioni fin qui avviate, soprattutto in applicazioni tecnologiche, hanno prodotto ottimi risultati, ma con investimenti soprattutto pubblici. Il rapporto più importante tra pubblico e privato sarebbe quello inerente alla creazione di infrastrutture per l'accesso; sarebbe anche desiderabile una sinergia tra ricerca ed applicazione. In

²² *Digicult Report: Technological Landscapes for Tomorrow's Cultural Economy. Unlocking the Value of Cultural Heritage*. Luxembourg: European Commission. Directorate General for the Information Society, 2001.

particolare, sembra prioritario sviluppare una sinergia tra pubblico e privato nel contesto di riferimento della Biblioteca Digitale Europea, al fine di individuare un modello economico della biblioteca digitale che coniughi la garanzia del servizio insieme alla sostenibilità di questo nel tempo. Questo modello, ed i correlati modelli di contratti di collaborazione tra pubblico e privato, attualmente non sono applicati in Italia. L'attuale fase di sviluppo della biblioteca digitale è ancora concentrata nell'organizzazione dell'infrastruttura di base e l'indagine realizzata dal Gruppo di studio evidenzia che va completato l'investimento iniziale.

In modo problematico, si può discutere se una vera apertura agli investimenti ed alla competizione tra più privati possa essere avviata fin da ora, o rinviata ad una fase più matura della biblioteca digitale. Vanno tuttavia considerati alcuni possibili rischi. Una necessaria precisazione è che la scelta del privato non deve essere legata all'incapacità di gestione del pubblico, che quindi delega al privato. L'ottica che qui ci si pone è quella dell'utente dei servizi, che ha il diritto di avere i servizi migliori al costo minore. Il pubblico spende nel modo migliore, il privato nel modo utile. Per l'integrazione corretta tra pubblico e privato occorre usare gli strumenti della sussidiarietà orizzontale, che implicano la necessità di definire con chiarezza i distinti ruoli del pubblico e del privato, nella tutela della libera concorrenza, di favorire la creazione di tavoli di confronto e di concertazione insieme agli strumenti per le decisioni partecipate.

Una ulteriore considerazione è che le istituzioni culturali hanno l'opportunità di rivalutare la loro missione, utilizzando le tecnologie. La tendenza verso l'affidamento all'esterno di certe funzioni, dovrebbe quindi orientare le scelte di cosa delegare ai privati tra quelle funzioni che sono di supporto alle attività istituzionali e non tra quelle che sono di stretta competenza delle finalità stesse. C'è stato il rischio in passato di delegare all'esterno funzioni che hanno impoverito le competenze del capitale umano, come ad esempio quando si sono delegate completamente all'esterno le applicazioni innovative delle tecnologie che si riferivano ad aree di crescita di importanza vitale per l'istituzione.

5. *Utenti ed uso*

Possiamo notare che il monitoraggio dell'uso e la valutazione costante del servizio non sono ancora parte di una prassi abituale e co-

stante nelle esperienze di biblioteca digitale. Si possono citare tuttavia alcune buone pratiche, come le indagini qualitative realizzate dall'Istituto e Museo di Storia della Scienza e dalla Biblioteca Nazionale Centrale di Firenze (su cui il Gruppo di studio si è basato per l'indagine dell'utenza). Le biblioteche universitarie usano misurazioni quantitative, attraverso l'applicazione delle misure di COUNTER, che è lo standard comune di editori e biblioteche per la quantificazione degli accessi. Sulla base di questa considerazione, i risultati dei singoli studi di caso, che vengono qui sintetizzati, vengono proposti solo come uno stimolo ad includere un'attività periodica di valutazione tra le attività gestionali necessarie per la biblioteca digitale.

Nell'indagine sono state considerate le seguenti risorse digitali delle biblioteche digitali: il catalogo in linea, i periodici elettronici, gli e-books, le banche dati, i Cd-Rom, il materiale didattico, gli audiovisivi, le tesi ed i lavori degli studenti. I servizi che sono basati su queste risorse sono: l'accesso remoto, il portale o il sito web, la promozione e l'assistenza del personale, l'addestramento all'uso della collezione ed eventuali tutorial ed aiuto in linea. La discussione del Gruppo di studio ha evidenziato come le funzioni di ricerca e l'accesso remoto potranno essere accompagnate da servizi aggiuntivi, come dare all'utente un proprio spazio nel server, estendere i percorsi tematici di approfondimento ed organizzazione semantica delle collezioni, favorire la possibilità di 'letture' e di interazioni a diversi livelli delle risorse digitali proposte. Tuttavia, dovendo prendere le mosse dalla situazione reale delle biblioteche digitali, la scelta fatta è stata quella di concentrarsi sull'attuale stato dell'arte di risorse e servizi.

5.1. *Chi sono gli utenti delle istituzioni culturali*

Nella metodologia di ricerca che è stata sviluppata, gli utenti sono stati classificati per età, sesso, nazionalità, professione ed attività. Il risultato che si prospetta da questa prima parte dell'indagine è quello dell'importanza della definizione dell'utente per la scelta delle priorità di servizio e la selezione delle risorse digitali. Pur nei limiti già evidenziati della ricerca, possiamo dire che gli utenti delle differenti istituzioni affermano diverse priorità di servizio, alla cui lettura rimandiamo al Rapporto in linea²³.

²³ Negli studi di caso esaminati, sono stati selezionati particolari comunità di utenti, appartenenti all'area umanistica. Nella Biblioteca Umanistica dell'Università di

5.2. Comparazione dei risultati

Pur nella diversità degli utenti, alcune priorità sono comuni e possono essere comparate. Si può affermare che il nuovo utente vuole essere indipendente nella ricerca, che vuole effettuare con accesso remoto: questo è dimostrato dalla generale aspettativa di un buon orientamento attraverso il portale, anche nel caso in cui si visiti abitualmente una biblioteca fisica. I servizi di cui è stata evidenziata la priorità sono quindi il portale e l'accesso remoto alla biblioteca digitale. C'è in genere un atteggiamento positivo per i servizi della biblioteca digitale, ma c'è anche poca capacità di usarli ed una mancanza di consapevolezza del servizio che si potrebbe avere. In modo che forse può sembrare paradossale, il supporto del personale rimane un'esigenza di servizio particolarmente sentita da parte degli utenti. Una possibile spiegazione è che non tutti hanno la capacità di saper usare le tecnologie disponibili; inoltre proprio i più esperti si rendono conto che non utilizzano appieno tutte le possibilità offerte. Questo è confermato dal risultato che evidenzia che gli utenti esaminati non si ritengono soddisfatti della promozione dei servizi, dei corsi che vengono offerti e soprattutto del supporto che hanno dallo staff. Alcuni utenti chiedono maggiore interazione e non solo la comunicazione unidirezionale ora disponibile attraverso i portali.

Una particolare evidenza va data al fatto che la promozione dei servizi è dichiarata come carente da tutti gli utenti. Non è quindi sufficiente la disponibilità di un servizio o di una risorsa, questa deve essere opportunamente diffusa e promossa perché divenga utilizzata e utile.

Firenze gli utenti scelti sono studenti e laureandi, con un'età che va dai 25 ai 40 anni; hanno una conoscenza media di Internet ed usano frequentemente il sito del Sistema bibliotecario di Ateneo. Gli utenti della Biblioteca dell'Istituto e Museo di Storia della Scienza sono professionisti ed impiegati con un titolo di studio post-laurea ed un'età che va dai 32 ai 76 anni; hanno una conoscenza buona o ottima di Internet ed usano molto frequentemente il sito della Biblioteca digitale. Gli utenti della Mediateca sono soprattutto studenti e tra i più giovani, dai 19 ai 25 anni; frequentano settimanalmente la Mediateca, anche da casa. Tutti gli utenti che hanno risposto all'indagine usano anche i servizi di altre biblioteche: gli utenti della Biblioteca dell'Istituto e Museo di Storia della Scienza e gli utenti della Mediateca accedono in modo remoto ad altre istituzioni nazionali ed internazionali con la stessa specializzazione; gli studenti dell'Università usano soprattutto i servizi locali del Sistema delle biblioteche pubbliche fiorentine. Cfr. TAMMARO, CASATI, LUZZI, *Biblioteche digitali in Italia* cit.

Le banche dati ed il catalogo in linea sono un'area in cui prestare attenzione particolare, per andare incontro alle priorità degli utenti. Ripetendo per le risorse digitali l'esame delle priorità correlate con la soddisfazione per le stesse risorse, si è potuto individuare che le risorse ritenute prioritarie sono:

- il catalogo OPAC;
- le banche dati in linea;
- i periodici elettronici.

Un servizio su cui gli utenti pongono particolare importanza riguarda quindi il recupero delle informazioni. Pur evidenziando che ci sono differenze tra gli utenti, si può notare che gli strumenti bibliografici mantengono il ruolo di strumenti di accesso, anche nell'epoca dei motori di ricerca. Nei suggerimenti ricevuti sembra particolarmente rilevante evidenziare la richiesta di un miglior funzionamento dell'OPAC. La ricerca e localizzazione delle risorse digitali si vorrebbe veloce e facile sul modello dei motori di ricerca.

L'impatto ottenuto dai servizi della biblioteca digitale che è stato evidenziato dagli utenti riguarda sostanzialmente alcuni vantaggi come: la velocità di recupero e scarico delle risorse digitali insieme ad un maggior numero di risorse accessibili. Tuttavia le risorse disponibili ancora non sono ritenute sufficienti e la gran parte degli utenti intervistati ne ha chiesto un'estensione. Altre richieste hanno riguardato la possibilità di ricerca integrata tra le diverse risorse digitali disponibili, con ad esempio un collegamento dall'OPAC alla disponibilità di anteprima della copertina, del frontespizio, dell'indice. Anche la personalizzazione del servizio è ritenuta importante: in questo caso sono state chieste funzionalità diverse dal recupero dell'informazione, come la navigazione tra diverse biblioteche digitali, uno spazio virtuale per collaborazioni e per gestire una collezione digitale personale, l'apertura a collegamenti con altre istituzioni.

Possiamo inoltre notare che gli utenti usano abitualmente i servizi di più di una istituzione culturale, anche se sono fisicamente collegati da una stazione di lavoro della biblioteca di appartenenza. Alla domanda sul bisogno di maggiore coordinamento tra le istituzioni culturali, l'indicazione è stata che questo è sicuramente insufficiente. I desiderata comprendono l'integrazione di collezioni di istituzioni culturali diverse con una modalità unificata di accesso.

6. Conclusioni

In conclusione possiamo evidenziare che occorrerà concentrare maggiori sforzi per migliorare l'accesso dell'utente alle biblioteche digitali, usando le opportunità esistenti per una maggiore interoperabilità ed interattività dei servizi di accesso. Questo comporta un'applicazione mirata sia di tecnologie sia di accordi cooperativi ed organizzativi per aggregazioni di collezioni e di servizi di istituzioni culturali diverse.

Un primo risultato che pare possa essere evidenziato è che nella costruzione della biblioteca digitale si è dato in Italia molta importanza all'infrastruttura tecnologica ed ad altri aspetti tecnici, legati alla produzione di contenuti e di metadata. Questo è sicuramente corretto ma, dal punto di vista degli utenti, deve essere accompagnata da azioni di sostegno sia al cambiamento dell'organizzazione delle istituzioni culturali, sia ad una attività di educazione e di promozione dei servizi agli utenti, sia da un supporto politico e legislativo fornito dagli organi politici competenti.

L'indagine ha rivelato che esistono alcune aree problematiche che riguardano:

- la necessità di approfondire gli ostacoli legislativi e culturali al libero accesso dei contenuti digitali;
- la promozione della cooperazione e dei possibili vantaggi per gli utenti dalla collaborazione intersettoriale tra archivi, biblioteche e musei;
- un approfondimento di alcuni aspetti della sinergia possibile tra pubblico e privato;
- il bisogno di architetture e sistemi di governance delle biblioteche digitali, con comprensione della dialettica centralizzazione/decentramento, soprattutto per la preservazione;
- la necessità della formazione dello staff;
- la diffusione di un insieme di indicatori per la valutazione continua delle biblioteche digitali.

La *mission* della biblioteca digitale è stata indicata dal Gruppo di studio nel permettere agli utenti l'accesso alla conoscenza nel rispetto dei diritti di tutti – con il richiamo ad un equilibrio tra interessi, spesso contrastanti. Per ottenere questo, occorre stimolare una cultura dell'eccellenza per i servizi, realmente orientata agli utenti che possono assumere un ruolo centrale nel progetto e nella realizzazione delle biblioteche digitali.

7. Riferimenti bibliografici

V. BERTINI, *Biblioteche pubbliche*: <<http://www.rinascimento-digitale.it>>.

J.C. BLIXRUD, *Measures for Electronic Use: The Arl E-Metrics Project*. IFLA Satellite Conference *Statistics in practice - Measuring & Managing*, Loughborough, UK, 13-25 august, 2002, Loughborough University <<http://www.lboro.ac.uk/>>, <<http://www.lboro.ac.uk/departments/dils/lisu/downloads/statsin-practice-pdfs/blixrud.pdf>>.

J.C. BLIXRUD, *Assessing Library Performance: New Measures, Methods, and Models*. IATUL Proceedings *Libraries and Education in the Networked Information Environment*, Ankara, Turkey, 2-5 june, 2003, IATUL <<http://www.iatul.org>>, <http://www.iatul.org/conference/proceedings/vol13/papers/BLIXRUD_fulltext.pdf>.

C.L. BORGMAN, *From Gutenberg to the Global Information Infrastructure (GII): Access to Information in the Networked World*, Cambridge MA, MIT Press 2000 <<http://www.cflr.beniculturali.it/>>.

D.N. SNOWDON, E.F. CHURCHILL, E. FRÉCON, *Inhabited Information Spaces – Living with your Data*, London, Springer-Verlag 2005.

CLIR, *Building and Sustaining Digital Collections: Models for Libraries and Museums*, 2001, CLIR <<http://www.clir.org>>, <<http://www.clir.org/pubs/reports/pub100/pub100.pdf>>.

Digicult Report: Technological Landscapes for Tomorrow's Cultural Economy. Unlocking the Value of Cultural Heritage. Luxembourg: European Commission. Directorate General for the Information Society, 2001.

A. GRANELLI, *Tecnologie e patrimonio culturale: riflessioni per una via italiana all'innovazione*, in «Sociologia. Rivista quadrimestrale di Scienze storiche e sociali», XXXIX, 2, pp. 149-155.

ICCD, *Normativa per l'acquisizione digitale delle immagini fotografiche*, 1998, ICCD <<http://www.iccd.beniculturali.it/>>, <<http://www.iccd.beniculturali.it/download/fotodig.pdf>>.

Manifesto AIB per le Biblioteche digitali: <<http://www.aib.it/aib/cg/gbdig-d05a.htm3>>.

S. McNICOL, *The Evalued Toolkit: A Framework for the Qualitative Evaluation of Electronic Information Services*, in «VINE», XXXIV, 4, 2004, pp. 172-175.

MICHAEL <http://www.michael-culture.org/copyright_i.html>.

R. MILLER, S. SCHMIDT, *E-metrics: Measures for Electronic Resources*, Key-note Delivered at the 4th Northumbria International Conference on Performance Measurement in Library and Information Services, Pittsburgh, 12-16 august, 2001, ARL <<http://www.arl.org/>>, <<http://www.arl.org/stats/newmeas/emetrics/miller-schmidt.pdf>>.

OTEBAC, l'Osservatorio tecnologico per i beni e le attività culturali è un centro del Ministero per i Beni e le Attività culturali: <<http://www.otebac.it/>>.

V. PONZANI, G. SOLIMINE, *Rapporto sulle biblioteche italiane 2005-2006*, 2006, AIB <<http://www.aib.it/>>, <<http://www.aib.it/aib/editoria/2006/pub144.htm>>.

TAMMARO, S. CASATI, D. LUZZI, *Biblioteche digitali in Italia*: <http://www.rinascimento-digitale.it/DLA/RapportoSintesi_DLA.pdf>.

S. THEBRIDGE, *Evaluating Electronic Information Services: A Toolkit for Practitioners*, in «Library and Information Research», XXVII, 87, 2003, pp. 38-46, E-LIS <<http://eprints.rclis.org/>>, <<http://eprints.rclis.org/archive/00003421/>>.

S. TOWN, *E-Measures: A Comprehensive Waste of Time*, in «VINE», XXXIV, 4, 2004, pp. 190-195.

B. USHERWOOD, K. WILSON, J. BRYSON, *Relevant Repositories of Public Knowledge? Libraries, Museums and Archives in the Information Age*, in «Journal of Librarianship and Information Science», XXXVI, 2, 2005, pp. 89-98.

E. VAGAGGINI, *Mediateche*: <<http://www.rinascimento-digitale.it/>>.

ANNA MARIA TAMMARO

Finito di stampare nel mese di dicembre 2010

in Pisa dalle

EDIZIONI ETS

Piazza Carrara, 16-19, I-56126 Pisa

info@edizioniets.com

www.edizioniets.com

