

Vom 23. bis 24. April 2013 fand in London der Workshop »Open Data on the Web«<sup>1</sup> als ein Gemeinschaftsunternehmen des World Wide Web Consortium (W3C), des Open Data Institute (ODI) und der Open Knowledge Foundation (OKFN) statt. Die Organisatoren hatten eingeladen, um von der Community eine Hilfestellung bei der Aufgabe zu erhalten, für die nächsten Jahre die Prioritäten des W3C im Bereich Data zu setzen. Ausgangspunkt der Veranstaltung war die These, Daten seien das »neue Öl« und der Zugang zu offenen Daten der Motor der digitalen Wirtschaft. Vor allem öffentliche Körperschaften wie Staaten, Länder und Gemeinden öffnen sog. Data Hubs (Datendrehscheiben), aber auch Web 2.0-Dienste wie OpenStreetMap und Open Signal erzeugen große Mengen an Daten, darüber hinaus verwenden inzwischen viele Vertriebsfirmen Vokabulare wie schema.org oder Good Relations, um Suchmaschinen den Zugang zu Daten über ihre Produkte zu vereinfachen.

Um sicher zu stellen, dass alle Teilnehmerinnen und Teilnehmer ein begründetes Interesse am Workshop hatten – und somit auch um die Teilnehmerzahl in Grenzen zu halten – war die Abgabe eines Position Papers aus dem Umfeld von Open Data eine Voraussetzung. Ausgehend von den Thesen bzw. den Fallstudien in den insgesamt 62 Papers<sup>2</sup> sollte in den zwei Tagen diskutiert werden, ob die These vom »neuen Öl« auf Realitäten beruht, bzw. wie sie realisiert und wie die Arbeit mit offenen Daten vereinfacht werden kann. Es würde hier den Rahmen sprengen, auf alle Inhalte einzugehen, insofern stellt das Folgende eine kurze Zusammenfassung der zwei Tage dar, gibt aber nicht die vollständigen Inhalte des Workshops wieder.

### **Auftakt: Open Data Meetup der OKFN**

Die OKFN richtet in London regelmäßig kleinere Treffen zum Thema Open Data aus, bei denen Interessenten aus verschiedenen Gebieten (Forschung, Verwaltung, Industrie, Kultur etc.) die Gelegenheit bekommen, sich kennen zu lernen und Erfahrungen auszu-

tauschen. Diesmal war eine Reihe von Lightning Talks von verschiedenen Stakeholdern vorgesehen.

Die Abendveranstaltung eröffnete Rufus Pollock von der OKFN, der in einer Zeit, wo alle über Big Data sprechen, die kontroverse These aufstellte, dass es eben nicht Big Data sein müsse, sondern dass oft kleinere, untereinander verlinkte Datensets nützlicher sind. Als Beispiel nannte er geobasierte Dienste, wo Daten von mehreren Anbietern – z. B. Läden, Restaurants, Museen – die in der Nähe eines bestimmten Orts sind, »on-the-fly«, d. h. beim Aufruf des Dienstes, aggregiert werden. Dies war später auch ein Thema auf der Hauptkonferenz.

Im Anschluss daran berichtete Antonio Acuña – Chef der Open-Data-Seite der britischen Verwaltung data.gov.uk – über die Pläne zur Weiterentwicklung der Plattform. Diese umfassen u. a. eine Seite, wo die Ministerien über ihre monatlichen Ausgaben berichten.<sup>3</sup> Carl Reynolds von Open Healthcare UK stellte die Plattform openprescribing.org vor, wo ausgehend von anonymisierten Daten des Gesundheitsministeriums die Verschreibung von Medikamenten (z. B. in welchen Gebieten Ärzte ungewöhnlich viel Methadon verschreiben) nachverfolgt werden kann.

Paul Maltby vom Cabinet Office wandte sich an die Open-Data-Community mit der Frage, wie der »übernächste Schritt« aussehen könnte. Als Nächstes würden sie sich auf Daten mit Georeferenzierung konzentrieren, um ortsbasierte Dienste zu ermöglichen, danach steht Tagging von Datasets auf der Agenda, doch was anschließend folgen könnte und den Nutzerinnen und Nutzern am meisten helfen würde, wisse er nicht.

### **Hauptkonferenz, Tag 1**

#### **Das Versprechen von Open Data**

Die Konferenz »Open Data on the Web« hatte ein strikt organisiertes Programm,<sup>4</sup> mit vielen Vorträgen und Diskussionen. Nach einer kurzen Einführung von Phil Archer (W3C) und Dan Brickley (Google) wurde das Programm mit einer Session über das Versprechen und die Erwartungen von Open Data eröffnet. John Sheridan von The National Archives betonte, wie wichtig es sei, Open-Data-Dienste nicht »auf Sand« zu bauen, besonders wenn es sich um Daten aus der Verwaltung handelt. Zentrale Qualitätsaspekte seien Auffindbarkeit, Persistenz und Nachhaltigkeit, die es ermöglichen, langfristig Dienste ausgehend von den Daten aufzubauen. Um die Datenqualität zu verbessern, können Communities gute Hilfe leisten – beispielsweise durch Tagging von Foto- und Bewegtbildmaterial aus Archiven und Museen.

Die nächsten zwei Präsentationen gingen besonders auf die Bedeutung von Open Data in Entwick-

geobasierte Dienste

Daten sind das »neue Öl«

zentrale Qualitätsaspekte

lungsländern ein. Millie Begovic von United Nations Development Programme untersuchte die Frage, wie die verfügbare Information verwendet werden kann, um die Auswirkung von Programmen zur Entwicklungshilfe zu evaluieren. So wurde untersucht, ob die Vergabe größerer Bauleistungen an regionale Unternehmen geht (Ergebnis: nein), und wenn nicht, wohin das Geld fließt (Ergebnis: u. a. nach Deutschland und Frankreich). Ziel des Projektes ist es, eine Datenbasis zu schaffen, mit der vorhersehbar wird, welchen Effekt ein bestimmtes Projekt haben wird. Tim Davies von Open Data in Developing Countries wies auf die Gefahr einer neuen »digital divide« (Digitale Spaltung) hin: diesmal nicht zwischen Regionen, in denen die Bevölkerung Zugriff auf digitale Information hat oder nicht, sondern zwischen Regionen, in denen die Verwaltung diese Information als Open Data zur Verfügung stellt oder eben nicht. Die Entwicklung scheint verhalten positiv zu sein: Open Data ist ein weltweites Thema und viele Regierungen aus Entwicklungsländern haben um Projektmittel gebeten, um entsprechende Projekte finanzieren zu können. Wichtig bei der Entwicklung der Projekte sei es, bei den Fragen, welche Daten man veröffentlichen will und welche Community dabei angesprochen werden soll, einem vorher festgelegten Plan zu folgen.

Hajo Schreier vom niederländischen Innenministerium griff die Vorlage von Tim Davies auf. Die niederländischen Open-Data-Projekte der Verwaltung seien besonders darauf ausgerichtet, Daten zu veröffentlichen, die sozio-ökonomische Information über den Bevölkerungsschwund in bestimmten Regionen, bzw. über soziale Brennpunkte zur Verfügung stellen. Dabei wolle man auch zeigen, dass die Bereitstellung der Daten als *Linked Open Data* auch für die Ministerien und Ämter durchaus vorteilhaft sein kann. Den Schluss der ersten Session machte Bob Schloss von IBM mit einem Position Paper über Designprinzipien für Open-Data-Dienste. Grundtenor war, dass derjenige, der (außer der öffentlichen Verwaltung) seine Daten öffnen will, auch einen handfesten Vorteil davon haben möchte. Wenn die Daten dann verfügbar sind, müssen entsprechende (transparente) Prozesse zur Qualitätssicherung vorhanden sein, damit die Nutzerinnen und Nutzer wissen, was sie von den Daten erwarten können. Drittens müsse Provenienzinformation vorhanden sein, damit es – auch durch längere Prozessketten hindurch – sichtbar ist, woher die Daten kommen und wer dafür verantwortlich ist.

Gegenstand der anschließenden Diskussion war v. a. das Thema Nachhaltigkeit. Einige der Teilnehmer waren Kleinunternehmer, die Dienste ausgehend von Open Data anbieten wollen. Für ihr Geschäftsmodell

ist es elementar, dass die Daten auch in Zukunft zur Verfügung stehen und dass die Anbieter die gleichen Formate verwenden, da jegliche Formatänderung eine Systemanpassung seitens der Datenkonsumenten zur Folge hat.

### Datenformate

Um die offenen Daten sinnvoll verwenden zu können, müssen sie in einem Format vorliegen, das der Kunde konsumieren kann. Da viele Berichte als PDF veröffentlicht werden, durfte Jim King – der bei Adobe PDF als Dateiformat maßgeblich mitentwickelt hat – präsentieren, was von PDF erwartet werden kann und was nicht. Im Zentrum seiner Argumentation stand, dass PDF ein Dokumentformat ist und nicht dafür entworfen wurde, um Daten zu transportieren. Dennoch sei es durch die Möglichkeiten, PDFs entsprechend zu strukturieren, durchaus möglich, Daten in PDF-Dokumente so zu integrieren, dass sie auch wieder herausgeholt und wiederverwendet werden können. Als Beispiel nannte er die Möglichkeit, ein XML- oder CSV-Dokument in ein PDF zu integrieren. In der folgenden Diskussion wies Peter Murray-Rust darauf hin, dass die Möglichkeit, PDFs gut zu strukturieren, nicht hilfreich sei, solange die Dokumenthersteller sie nicht verwendeten. Als Beispiel nannte er die großen Wissenschaftsverlage, deren PDF-Publikationen weder strukturiert sind noch Metadaten mitgeben; dadurch seien die Ergebnisse von schätzungsweise 100 Milliarden US-Dollar Forschungsgeldern nicht nachnutzbar. Ob die schlechte Qualität aus der mangelnden Nachfrage nach gut strukturierten Dokumenten mit Metadaten resultiere oder die Erstellung nachnutzbarer PDFs ein finanzielles Problem darstelle, ist noch offen.

Es folgte eine Podiumsdiskussion über tabellarische Datenformate, die bei den Kunden von Open Data sehr beliebt sind – besonders im Bereich der statistischen Daten. Das Problem hierbei sei jedoch, dass z. B. CSV (Comma Separated Values) oft für Information verwendet wird, die nicht wirklich gut tabellarisch darzustellen ist. Es gab viele Vorschläge, wie CSV verbessert werden könnte,<sup>5</sup> die aber dann auch verlangen, dass die Entwickler bereit sind, sich mit den Formaten auseinanderzusetzen. Ähnlich wie viele bei ihrer Literaturrecherche vorgehen – lieber Publikationen niedrigerer Qualität zu verwenden, an die man leicht kommt, als hochwertige, für die man in die Bibliothek gehen muss –, scheint bei vielen Konsumenten von Open Data zu gelten, dass man lieber schlechte Daten nimmt, die einfach zu interpretieren sind, als wirklich gute, aber schwer zu bearbeitende Daten. Eine größere Diskussion entbrannte auch darüber, wie man maschinenlesbare Metainformation über die Struktur einer

Gefahr einer neuen  
»digital divide«

keine Nachnutzung  
möglich

Bequemlichkeit der User

CSV-Datei (welche Information ist in einer bestimmten Spalte bzw. Zeile enthalten) zur Verfügung stellen kann. Die Meinung tendierte dahin, diese Information in einer separaten Datei bereitzustellen und die CSV-Datei mit der Metainformation mithilfe der Möglichkeiten, die das http-Protokoll bietet, zu verknüpfen. Es wurde auch darauf hingewiesen, dass es wichtig ist, eine gute, suchmaschinenoptimierte Homepage für die Datendienste bereit zu halten, um die Auffindbarkeit und den Zugriff zu optimieren.

### **(Linked) Open Data**

In einer Reihe von Lightning Talks mit anschließender Diskussion wurde dann die Schnittmenge zwischen Open Data und Linked Data diskutiert. Tenor sowohl der Vorträge als auch der Diskussion war, Linked Data als »nice-to-have« zu betrachten – für viele Anwendungsfälle sind tabellarische Daten wie z. B. CSV das präferierte Format, da es technisch einfach zu interpretieren ist –, während die Offenheit der Daten ein absolutes Muss ist. Linked Data scheint am ehesten dort Anwendung zu finden, wo es z. B. um die Aggregation von Information über Firmen geht und um die anschließende Verknüpfung der verschiedenen Datentöpfe beispielsweise über die Steuernummer. Für Datenbereinigung und Verlinkung ist eine gute Toolunterstützung wesentlich, da viele Anbieter nicht auf sauber strukturierte und verlinkte Rohdaten aufsetzen. Als ein Beispiel wurde die DataLift-Plattform vorgestellt.<sup>6</sup> Auch für interne Datenintegration scheint Linked Data eine Option zu sein, wie ein Projekt aus dem Science Museum zeigt. In diesen Fällen müssen die verlinkten Daten nicht immer offen sein, sondern können teilweise auch nur für den internen Gebrauch genutzt werden. Die neu entstehende W3C Open Linked Education Community Group hat zum Ziel, disparate Repositorien und Sammlungen von Online-Kursen thematisch zu verlinken. Auch hier ist neben der Erstellung eines geeigneten Themenvokabulars – ausgehend von Themen aus Wikipedia – die Erstellung eines Verlinkungstools – ausgehend von Google Refine – eines der wichtigeren Arbeitsergebnisse. Dass die Verlinkung besser funktioniert, wenn die Daten innerhalb einer eng gefassten Domäne vorgehalten werden, wurde erwähnt, ist aber keine wirklich neue Erkenntnis.

### **Geodaten**

Nach einer kurzen Session über die Bedeutung von interoperablen Standards (Botschaft: wer Daten veröffentlicht, soll Vokabulare / Datenformate / Schnittstellen verwenden, die der Zielgruppe bekannt sind, denn die Konsumenten sollen es einfach haben), kam eine noch kürzere Session über offene Geodaten. Es

war Konsens unter den Rednern und auch im Publikum, dass die Georeferenzierung eine wichtige und leicht zu visualisierende Komponente von Anwendungen ist, aber auch, dass die gängigen GIS-Tools nicht in der Lage sind, semantisch kodierte GIS-Information zu verarbeiten. Im Rahmen von INSPIRE sind EU-weit verwendete Datenmodelle entstanden, die den Datenaustausch erleichtern, es gibt jedoch noch kein Konzept für die persistente Referenzierung der INSPIRE-Daten, was dazu führt, dass sie beispielsweise für Linked-Datendienste nur begrenzt nutzbar sind. Desweiteren sehen viele Datenlieferanten keinen Nutzen darin, die Daten offen zur Verfügung zu stellen, so dass die zur Verfügung stehenden Datensets oft lückenhaft sind. Im EU-Projekt GeoKnow wird daher Information aus mehreren Quellen, inklusive Crowdsourcing-Plattformen wie Wikipedia, gesammelt, um eine größtmögliche Deckung zu bekommen.

### **Geschäftsmodelle**

Das Thema Geschäftsmodelle für Open Data bekam eine eigene Podiumsdiskussion. John Sheridan leitete mit einer provokativen Frage ein: »It's 2020 and we have seen the first multi-billion dollar company based on open data; how did that happen?« Eine Antwort auf diese Frage erhielt er nicht, stattdessen waren die Probleme, die Firmengründer haben, wenn sie ihr Geschäft auf Open Data aufbauen wollen, ein Thema. Ein Problem stellt die rasante technische Entwicklung dar, die es fast unmöglich macht, mit ihr Schritt zu halten. Ein weiteres Problem ist der fehlende Konsens darüber, welche Datasets für welche Anwendungsfälle in Zukunft wichtig sein werden, und es somit schwierig ist, sich aktuell auf bestimmte Daten zu konzentrieren. Seitens der Entwickler/Firmengründer wird sehr häufig der Wunsch geäußert, dass eine datenzentrierte öffentliche Verwaltung wie eine Programmierschnittstelle funktionieren soll und dass alle Teile der Verwaltung gemeinsame Datenmodelle verwenden, damit Applikationen auch kontextübergreifend verwendbar sind. Die Qualität der Daten wurde ebenfalls diskutiert. Viele Stimmen fordern, dass öffentliche Einrichtungen ihre Daten öffnen, »da sie ohnehin Geld für deren Erstellung erhalten«. Auf der anderen Seite stehen die Forderungen von übergeordneten Behörden, dass sich beispielsweise Digitalisierungsaktionen rechnen und die Einrichtungen dafür Sorge tragen, durch Verkauf und Lizenzierung wieder Einnahmen zu erzielen. Einige setzen auf gemischte Modelle, wo Digitalisate in niedrigerer Qualität (»gut genug für das Web«) frei zur Verfügung stehen, während hochauflösende TIFFs nur gegen Gebühr und entsprechende Lizenz einschränkung zu beziehen sind. Bezogen auf

**Lightning Talks**

**gute Toolunterstützung  
wesentlich**

**gemeinsame  
Datenmodelle**

die Verbraucherseite kam die Frage auf, wie ein Datenlieferant reagiert, wenn ein Kunde ihn wegen mangelnder Datenqualität vor Gericht zitiert; eine lebhaft Diskussion folgte mit dem Fazit, dass entsprechende Informationen zur rechtlichen Haftung bei Datenlieferungen in den allgemeinen Geschäftsbedingungen stehen müssen.

### Benutzungsoberflächen

In der letzten Session des ersten Tages stand der Endnutzer im Fokus. Die Vortragenden wiesen ausnahmslos darauf hin, dass offene Daten nur dann sinnvoll sind, wenn sie auch entsprechend aufbereitet und visualisiert werden. Damit auch »normale« Endnutzer davon profitieren können, müssen sie über entsprechende Tools verfügen, wie z. B. Visualbox<sup>7</sup> und OpenDataVis<sup>8</sup>. Ziel müsse es sein, dass ausgehend von einfach zu erstellenden Daten wie CSV oder KML gute Visualisierungen ohne tiefe technische Vorkenntnisse erstellt werden können. Den Abschluss lieferte Benedikt Groß vom Royal College of Arts mit einigen sehr ansprechenden Visualisierungen basierend auf Open Data: Zum einen hatte er NASA-Daten über den Anstieg des Meeresspiegels auf eine Europakarte übertragen, zum anderen zeigte er anhand einer London-Karte und dem Plan der Londoner U-Bahn, dass es manchmal notwendig ist, komplexe Datentransformationen durchzuführen, damit die Daten auch wirklich zueinander passen (der Plan der Londoner U-Bahn ist geografisch nicht korrekt, sondern musste transformiert werden, bevor er auf die »richtige« Karte passte).

### Hauptkonferenz, Tag 2

Den Auftakt am Mittwochmorgen bildeten unterschiedliche Themen, die die Organisatoren zwar interessant fanden, die jedoch nicht wirklich in eine der anderen Sessions passten; entsprechend groß war die Spannweite der Beiträge. Neben einem Vortrag des Verfassers über die Lizenzierung der Metadaten der Deutschen Nationalbibliothek unter CC0<sup>9</sup> stellte Fiona Nielsen von DNA Digest den Plan einer weltweiten, offen lizenzierten Genomdatenbank vor – es gibt zwar viel Daten über Genome, die meisten sind aber nicht offen. Florian Bauer von Reegle referierte über die Bedeutung von offenen Daten über Energienetze, um die Energiewende vorantreiben zu können, und Dan Brickley stellte die Vorteile von schema.org für eine verbesserte Integration der Daten in Suchmaschinen dar. Über einen Fall, bei dem offene Daten dem realen Bevölkerungsschutz gedient haben, berichtete Shuichi Tashiro von der Information-technology Promotion Agency (IPA) in Japan. Eins der Datasets, die die japanischen Behörden zur Verfügung stellen, sind die

Daten des Erdbebenwarnsystems. Die Eisenbahn wertet diese Daten in Echtzeit aus; somit konnten bei dem letzten größeren Erdbeben sämtliche Hochgeschwindigkeitszüge rechtzeitig zum Stillstand gebracht werden, so dass keiner der Züge entgleist ist. Open Data rettet Leben ...

### Produktdaten

Eine Reihe von kürzeren Präsentationen fokussierte offenen zugängliche Produktdaten. Die Vortragenden, die von NXP Semiconductors, Tesco und der Plattform Product Open Data kamen, waren sich alle darüber einig, dass es für die Verbraucher gut ist, wenn sie Zugriff auf strukturierte Produktinformation haben. Dies gilt sowohl für Information über die technischen Eigenschaften (Spannungen, Schaltfrequenzen etc.) als auch für Allergene in Lebensmitteln (sind in dem Kuchen Nüsse drin und wenn ja, welche?). In der Diskussion wurde darauf hingewiesen, dass die Freigabe der Daten auch Marketingzwecken dienen könne und Potenzial für neue Dienste berge (z. B. welche Läden in meiner Nähe verkaufen laktosefreie Milch?). Interessant ist dabei auch der Ansatz von Product Open Data, die gTin-Codes (die von GS1 standardisiert werden) als URI zu kodieren, damit Produkte eindeutig referenziert und die Daten durch die Verwendung eines Barcode-scanners abgerufen werden können.

### Suche (Discovery) und Crowd Wisdom

Nach einer kurzen Session über den Nutzen von guten Identifiern und wohlgedachten Namenskonventionen sowie deren Bedeutung für die Anwendungen, die auf die Daten zugreifen und sie verarbeiten, wurde das Thema Suche (Discovery) in einer Podiumsdiskussion aufgegriffen. Hier wurde auf die Wichtigkeit von »Small Data« verwiesen: Wenn es z. B. um lokale Information geht, ist es häufig viel effizienter, die Datasets von kleineren Anbietern auszuwerten. Dies erfordert zwar strukturierte und gut greifbare Daten (z. B. über die Webseite), aber gerade durch das Einbetten von schema.org in HTML durch RDFa ist das (technisch) kein größeres Problem. Da inzwischen fast alle Suchen über die großen Suchmaschinen laufen, sollte es ein Hauptanliegen der Anbieter sein, ihre Information dort einzubringen, denn die Bedeutung von manuell gepflegten Listen etc. wird rapide abnehmen.

Die letzten Präsentationen waren dem Thema Crowd Wisdom gewidmet. Vorgestellt wurden z. B. die griechische Plattform publicspending.gr, wo öffentliche Ausgaben verzeichnet werden, inklusive der Firmennamen, die die Aufträge erhielten, und die der Entscheider, die diese Ausgaben genehmigten. Vor allem Journalisten greifen für ihre Hintergrundartikel gerne auf

diese Informationen zurück. In Japan ist ein großer Teil der Open-Data-Aktivitäten der Regierung von einzelnen Communities getrieben, die Vorschläge liefern, welche Daten die Verwaltung als nächstes öffnen soll. Das Thema Social Media fehlte natürlich auch nicht. Hier wies Deirdre Lee von DERI darauf hin, dass die öffentliche Verwaltung zwar Social-Media-Plattformen zur Informationsverbreitung einsetzt, die Industrie jedoch deutlich besser ausgerüstet ist, wenn es darum geht, die dort verfügbare Information auch auszuwerten, um z. B. Stimmungslagen frühzeitig wahrzunehmen.

### Fazit

Um Bibliotheksdaten für Communities außerhalb des Bibliothekswesens interessant zu machen, muss es Verknüpfungspunkte geben, an denen diese Communities einsteigen können. Ein Beispiel sind Geokoordinaten: Das Projekt GeoKnow sieht gute Möglichkeiten, Norm- und Titeldaten mit Geokoordinaten zu verwerthen. Weitere Chancen bieten auch Themenportale, die Ressourcen zu bestimmten Themengebieten sammeln: hier können Bibliotheken mit Verschlagwortung und Klassifikation (halb-)automatische Themencluster anbieten. Wichtig sind natürlich auch die Normdaten, hierbei ist jedoch zu bedenken, dass der Umfang der GND zu gering sein könnte oder nicht die relevanten Daten (z. B. Autoren wissenschaftlicher Aufsätze) enthält. Fest steht aber, dass die Qualität der Links, die in den Bibliotheksdaten vorhanden sind, mit der Erschließung steht und fällt.

Viele Bibliotheken stellen inzwischen ihre Daten unter offenen Lizenzen zur Verfügung, und diese Entwicklung wurde sehr positiv aufgenommen. Es ist durchaus akzeptiert, dass nicht alle Daten freigegeben werden, sondern offene Daten (von außen) mit geschlossenen internen Daten gemischt werden, um die eigenen Geschäftsprozesse voranzutreiben. Offene Daten scheinen eine sichere Zukunft zu haben.

- 1 [www.w3.org/2013/04/odw/](http://www.w3.org/2013/04/odw/)
- 2 [www.w3.org/2013/04/odw/papers](http://www.w3.org/2013/04/odw/papers)
- 3 <http://data.gov.uk/data/openspending-report/>
- 4 Komplette Agenda unter [www.w3.org/2013/04/odw/agenda](http://www.w3.org/2013/04/odw/agenda)
- 5 Vgl. z. B. den Vorschlag über Linked CSV von Jeni Tennison: <http://jenit.github.io/linked-csv/>
- 6 <http://datalift.org/en>
- 7 <http://alangrafu.github.io/visualbox/>
- 8 <https://github.com/alangrafu/pendatavis>
- 9 [www.w3.org/2013/04/odw/odw13\\_submission\\_57.pdf](http://www.w3.org/2013/04/odw/odw13_submission_57.pdf)

### DER VERFASSER

**Dr. Lars G. Svensson**, Referent für Wissensvernetzung in der Abteilung Informationstechnik, Deutsche Nationalbibliothek, Adickesallee 1, 60322 Frankfurt am Main, E-Mail: [l.svensson@dnb.de](mailto:l.svensson@dnb.de)

## BERICHT ZUR 34. ASPB-TAGUNG VOM 11. BIS 13. SEPTEMBER 2013 IN KIEL

Vom 11. bis 13. September 2013 fand in Kiel die 34. ASPB-Tagung ([www.aspb2013.de](http://www.aspb2013.de)) statt. Der Titel der diesjährigen Konferenz lautete: »Leinen los! Innovationen und strategische Turn Arounds in Spezialbibliotheken«. Veranstaltet wurde die Tagung in Kooperation mit der Deutschen Zentralbibliothek für Wirtschaftswissenschaften – Leibniz-Informationszentrum Wirtschaft (ZBW).

Die 34. ASPB-Tagung in Kiel startete am 11. September vor der offiziellen Eröffnung mit drei Workshops. Dr. Ingeborg Rubbert, Coach und Organisationsberaterin, stellte in ihrem Workshop »Selbstmarketing in Bibliotheken« Instrumente für insbesondere Institutsbibliotheken und OPLs (One-Person-Libraries) vor, um die Leistungen ihrer Bibliothek vor ihren Stakeholdern ins rechte Licht rücken können. Jan Lüth, Leiter des GIGA Informationszentrums in Hamburg und Teilnehmer des Workshops sagte nach dem Kurzseminar: »Der Workshop ›Selbstmarketing‹ bot eine gute Möglichkeit, die eigenen Strategien der Selbstvermarktung weiter zu entwickeln. Denn wissenschaftliche Spezialbibliotheken, wie das GIGA Informationszentrum, haben sich längst zu modernen Dienstleistern weiterentwickelt, werden aber als solche noch nicht in ausreichendem Maße wahrgenommen.« Dr. Anna Maria Koeck und André Vatter von der ZBW führten in einem weiteren Workshop in das Thema Social Media für Bibliotheken ein und diskutierten mit den Workshopteilnehmerinnen und -teilnehmern die Möglichkeiten der unterschiedlichen Instrumente für mehr Sichtbarkeit bei Kundinnen und Kunden. Birgit Fingerle, Innovationsmanagerin der ZBW, erörterte in ihrem Workshop mit den Teilnehmenden vor allem, wie gute Ideen, an denen es in der Regel ja nicht mangelt, auch in die Tat umgesetzt werden können. Die Anwesenden bekamen Anregungen zur systematischen Förderung ihrer Innovationsfähigkeit an die Hand. Fabian Gail, Projektassistent »Strategie- und Neuorganisationprozess ZB MED« äußerte: »Das offene Format ist der Vielfalt der vertretenen Institutionen gerecht geworden und hat gezeigt, aus wie vielen Winkeln das Thema ›Innovation‹ beleuchtet werden muss. Dabei ist meiner Ansicht nach klar geworden, dass sich Kreativität zwar nicht zwingen, aber organisieren lässt: durch die gezielte Schaffung von Freiräumen und Verfahren, aus denen auch Wertschätzung für das schöpferische Potential aller Mitarbeiterinnen und Mitarbeiter spricht.«

Innovationsmanagement, Strategieentwicklung und Technologieentwicklung – das waren dann auch die Themen, die die Tagung selbst mit insgesamt 17

**Industrie nutzt  
Social Media besser**

**Qualität steht und fällt  
mit der Erschließung**

**Innovationsmanagement,  
Strategie- und  
Technologieentwicklung**