



Anno III, Numero 2 - **2008**

Rivista del digitale nei beni culturali

ICCU-ROMA

# Convegno “DC 2008. International Conference on Dublin Core and Metadata Applications”

Berlino 22-26 settembre 2008

**Luigi Siciliano**

*Biblioteca Universitaria di Bolzano*

Dal 22 al 26 settembre 2008 si è tenuta a Berlino “DC 2008. International Conference on Dublin Core and Metadata Applications”<sup>1</sup>. Ospitata presso la Humboldt-Universität di Berlino<sup>2</sup>, e organizzata dalla Max Planck Digital Library<sup>3</sup>, dalla Göttingen State and University Library<sup>4</sup>, dal Kompetenzzentrum Interoperable Metadaten (KIM)<sup>5</sup> e dalla Deutsche Nationalbibliothek<sup>6</sup> con il supporto di Wikimedia Deutschland<sup>7</sup>, l’edizione berlinese è stata significativamente intitolata *Metadata for Semantic and Social Applications*. Ogni conferenza Dublin Core (DC) infatti, oltre a contribuire alla diffusione del noto set di metadati descrittivi e a permettere un’occasione di confronto e bilancio ai singoli gruppi di lavoro interni, presenta uno specifico taglio tematico e approfondisce argomenti di particolare attualità. Dopo i saluti di apertura, la conferenza è en-

trata nel vivo con la *keynote* di Kurt Mehlhorn. Presentando *eResearch. A Max Planck Perspective*, il relatore ha enfatizzato l’impegno della Max Planck Digital Library (MPDL) nella realizzazione di strumenti per la gestione della conoscenza, quali per esempio eSciDoc, un ambiente integrato a supporto di tutte le fasi della ricerca, dalla definizione del progetto, alla programmazione del flusso di lavoro. Un altro esempio è rappresentato dal progetto YAGO: Yet Another Great Ontology, con cui MPDL si proietta nel mondo del *semantic Web*<sup>8</sup> usando la versione tedesca di Wikipedia come gigantesca *knowledge base* e progettando di consultare questa massa di informazioni tramite NAGA: Graph Search with Ranking, un sistema di navigazione visuale. Il successivo intervento del direttore esecutivo di DCMI<sup>9</sup>, Makx Dekkers, ha ripercorso le tappe fondamentali di DC. Una vicenda che inizia

<sup>1</sup> Sul ricchissimo sito della conferenza sono disponibili le slides di presentazione dei vari interventi: International Conference on Dublin Core and Metadata Applications, <http://dc2008.de>.

<sup>2</sup> Humboldt-Universität zu Berlin, <http://www.hu-berlin.de>.

<sup>3</sup> Max Planck Digital Library, <http://www.mpd.lmpg.de>.

<sup>4</sup> Göttingen State and University Library, <http://www.sub.uni-goettingen.de/index-e.html>.

<sup>5</sup> Kompetenzzentrum Interoperable Metadaten (KIM), <http://www.kim-forum.org>.

<sup>6</sup> Deutsche Nationalbibliothek, <http://www.d-nb.de>.

<sup>7</sup> Wikimedia Deutschland, <http://www.wikimedia.de>.

<sup>8</sup> Per semantic web si intende un web in cui a ogni risorsa sia assegnato uno stabile identificatore univoco, o *Uniform Resource Identifier* (URI). Una URI permette non solo di identificare e linkare ogni risorsa senza ambiguità da qualunque punto del world wide web, ma è anche processabile automaticamente da una macchina. Il linguaggio standard W3C che permette di definire, tramite espressioni speciali dette triple, le relazioni tra le risorse del semantic web, è *Resource Description Framework* (RDF).

<sup>9</sup> Dublin Core Metadata Initiative, <http://dublincore.org>.

prima dell'affermazione trionfale dei motori di ricerca, sull'onda dell'esigenza di "catalogare" le risorse disponibili nel Web. Oggi però DC deve contribuire a superare le barriere che mantengono le informazioni in silos non comunicanti fra loro, partecipando alla realizzazione del Web semantico.

La *keynote* di Jennifer Trant, *Access to art museums on-line: a role for social tagging and folksonomy?*, ha trattato invece il problema delle *folksonomy* e del *social tagging* in ambito museale. Una sperimentazione condotta presso il Metropolitan Museum of Art ha dimostrato quanto consistente sia la differenza tra i *tag* assegnati dagli utenti, attenti soprattutto al contenuto figurativo dell'opera, rispetto ai termini riportati nella scheda dei cataloghi. Pertanto se il *social tagging* mostra ancora dei limiti consistenti per essere usato direttamente nella descrizione della risorsa, è sin d'ora di straordinaria utilità per cercare di ridurre il *gap* spesso eccessivo tra la descrizione specialistica delle opere e la percezione degli utenti.

Carol Jean Godby nella sua relazione *Encoding Application Profiles in a Computational Model of the Crosswalk* ha presentato il servizio di crosswalk, ovvero di mappatura automatica tra set di metadati diversi, implementato come *web service* da OCLC (Online Computer Library Center)<sup>10</sup>, mentre il *paper Relating Folksonomies with Dublin Core*, di Maria Elisabete Catarino e presentato da Ana Alice Baptista, ha proposto i risultati di un'analisi manuale dei *tag* assegnati a diverse risorse dagli utenti in un contesto bibliotecario, allo scopo di ricondurli a proprietà DC già esistenti.

Ed Summers con la sua relazione *LCSH, SKOS and Linked Data* ha illustrato le straordinarie potenzialità dei Library of Congress Subject Headings (LCSH) nella realizzazione del Web se-

mantico. Il progetto è infatti quello di rendere LCSH disponibile in Simple Knowledge Organization System (SKOS), un linguaggio formale per la rappresentazione di thesauri basato su RDF, rendendolo così fruibile nella forma di *linked data*. Il progetto prevede infatti di assegnare un Uniform Resource Identifier (URI) a ogni soggetto, così che, per esempio, ovunque nel Web si possa fare riferimento univoco al soggetto World Wide Web con la stringa <http://lcsch.info/sh95000541#concept>. In tale forma il soggetto potrebbe diventare membro di una tripla espressa in SKOS, formando una gigantesca ragnatela di significati connessi tra loro. Nello stesso ambito tematico si colloca l'intervento di Xia Lin, *Theme Creation for Digital Collections*. Va notato che nel contesto delle *topic maps* il termine *theme* indica l'estensione semantica di un determinato *topic*. Il prototipo descritto rappresenta un tentativo di arricchire di relazioni semantiche le risorse digitali contenute nella Internet Public Library (IPL)<sup>11</sup>. A tal fine si sono sperimentati metodi di generazione di metadati a partire dal contesto, e il contesto è stato a sua volta estratto automaticamente tramite analisi delle collezioni digitali e dei *log* di ricerca degli utenti. Nel sistema le mappature semantiche generate automaticamente con questi strumenti possono poi essere riviste manualmente tramite un'apposita interfaccia grafica.

Anche il *mapping* tra diversi sistemi semantici è stato uno dei temi maggiormente trattati nel corso della conferenza. Philipp Mayr (*Comparing human and automatic thesaurus mapping approaches in the agricultural domain*) ha illustrato i risultati di un confronto fra gli esiti di un *mapping* manuale e quelli di un *mapping* automatico, usando il thesaurus per l'agricoltura AGROVOC come banco di prova<sup>12</sup>.

<sup>10</sup> OCLC Crosswalk Web Service Demo, <http://www.oclc.org/research/researchworks/xwalk>.

<sup>11</sup> Internet Public Library è un repertorio di migliaia di risorse web qualificate su diversi argomenti e una delle più antiche digital library esistenti, <http://www.ipl.org>.

<sup>12</sup> Il confronto è stato in particolare fra la mappatura manuale tra AGROVOC e il sistema di soggettazione tedesco *Schlagwortnormdatei* (SWD) e la mappatura generata automaticamente tra AGROVOC e il *National Agricultural Library Thesaurus* (NALT).

La mappatura manuale è infatti un'operazione estremamente onerosa in termini di risorse umane ed economiche, e pertanto le potenzialità dei *mapping* automatici sono oggetto di attenta indagine. I test mostrano tuttavia come i *mapping* automatici incontrino notevoli difficoltà nell'ambito delle scienze umane e sociali e come ottengano risultati migliori nelle discipline scientifiche che offrono tassonomie rigorose, come per esempio la biologia, ma comunque inferiori rispetto ai *mappings* manuali. Un'integrazione delle due tipologie di *mapping* risulta pertanto la strada attualmente più proficua.

La successiva sessione parallela ha proposto due *report* su progetti in corso (un *paper* di Martin Malmsten sull'approccio al Web semantico della Biblioteca Nazionale Svedese<sup>13</sup> e uno di Eva Méndez sull'iniziativa DC Microformats<sup>14</sup>) e quattro workshop per altrettanti gruppi di lavoro: *Education*<sup>15</sup>, *Government*<sup>16</sup>, *Social tagging*<sup>17</sup> e *Libraries*<sup>18</sup>. Quest'ultimo, coordinato da Christine Frodl, si è soffermato principalmente sullo stato attuale del DC Libraries Application Profile (DC-Lib)<sup>19</sup>. Il suo sviluppo risente tuttavia di un certo ritardo, e la sua evoluzione avverrà in stretto raccordo con quella di un altro DCAP, Scholarly Works Application Profile (SWAP). Come sottolineato da Corey Harper, DC-Lib oltre a tenere ben presente i *Functional Requirements for Bibliographic Records* (FRBR)<sup>20</sup> deve tenere conto dell'imminente pubblicazione di *Resource Description and*

*Access* (RDA). È inoltre necessario rimuovere da DC-Lib i termini *DateCaptured*, *Edition* e *Location* provenienti da Metadata Object Description Schema (MODS)<sup>21</sup>, schema incompatibile con il modello logico di DCAM (DC Abstract Model). Nella stessa riunione Sally Chambers ha illustrato ai presenti l'architettura dei metadati di The European Library (TEL)<sup>22</sup>.

Nella successiva *keynote* – *The Ultimate Question* – Ute Schwens si è chiesta provocatoriamente se nel mondo dominato da Google siano ancora necessarie regole per la catalogazione. La risposta è affermativa, ma le biblioteche potranno raccogliere la sfida e fornire conoscenza selezionata e qualificata solo se verrà soddisfatta la domanda di standard internazionali condivisi. Per questo la Biblioteca Nazionale tedesca ha annunciato la decisione strategica di adottare RDA come norma per la catalogazione.

La relazione *Automatic Metadata Extraction from Museum Specimen Labels* presentata da P. Bryan Heidorn ha illustrato il progetto di digitalizzazione delle schede degli *specimina* botanici e zoologici conservati nei musei americani. Il lavoro risulta distribuito nelle diverse sedi, che trasmettono le scansioni delle schede cartacee al server centrale. La lettura OCR delle schede avviene facendo riferimento a dei modelli realizzati partendo dalle tipologie di scheda più diffuse. Una revisione manuale permette di identificare nuovi tipi di scheda e istruire così il software che migliora mano in mano in precisione.

<sup>13</sup> LIBRIS, <http://libris.kb.se>.

<sup>14</sup> DCMF, <http://www.webposible.com/microformats-dublincore>.

<sup>15</sup> DCM Education Community, <http://dublincore.org/groups/education>.

<sup>16</sup> DCM Government Community, <http://dublincore.org/groups/government>.

<sup>17</sup> DCM Social Tagging Community, <http://dublincore.org/groups/social-tagging>.

<sup>18</sup> DCM Libraries Community, <http://dublincore.org/groups/libraries>.

<sup>19</sup> Un *DC Application Profile* (DCAP) è un documento che definisce e descrive i metadati DC usati in una particolare applicazione per un determinato scopo. Esistono quindi diversi DCAP, sorta di dialetti di DC, sviluppati dalle diverse comunità sulla base di specifiche esigenze. Si veda: *Guidelines for Dublin Core Application Profiles*, <http://dublincore.org/documents/profile-guidelines>.

<sup>20</sup> *Functional Requirements for Bibliographic Records* (FRBR), <http://www.ifla.org/VII/s13/frbr>.

<sup>21</sup> Metadata Object Description Schema (MODS), <http://www.loc.gov/standards/mods>.

<sup>22</sup> The European Library (TEL), <http://www.theeuropeanlibrary.org>.

L'intervento di Stuart A. Sutton, *Achievement Standards Network (ASN): An Application Profile for Mapping K-12 Educational Resources to Achievement Standards*, presenta l'imponente progetto di un Application Profile che permetta di superare la varietà di standard esistenti nel settore dell'educazione. Il risultato è ASN, un repository istituzionale in cui tutti gli standard esistenti (più di settecento) sono espressi in formato RDF. Questo sistema facilita valutazioni complessive e permette di stabilire correlazioni (tra singole risorse educative e standard) e allineamenti (relazioni tra diversi standard).

In un mondo in cui i motori di ricerca sono in grado di raggiungere direttamente ogni singola risorsa, si tende purtroppo a dimenticare il valore del contesto, e in particolare della collezione. Il progetto Collection/Item Metadata Relationships (CIMR), illustrato da Richard J. Urban, prevede invece:

- lo sviluppo di un modello logico di metadati per le relazioni fra collezione e item, corredato di regole di inferenza;
- la realizzazione di analisi empiriche per valutare la rispondenza del modello alle esigenze delle comunità di sviluppatori di metadati e di utenti;
- l'implementazione di prototipi basati su RDF/OWL<sup>23</sup> per la ricerca, la consultazione e la navigazione delle collezioni.

La successiva sessione di *project reports* ha visto la presentazione di tre progetti relativi a

*Metadata Scheme Design, Application, and Use*<sup>24</sup>, e di tre workshop. Il primo, curato da Traugott Koch, dedicato a NKOS (Networked Knowledge Organization Systems), il secondo, diretto da Douglas Campbell si è occupato del tema degli *identifiers*<sup>25</sup>, il terzo infine, di particolare interesse per il mondo delle biblioteche, si è concentrato su RDA. Nel corso del workshop Diane Hillman ne ha annunciato la pubblicazione per giugno 2009, sebbene sia già disponibile *on-line* un *Metadata registry*<sup>26</sup>. Autenticamente multilingue, RDA si propone come primo standard internazionale a livello mondiale direttamente pensato per la diffusione *on-line* (sebbene sia prevista anche una versione a stampa). Tutti potranno usare tale standard ma le linee guida saranno diffuse a pagamento. Come spiegato da Gordon Dunsire nella stessa occasione, RDA condivide il modello Entità/Relazioni di FRBR e un gruppo di lavoro DC si occupa attivamente della sua adozione e diffusione<sup>27</sup>.

Una nuova sessione parallela ha visto riunirsi il gruppo Scholar<sup>28</sup> coordinato da Rosemary Russell, e il gruppo DCMI/IEEE<sup>29</sup> coordinato da Mikael Nilsson, mentre Thomas Severiens ha diretto la riunione del workgroup Tools<sup>30</sup>.

Alla ripresa dei lavori Seth van Hooland e Seth Kaufman hanno introdotto il problema della valutazione qualitativa dei metadati nel loro intervento *Answering the call for more accountability: Applying data profiling to museum metadata*. Una prima soluzione sperimentata nell'ambito della *cultural heritage* consiste nell'usare dei *tool* che effettuano

<sup>23</sup> Web Ontology Language (OWL) è un'estensione di RDF, usata per rappresentare esplicitamente significato e semantica di termini. Si veda: <http://www.w3.org/2004/OWL/> e [http://en.wikipedia.org/wiki/Web\\_Ontology\\_Language](http://en.wikipedia.org/wiki/Web_Ontology_Language).

<sup>24</sup> Gli interventi sono stati rispettivamente: *The Dryad Data Repository: A Singapore Framework Metadata Architecture in a DSpace Environment*; *Applying DCMI Elements to Digital Images and Text in the Archimedes Palimpsest Program*; *Assessing Descriptive Substance in Free-Text Collection-Level Metadata*.

<sup>25</sup> DCMI Identifiers Community, <http://dublincore.org/groups/identifiers>.

<sup>26</sup> National Science Digital Library (NSDL) Metadata Registry, <http://metadataregistry.org>.

<sup>27</sup> DCMI RDA, <http://dublincore.org/dcmirdataskgroup>.

<sup>28</sup> DCMI Scholarly Communications Community, <http://dublincore.org/groups/scholar>.

<sup>29</sup> DCMI/IEEE LTSC Taskforce, <http://dublincore.org/educationwiki/DCMIIEEELTSCTaskforce>.

<sup>30</sup> DCMI Tools Community, <http://dublincore.org/groups/tools>.

controlli di consistenza dei metadati. Si può parlare in questo senso di tecniche di *data profiling*, ma è anche possibile avere un approccio proattivo e investigare invece quali siano le reali esigenze degli utenti. A tale scopo si usano di solito analisi dei *log* delle *query* degli utenti oppure questionari, ma i relatori propongono di offrire agli utenti interfacce di ricerca personalizzabili e visualizzazioni dinamiche dei metadati. Sarà il monitoraggio costante delle scelte operate dagli utenti nelle fasi di personalizzazione delle interfacce a fornire le necessarie indicazioni su quali metadati siano maggiormente utili, quale sia il loro uso reale, e quali siano invece trascurati o inutili. Un prototipo in tal senso è stato sviluppato in ambito museale con il software open source OpenCollection<sup>31</sup>.

Un approccio diverso allo stesso tema è quello seguito da Thomas Margaritopoulos, nella sua riflessione teorica *A Conceptual Framework for Metadata Quality Assessment*, secondo la quale i criteri di riferimento per valutare la qualità dei metadati sono i concetti di *correctness*, *completeness*, *relevance*. Paragonando i metadati alle deposizioni dei testimoni in un tribunale, si intende per correttezza il rigore della sintassi usata e l'esattezza dei dati; per completezza la capacità di descrivere la risorsa in maniera esauriente; per rilevanza infine il riferimento corretto della risorsa al contesto per il quale è lecito un suo utilizzo. Completano lo schema teorico proposto la definizione di alcune regole logiche per la definizione delle relazioni fra i metadati e le risorse stesse.

La *keynote* di Paul Miller, dal titolo *Why the Semantic Web matters*, si è soffermata sulla necessità di abbattere le barriere tra i silos informativi facilitando invece i collegamenti fra i dati all'interno del Web semantico.

Un esperimento condotto presso i laboratori

Talis<sup>32</sup>, in cui in due settimane, lavorando sui riferimenti contenuti in soli 15 articoli scientifici, è stato possibile elaborare una gigantesca ragnatela di svariate migliaia di link semantici, mostra le potenzialità straordinarie di questo approccio.

Anche il *social tagging* è una realtà in grande espansione, ma come identificare relazioni semantiche a partire dalla disordinata congerie di *tag*? Miao Chen, nella relazione *Semantic Relation Extraction from Socially-Generated Tags: A Methodological Exploration* ha affrontato la questione di petto e in maniera estremamente pragmatica. È l'identificazione del più verisimile contesto d'uso di un termine usato come *tag* che può chiarircene il significato. E nella sperimentazione proposta tale contesto viene identificato riferendosi alla più ampia *knowledge base* attualmente esistente: Google. I risultati ottenuti sono infine rielaborati grazie ad altre tecnologie, quali *machine learning* e Natural Language Processing (NLP)<sup>33</sup>.

L'approccio pragmatico è condiviso da Simon Scerri che presenta *The State of the Art in Tag Ontologies: A Semantic Model for Tagging and Folksonomies*. In fondo nel pur confuso mondo dei *tag* i concetti principali restano tre: l'utente, il *tag*, la risorsa. Perché non gestire tutto questo in maniera standard? Occorre una ontologia per i *tag* e gli autori propongono una comparazione tra quelle esistenti. Dal confronto emerge come l'utilizzo combinato di MOAT (Meaning Of A Tag) e SCOT (Social Semantic Cloud of Tags) permetta di rappresentare *collaborative tagging* e *folksonomy* in maniera suscettibile di trattamento automatico e condivisibile attraverso le più varie piattaforme di *social tagging*.

La quarta sessione parallela ha visto quattro eventi in contemporanea. Mentre si teneva una sessione speciale dedicata a *Wikis and*

<sup>31</sup> OpenCollection, <http://www.opencollection.org>.

<sup>32</sup> Talis Inc., <http://www.talis.com>.

<sup>33</sup> Nel contesto italiano si parla spesso anche di Trattamento Automatico della Lingua (TAL). Sebbene non recente, un libro bianco sul tema è: <http://forumtal.fub.it/LibroBianco.php>.

*Metadata*<sup>34</sup>, si riunivano i gruppi di lavoro Knowledge Management<sup>35</sup>, Registry<sup>36</sup> e Architecture. Quest'ultimo workshop, coordinato da Mikael Nilsson, ha chiarito le prospettive di DC sul medio periodo, soffermandosi su due punti: la necessità di seguire le regole indicate nel cosiddetto Singapore Framework<sup>37</sup> quando si intenda elaborare un DC AP, e l'importanza dell'imminente rilascio della specifica *DC Description Set Profile*<sup>38</sup>, che formalizzerà in maniera *machine readable* regole e vincoli indicati nei DCAP creati dalle diverse comunità. Si possono pertanto distinguere oggi quattro livelli di interoperabilità DC:

- condivisione di vocabolario;
- condivisione della codifica in RDF;
- conformità al DC Abstract Model;
- conformità al *Description Set Profile* (DSP).

L'ultima sessione parallela ha visto addirittura cinque eventi svolgersi parallelamente. Mentre si tenevano i workshop dei gruppi Accessibility<sup>39</sup>, Localization and Internationalization<sup>40</sup>, Metadata for Scientific Datasets<sup>41</sup>, Architecture Technical Session, la Project Report Session 3 si è soffermata sul tema *Vocabulary Integration and Interoperability*.

Philipp Mayr (*Building a terminology network for search: the KoMoHe project*) ha presentato i risultati del *mapping* manuale condotto

dal progetto KoMoHe<sup>42</sup>. Nato per permettere una ricerca trasparente attraverso collezioni multiple ed eterogenee, il progetto ha realizzato una mappatura manuale per un totale di 25 fra thesauri, sistemi di classificazione e di soggettazione. Una serie di test ha evidenziato come sia una ricerca per soggetto, sia una ricerca a testo libero, forniscano risultati migliori laddove facciano uso di tali mappature. Michael Panzer, nel suo intervento *Cool URIs for the DDC: Towards Web-scale accessibility of a large classification system*, ha illustrato lo sforzo di OCLC per portare la Dewey Decimal Classification (DDC) nel Web semantico. La realizzazione di identificatori univoci (URI) è resa particolarmente complessa dal fatto che DDC cambia nel tempo (a seconda delle versioni) e nello spazio (a seconda delle diverse localizzazioni). Viene così presentata l'elaborata sintassi attualmente allo studio per permettere di linkare univocamente una classe Dewey tenendo conto di tutti questi fattori.

A Barbara Levergood è toccato l'onere di chiudere la sessione con il suo intervento *The Specification of the Language of the Field and Interoperability: Cross-language Access to Catalogues and Online Libraries (CACAO)*. Sfruttando l'esperienza accumulata in seno al progetto CACAO, che prevede la realizzazione di una piattaforma software per le biblioteche che permetta a un utente di inoltrare una *query* nella propria lingua e di ottenere auto-

<sup>34</sup> I tre progetti presentati sono stati: *Facilitating Wiki/Repository Communication with Metadata*; *Aratta: A Web-based Research Tool for Collaborative Research on Iranian Architectural History*; *Wikipedia as Controlled Vocabulary*.

<sup>35</sup> DCM Knowledge Management Community, <http://dublincore.org/groups/km>.

<sup>36</sup> DCM Registry Community, <http://dublincore.org/groups/registry>.

<sup>37</sup> The Singapore Framework for Dublin Core Application Profiles, <http://dublincore.org/documents/singapore-framework>.

<sup>38</sup> Description Set Profiles: A constraint language for Dublin Core Application Profiles, <http://dublincore.org/documents/dc-dsp>.

<sup>39</sup> DCM Accessibility Community, <http://dublincore.org/groups/access>.

<sup>40</sup> DCM Localization and Internationalization Community, <http://dublincore.org/groups/languages>.

<sup>41</sup> Sono stati illustrati diversi progetti. Fra questi in particolare il DISC DataShare Project, <http://www.disc-uk.org/datashare.html>.

<sup>42</sup> KoMoHe Project, [http://www.gesis.org/en/research/information\\_technology/komohe.htm](http://www.gesis.org/en/research/information_technology/komohe.htm).

maticamente risultati pertinenti anche in altre lingue, Barbara Levergood ha illustrato come sia possibile utilizzare la specificazione della lingua nei metadati per migliorare la performance dell'*information retrieval* in casi critici, in particolare nella traduzione e trattamento di termini ambigui e polisemici.

La presentazione di dodici poster<sup>43</sup>, due *tutorial* introduttivi opzionali e quattro seminari paralleli dedicati rispettivamente a *User Generated Metadata, Using the TEI for Documenting Describing Documents, PREMIS Metadata Tutorial* e *Ontology Design and Interoperability* hanno ulteriormente arricchito un programma che si è rivelato densissimo e che è stato seguito da trecento partecipanti provenienti da poco meno di quaranta paesi del mondo.

A distanza di anni DC si conferma quindi più che mai vitale nel panorama mondiale dei metadati descrittivi. Sembra tuttavia di poter cogliere una tensione tra l'evoluzione di DC verso un modello astratto potenzialmente molto effi-

cace, ma non ancora completo in ogni sua parte ed estremamente complesso, e dall'altro lato l'eredità di un successo legato proprio alla stabilità e semplicità dei quindici elementi storici e alla documentazione chiara e centralizzata. Nell'enfasi sulla modularità del sistema e sulla possibilità di implementare elementi mutuati da altri standard, purché compatibili a livello di *abstract model*, DC acquista in flessibilità ma sembra ridursi a una sorta di contenitore. Potrebbe in futuro delinearsi uno scarto fra la pratica e la teoria? Gli utilizzatori, primi fra tutti le biblioteche, sapranno accettare nuovi modelli di metadati flessibili e rigorosi ma difficili da comprendere, implementare e gestire? O sarà forse DC in grado di offrire alle diverse comunità strumenti in grado di gestire in maniera trasparente la complessità del sottostante modello logico? Si tratta di problemi in fondo ricorrenti, e gli esiti potranno essere positivi solo se le diverse comunità, prima fra tutte quella delle biblioteche, saranno parte attiva nella realizzazione del DC di domani.

<sup>43</sup> Tra tutti è stato particolarmente apprezzato, e premiato quale miglior poster della conferenza, quello realizzato da Marcia L. Zeng, dedicato alla realizzazione in SKOS di un *Chinese Classified Thesaurus* (CCT).