

Grundlagen der Informationssuche, Informationsvisualisierung und Informationsverarbeitung für die Integration von interaktiven Visualisierungen in die Websuche

Hienert, Daniel

Veröffentlichungsversion / Published Version

Zeitschriftenartikel / journal article

Zur Verfügung gestellt in Kooperation mit / provided in cooperation with:

GESIS - Leibniz-Institut für Sozialwissenschaften

Empfohlene Zitierung / Suggested Citation:

Hienert, Daniel: Grundlagen der Informationssuche, Informationsvisualisierung und Informationsverarbeitung für die Integration von interaktiven Visualisierungen in die Websuche. In: *Historical Social Research* 39 (2014), 3, pp. 193-285. DOI: <http://dx.doi.org/10.12759/hsr.39.2014.3.193-285>

Nutzungsbedingungen:

Dieser Text wird unter einer CC BY-NC Lizenz (Namensnennung-Nicht-kommerziell) zur Verfügung gestellt. Nähere Auskünfte zu den CC-Lizenzen finden Sie hier:
<https://creativecommons.org/licenses/by-nc/4.0/deed.de>

Terms of use:

This document is made available under a CC BY-NC Licence (Attribution-NonCommercial). For more information see:
<https://creativecommons.org/licenses/by-nc/4.0>

Grundlagen der Informationssuche, Informations- visualisierung und Informationsverarbeitung für die Integration von interaktiven Visualisierungen in die Websuche

Daniel Hienert*

Abstract: »Foundations of Information Retrieval, Information Visualization and Information Processing for the Integration of Interactive Visualizations in the Web Search«. This article gives an overview of the foundations in the areas of Information Search, Information Visualization and Information Processing. They form the basis for developing the model in the previous article (Hienert 2014). The field of *Information Search* provides various models which describe how users search for information. Several characteristics and methods are presented which are part of the search process. A major challenge is the heterogenous information basis on the Web. The section *Information Visualization* describes the goals, benefits, processes, models and techniques of interactive visualizations. It can be shown that interactive visualizations are beneficial for the representation of large and complex information collections. The following section *Information Processing* shows how users process information. For this purpose, basic mechanisms of cognitive processing and properties are presented. Based on this, the process of cognitive processing of visualizations is described. Interactive Visualizations can expand the cognitive process in which an ongoing exchange between external and internal representation takes place. This article is a shortened and revised summary of the chapter in the dissertation of Hienert (2013).

Keywords: Information visualization, information retrieval, information processing.

1. Einleitung

Dieser Artikel behandelt die Grundlagen, die für die Modellbildung im vorherigen Artikel (Hienert 2014) benötigt werden. Eigenschaften, Modelle und Methoden der Informationssuche sind wesentlich für die Einbettung von interaktiven Visualisierungen in den Suchprozess. Der Abschnitt stellt die verschiedenen Aspekte und Eigenschaften der Informationssuche auf Metaebene vor. Der zweite große Bereich ist die Informationsvisualisierung, auf der interaktive

* Daniel Hienert, GESIS – Leibniz Institute for the Social Sciences, Unter Sachsenhausen 6-8, 50667 Cologne, Germany; daniel.hienert@gesis.org.

Visualisierungen beruhen. Im Gegensatz zur Informationssuche hat die Informationsvisualisierung eigene Ziele, Eigenschaften, Methoden und Ansätze, um Information verstehbar und durchsuchbar zu machen. Beide Bereiche sind als Forschungszweige noch stark getrennt, obwohl sie beide dazu dienen, Nutzern Informationen leichter zugänglich zu machen. Der Abschnitt Informationsverarbeitung zeigt Strukturen, Modelle und Prozesse, wie Informationen und Visualisierungen beim Nutzer verarbeitet werden. Es ist erkennbar, dass Nutzer mit heterogenen Informationen und Darstellungsweisen sehr gut zurechtkommen, aber diese Heterogenität durch erhöhten Zeitaufwand im Verarbeitungsprozess belastet ist.

2. Informationssuche

Der Abschnitt *Informationssuche* zeigt, welche verschiedenen Modelle, Methoden und Eigenschaften die Informationssuche hat. Dafür werden zuerst verschiedene Suchmodelle ausgehend vom klassischen IR-Modell, über Berrypicking, zu Exploratory Search und der Suche im Web vorgestellt. Alle Modelle versuchen, den Suchprozess zu modellieren und zeigen verschiedene Eigenschaften der Informationssuche auf. Verschärft werden die Anforderungen an die Informationssuche durch die Eigenschaften von Informationen im Web. Ein erster Lösungsansatz dazu soll der Ansatz des Semantic Web bieten. Für die Suche werden verschiedene Retrieval-Techniken angewandt, die hier vorgestellt werden und später in den Bereich der Informationsvisualisierung transformiert werden sollen.

2.1 Modelle

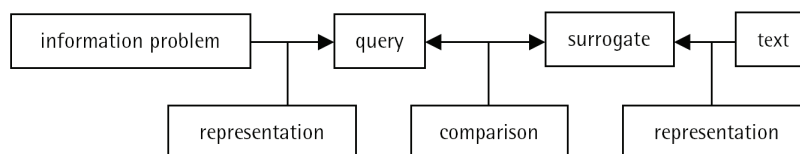
Für den Prozess der Informationssuche existieren verschiedene Modelle auf unterschiedlichen Abstraktionsschichten. Der folgende Abschnitt führt verschiedene Konzepte wie das klassische IR-Modell, das Berrypicking-Modell, den Exploratory Search-Ansatz, die Informationssuche im Web und den Ansatz des Semantic Webs auf und zeigt, welche spezifischen Eigenschaften der Informationssuche die Modelle herausarbeiten und damit die Informationssuche insgesamt hat. Die Modelle beziehen sich zudem auf verschiedene Informationstypen als Grundlage, die sich beispielsweise auf Text oder Faktendaten konzentrieren. Dabei werden die Konzepte nur auf abstrakter Ebene vorgestellt, um eine Gesamtübersicht der verschiedenen Modelle zu erhalten.

2.1.1 Information Retrieval-Modelle

Das prinzipielle Konzept des Information Retrievals (IR) stellen Belkin und Croft (1987) dar (vgl. Abbildung 1). In diesem Modell wird das Informationsbedürfnis eines Anwenders in einer Anfrage repräsentiert und mit der Reprä-

sensation von Dokumenten aus einer Datenbank verglichen. Das Information Retrieval-System (IR-System) gibt dann die Dokumente aus, die am ehesten der Anfrage entsprechen. Das klassische IR-Modell ist dabei stark verbunden mit dem Cranfield-Paradigma, das die Effektivität von IR-Systemen und Algorithmen für das Finden relevanter Dokumente anhand von Messgrößen wie Precision und Recall misst. Dabei wird ausgegangen von einer festen Dokumentsammlung, einem feststehenden Informationsbedürfnis (verschiedenen Topics) und Relevanzbeurteilungen durch Juroren (vgl. Baeza-Yates und Ribeiro-Neto 1999).

Abb. 1: Prinzip des Information Retrieval



Nach Belkin und Croft (1987, Abb. 1).

Belkin und Croft gehen vornehmlich von Dokumenten als Datenbasis aus:

the means for identifying, retrieving, and/or ranking texts (or text surrogates or portions of texts) in a collections of texts, that might be relevant to a given query (Belkin und Croft 1987, 109).

Die Datenbasis wird also durch Textdokumente gestellt, die mithilfe automatischer oder manueller Indexierung repräsentiert werden. Van Rijsbergen (1979) nimmt anhand Tabelle 1 eine Abgrenzung von Data Retrieval (Faktenretrieval) zu Information Retrieval vor.

Tab. 1: Information Retrieval versus Data Retrieval

	Data Retrieval	Information Retrieval
Matching	Exact match	Partial match, best match
Inference	Deduction	Induction
Model	Deterministic	Probabilistic
Classification	Monosthetic	Polythetic
Query Language	Artificial	Natural
Query specification	Complete	Incomplete
Items wanted	Matching	Relevant
Error response	Sensitive	Insensitive

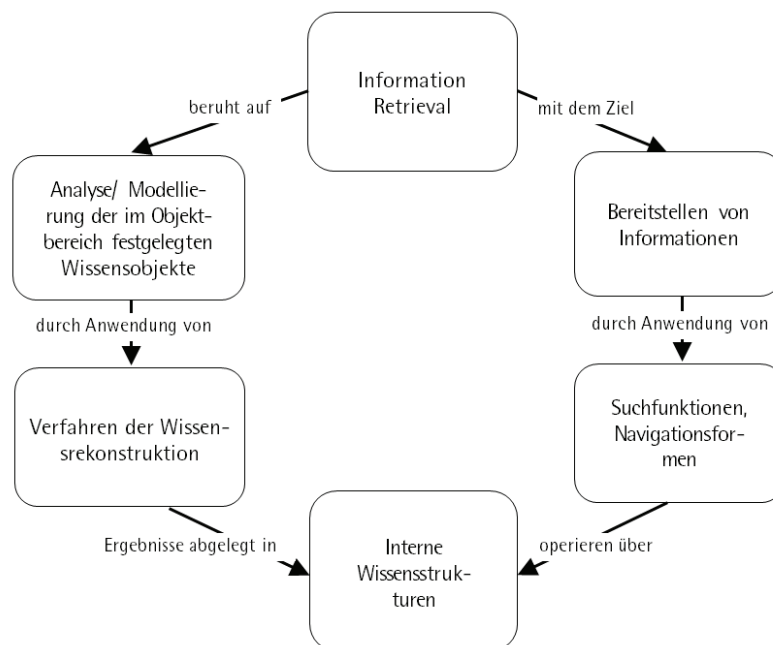
Aus van Rijsbergen (1979).

Van Rijsbergen geht dabei von einer dedizierten Anfrage in einer Datenbank-abfragesprache aus. Diese berücksichtigt alle in der Tabelle aufgeführten Punkte wie zum Beispiel Vollständigkeit und eindeutiges Ergebnis. Nicht berücksichtigt wird jedoch bei dieser Aufstellung, dass auch bei einer Suche nach

Faktendaten ein vages Informationsbedürfnis des Anwenders bestehen kann und sich die explizite Anfrage erst durch den iterativen Retrievalprozess bildet.

Kuhlen (1995) vereint auf abstrakter Ebene die unterschiedlichen Ansätze von Informationssuche und geht von einem allgemeinen Grundmodell des Information Retrievals aus (vgl. Abbildung 2). Daraus leitet er Modelle für das Dokumentenretrieval, Faktenretrieval, Information Retrieval als Hypertext und Information Retrieval als wissensbasiertes System ab. Datenbasis können hier also nicht nur Dokumente, sondern auch Faktendaten, Volltexte, Multimediaobjekte, bibliographische Angaben, semantische Netze usw. sein. Die Modelle des Information Retrievals sind dabei immer zweigeteilt: Auf der linken Seite der Abbildung finden sich die Prozesse der Inhaltserschließung, Modellierung und Wissensrepräsentation, auf der rechten Seite die Prozesse des Retrievals, wie Suche und Navigation (vgl. Kuhlen 1995, 276).

Abb. 2: Allgemeines Modell des Information Retrieval (nach Kuhlen 1995, Abb. 7-1)



Nach Kuhlen (1995, Abb. 7-1).

Auch die Fachgruppe Information Retrieval der Gesellschaft für Informatik gibt eine Definition für den Begriff Information Retrieval:

Im Information Retrieval (IR) werden Informationssysteme in Bezug auf ihre Rolle im Prozess des Wissenstransfers vom menschlichen Wissensproduzen-

ten zum Informations-Nachfragenden betrachtet. Die Fachgruppe „Information Retrieval“ in der Gesellschaft für Informatik beschäftigt sich dabei schwerpunktmäßig mit jenen Fragestellungen, die in Zusammenhang mit vagen und unsicheren Wissen entstehen. Vage Anfragen sind dadurch gekennzeichnet, dass die Antwort a priori nicht eindeutig identifiziert ist. Hierzu zählen neben Fragen mit unscharfen Kriterien insbesondere auch solche, die nur im Dialog iterativ durch Reformulierung (in Abhängigkeit von den bisherigen Systemantworten) beantwortet werden können; häufig müssen zudem mehrere Datenbasen zur Beantwortung einer einzelnen Anfrage durchsucht werden. Die Darstellungsform des in einem IR-System gespeicherten Wissens ist im Prinzip nicht beschränkt (z. B. Texte, multimediale Dokumente, Fakten, Regeln, semantische Netze). Die Unsicherheit dieses Wissens resultiert meist aus der begrenzten Repräsentation von dessen Semantik (z. B. bei Texten oder multimedialen Dokumenten); darüber hinaus werden auch solche Anwendungen betrachtet, bei denen die gespeicherten Daten selbst unsicher oder unvollständig sind (wie z. B. bei vielen technisch-wissenschaftlichen Datensammlungen). Aus dieser Problematik ergibt sich die Notwendigkeit zur Bewertung der Qualität der Antworten eines Informationssystems, wobei in einem weiteren Sinne die Effektivität des Systems in Bezug auf die Unterstützung des Benutzers bei der Lösung seines Anwendungsproblems beurteilt werden sollte (Fuhr 1996).

Auch in dieser Definition wird die Datenbasis von rein textuellen Daten auf multimediale Dokumente, Fakten, Regeln und semantische Netze erweitert. Weiterhin wird in dieser Definition ein Schwerpunkt auf die vage Anfrageformulierung gelegt. Die Anfrage durch den Anwender ist nicht a priori fest vorgegeben, sondern entsteht erst nach und nach im Prozess des iterativen Retrievals. Das Informationsbedürfnis des Anwenders steht also im Mittelpunkt und die Bewertung des IR-Systems muss durch die Qualität der Antworten erfolgen.

2.1.2 Berrypicking

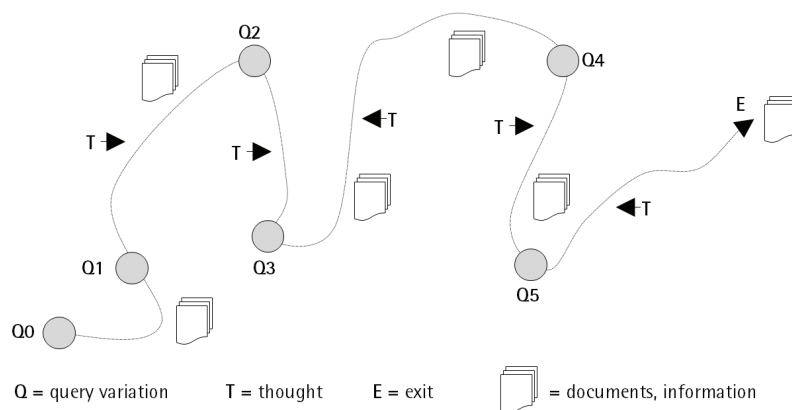
Bates (1989) schlägt das *Berrypicking-Modell* für die Suche in Online- und anderen Informationssystemen vor, das auch sehr gut den Prozess des Web-Browsing abbildet. Das Modell soll näher am echten Verhalten von Informationssuchenden sein als das klassische IR-Modell. Bates geht dabei von folgenden Limitierungen des klassischen IR-Modells aus:

- Das klassische IR-Modell deckt nur einen Teil der Suchprozesse der Wirklichkeit und diese nur unvollständig ab.
- Es ist ein formales Modell, das gut für die wissenschaftliche Evaluation von verschiedenen Retrieval- und Ranking-Modellen ist, aber nicht die Wirklichkeit repräsentiert.
- Die Anfrage wird als einzige, einheitliche nicht veränderliche Auffassung des Informationsbedürfnisses behandelt.
- Bei der Anfrage-Repräsentation stellt sich die Frage, wieso die Anfrage als passend zum System formuliert werden muss und nicht umgekehrt.

- Bei der Dokumenten-Repräsentation machen es Fortschritte in der Informatik möglich, Volltexte zu repräsentieren und zu finden; auf klassische kontrollierte Vokabulare kann weitgehend verzichtet werden.

Abbildung 3 zeigt den Verlauf des Berrypicking-Prozesses. Zu Beginn steht wie beim klassischen IR-Modell das Informationsbedürfnis eines Nutzers. Durch eine Abfolge von Denkprozessen, Aktionen und resultierenden Dokumenten und Informationen innerhalb des Suchprozesses wird die Anfrage immer weiter adaptiert bis zum Abschluss des Prozesses die gewünschten Informationen gefunden wurden. Der wellenartige Prozesspfeil symbolisiert dabei die Veränderungen einer sich entwickelnden Suche. In jedem Prozessschritt führen Anfragen zu Dokumenten und Informationen, die durch Denkprozesse wieder zu adaptierten Anfragen und neuen Dokumenten führen, bis das Informationsbedürfnis des Nutzers befriedigt ist.

Abb. 3: Der Berrypicking-Prozess



Nach Bates (1989, Abb. 2).

Das Berrypicking-Modell erweitert in vier Bereichen das klassische IR-Modell (Bates 1989, 421):

- 1) Die Art der Suchanfrage entwickelt sich im Suchprozess und wird immer wieder adaptiert.
- 2) Der Suchprozess verhält sich nach dem Berrypicking-Muster wie oben beschrieben, anstatt dass er wie im IR-Modell zu einem einzigen, festen Suchergebnis führt.
- 3) Die Bandbreite der Suchtechniken ist größer und die eingesetzten Techniken wechseln im Suchprozess. Bates listet als Beispiele verschiedene Techniken wie „footnote chasing“, „citation searching“, „journal run“, „area scanning“, „subject searches in bibliographics and abstracting and indexing service“ und „author searching“ auf. Nutzer wechseln die Strategie, je nach der mo-

mentanen Anforderung. Wechselt die Anforderung teilweise oder ganz, wird auch die Suchstrategie adaptiert.

- 4) Die Quellen, in denen gesucht wird (Dokumente, Informationen, Zitationen etc.), ändern sich sowohl in Form als auch im Inhalt.

2.1.3 Exploratory Search

Eine Weiterentwicklung von Berrypicking ist *Exploratory Search* (Marchionini 2006; White u.a. 2006). Exploratory Search kombiniert Anfrage- und Browsing-Strategien, um den Lern-, und Untersuchungsschritt im Suchprozess zu fördern. Es grenzt sich damit vom reinen Retrieval oder Lookup ab, das auf analytischen Strategien wie präzise gestellten Anfragen beruht. Exploratory Search unterscheidet drei Suchaktivitäten: (1) Lookup, (2) Learn und (3) Investigate, wobei sich Exploratory Search hauptsächlich auf die letzten beiden Schritte fokussiert.

Lookup ist dabei die Sammlung elementarer Suchaktivitäten und der vorherrschende Untersuchungsgegenstand der letzten Dekaden. Vertreten ist sie zum Beispiel durch Faktenretrieval in Datenbanksystemen oder durch die Suche mit Websuchmaschinen. Dabei werden durch den Nutzer Anfragen gestellt und die Systeme liefern diskrete und strukturierte Objekte wie Texte, Dokumente, kurze Beschreibungen oder andere Medientypen zurück. In dieser Art der analytischen Suche liefern sorgfältig spezifizierte Anfragen präzise Resultate zurück, bezüglich derer auf Nutzerseite nur wenig Bedürfnis besteht, diese aufwendig weiter zu untersuchen oder zu vergleichen.

Learn erweitert diese elementare Suchaktivität um einen weiteren Schritt. Das Web hat sich zur primären Informationsquelle entwickelt und verlangt durch seine heterogenen Inhalte und komplexeren Informationsbedürfnisse der Nutzer nach einer Erweiterung des Suchprozesses. Informationsobjekte liegen in verschiedenen Medientypen und Modalitäten vor, wie zum Beispiel Text, Zahlen, Diagrammen, Karten, Videos, Bilder. Ein Großteil der Zeit im Suchprozess wird darauf verwendet die Resultate einer Suche zu untersuchen und zu vergleichen und daraus folgernd die Suchanfragen anzupassen. Der Nutzer verbringt dabei viel Zeit mit dem „scannen/betrachten, vergleichen und qualitativ beurteilen“ mit dem Ergebnis des „Erwerbs von Wissen, Verstehen von Konzepten oder Fähigkeiten, der Auslegung von Ideen und Vergleichen oder Zusammenfassungen von Daten und Konzepten“ (Marchionini 2006). Die Informationsobjekte verlangen also multiple Iterationen und kognitive Verarbeitung und Interpretation, wobei durch einen Lernprozess die Bildung von neuem Wissen generiert werden soll.

Investigate ist der dritte Schritt im Suchprozess. Er verlangt mehrere Iterationen über längere Zeitabstände. Die Informationen werden kritisch überprüft, bevor sie in persönliche oder professionelle Wissensbestände integriert werden. Es benötigt vorhandenes Wissen für die Analyse, Synthese und Auswertung der Informationen. Dabei kann neu generiertes Wissen die Planung, Prognose

und die Transformation von existierenden Informationen in das Entdecken von neuen Informationen oder Wissen unterstützen. Auch können dadurch Wissenslücken entdeckt werden.

Der Lern- und Untersuchungsschritt verlangt eine stärkere Beteiligung des Nutzers im Suchprozess. Anstatt den Suchprozess als Übereinstimmung von Anfrage und Dokumentenrepräsentation zu sehen, wird im Feld Human-Computer Information Retrieval (HCIR) eine aktive Beteiligung des Nutzers mit einem Informationsbedürfnis, Informationsfertigkeiten, weiter entwickelten Digitalen Bibliotheken, und lokalen und weltweiten Communities mit dynamisch veränderlichen Aspekten gesehen. Daraus ergibt sich die Anforderung, hoch interaktive Suchsysteme zu gestalten, um den Nutzer aktiv in den Suchprozess zu integrieren. Die Systeme können durch verschiedene interaktive Konzepte und Komponenten unterstützt werden, welche teilweise bereits zur Verfügung stehen. Marchionini nennt hier beispielsweise den Hyperlink-Mechanismus im Web, Query-By-Examples-Interfaces und Techniken wie Dynamic Queries oder Brushing-Techniken.

2.1.4 Informationssuche im Web

Die Suche von Informationen im Web unterscheidet sich stark von der Informationssuche in klassischen IR-Systemen. Lewandowski (2005, 71ff) fasst die Unterschiede zwischen Information Retrieval und Web Information Retrieval zusammen und unterteilt sie in die vier Kategorien: (1) Merkmale des Dokumentenkorpus, (2) Inhalte, (3) Nutzer und (4) IR-System (vgl. Tabelle 2).

Tab. 2: Unterschiede zwischen Web-IR und klassischem Information Retrieval

Unterscheidungsmerkmal	Web	Klassische Datenbanken
<i>Merkmale des Dokumentenkorpus</i>		
Sprachen	Dokumente liegen in einer Vielzahl von Sprachen vor; aufgrund der Volltexterschließung keine einheitliche Erschließung über Sprachgrenzen hinweg.	Einzelne Sprache oder Dokumente liegen in vorher definierten Sprachen vor; Erschließung von Dokumenten verschiedener Sprachen mittels einer einheitlichen Indexierungssprache.
Medienarten	Dokumente in unterschiedlichen Formaten.	Dokumente liegen in der Regel in nur einem Format vor.
Länge und Granularität der Dokumente	Länge der Dokumente variiert, große Dokumente werden oft aufgeteilt.	Länge der Dokumente variiert innerhalb eines gewissen Rahmens; pro Dokument eine Dokumentationseinheit.
Spam	Problem der von den Suchmaschinen unerwünschten Inhalte.	Beim Aufbau der Datenbank wird vorab definiert, welche Dokumente erschlossen werden.

<i>Tabelle 2 Fortsetzung...</i>		
Hyperlink-Struktur	Dokumente sind miteinander verbunden.	Dokumente sind in der Regel nicht miteinander verknüpft; keine Notwendigkeit, aus Verlinkungsstrukturen auf die Qualität der Dokumente zu schließen.
<i>Inhalte</i>		
Datenmenge / Größe des Datenbestands	Genaue Datenmenge nicht bestimmbar; keine vollständige Indexierung möglich.	Genaue Datenmenge aufgrund formaler Kriterien bestimmbar.
Abdeckung des Datenbestands	Abdeckung der Zielmenge unklar.	Abdeckung gemäß dem bei der Planung der Datenbank gesteckten Ziel in der Regel vollständig.
Dubletten	Dokumente können mehrfach/ vielfach vorhanden sein; teils auch in unterschiedlichen Versionen.	Dublettenkontrolle bei der Erfassung der Dokumente. Versionskontrolle in der Regel nicht notwendig, da jeweils eine endgültige Fassung existiert und diese in die Datenbank eingestellt wird.
<i>Nutzer</i>		
Art der Anfragen	Aufgrund heterogener Informationsbedürfnisse der Nutzer sehr unterschiedlich.	Genaue Zielgruppe mit klarem Informationsbedürfnis.
Ill-formed queries	Geringe Kenntnis der Nutzer über angebotene Suchfunktionen/ Recherche allgemein.	Nutzer sind mit der jeweiligen Abfragesprache vertraut.
<i>IR-System</i>		
Interface	Einfache, intuitiv bedienbare Interfaces für Laien-Nutzer.	Oft komplexe Interfaces; Einarbeitung notwendig.
Ranking	Relevance Ranking aufgrund der großen Treffermengen notwendig.	Relevance Ranking aufgrund genau formulierter Suchanfragen und dadurch geringerer Treffermengen meist nicht nötig.
Suchfunktionen	Beschränkte Suchfunktionen	Komplexe Abfragesprachen
Modifikation der Suche	In der Regel nur Möglichkeiten zur weiteren Einschränkung der Suchanfrage.	Umfangreiche Modifikationsmöglichkeiten.
Strukturierung der indexierten Dokumente	Schwache Strukturierung; Feldsuche nur bedingt für die Recherche geeignet.	Starke Strukturierung; Suche innerhalb einzelner Felder gut für die Recherche geeignet.
Auswahl der Dokumente	Abgesehen vom Ausschluss von Spam keine weitere Auswahlkriterien.	Klare Auswahlkriterien werden schon bei der Planung der Datenbank bestimmt.

Aus Lewandowski (2005, Tabelle 5.1).

Die Auflistung von Lewandowski lässt sich in den verschiedenen Kategorien noch erweitern bzw. für den Rahmen dieser Arbeit noch weiter spezifizieren.

Merkmale des Dokumentenkorporus

Die Hyperlinking-Struktur ist eines der wesentlichsten Merkmale des Webs und Grundvoraussetzung für die Retrievaltechnik *Browsing*. Da sich verteilte und inhomogene Datenbestände nur durch Indexierung durchsuchbar machen lassen, ist Adhoc-Retrieval mit intuitiv bedienbaren Benutzungsoberflächen die zweite vorherrschende Retrievaltechnik. Browsing und Adhoc-Retrieval bilden als Verbund die Techniken für das Retrieval im Web (vgl. auch Olston und Chi 2003).

Inhalte

Informationen im Web liegen in (1) verteilten Datenbeständen, (2) heterogenen Modalitäten, (3) unterschiedlicher Strukturiertheit, (4) unterschiedlicher Granularität und (5) unterschiedlicher Qualität vor.

- 1) Daten und Informationen können aus unterschiedlichen Quellen wie Webseiten, Web-Enzyklopädien (bspw. Wikipedia), Webdatenbanken, Web-APIs oder Webarchiven stammen.
- 2) Die Informationen können auch in verschiedenen Modalitäten wie Text, Bild, Video, Audio aber auch in verschiedenen Medienformaten wie HTML-Seiten, als Worddokument, im PDF-, JPG-, GIF-, WAV-, MP3-, FLASH-, DIVX-Format vorliegen.
- 3) Die Informationen sind unterschiedlich strukturiert (vgl. Lewandowski, 2005, 59ff). Webseiten werden über ihren Textinhalt indiziert, Ton- und Bilddokumente werden mit Schlagwörtern ausgezeichnet. Inhalte in Webdatenbanken (z. B. Nachrichtenartikel) sind oft feiner mit Metadaten ausgezeichnet.
- 4) Daten und Informationen können unterschiedliche Granularitäten aufweisen. Dabei kann die Granularität von Rohdaten bis zu komplexen Informationstypen oder Objekten reichen. Informationstypen im Sinne der Informatik sind Entitäten und ihre Beziehungen, die sich durch Auszeichnungssprachen wie Metadaten, Metadatenmodelle, Mikroformate oder Ontologien beschreiben lassen.
- 5) Auch die Qualität der Webressourcen spielt eine wesentliche Rolle für das Ranking und die wahrgenommene Qualität eines Suchergebnisses. Eine große Anzahl an Merkmalen von Webseiten wie extrahierte Features der HTML-Struktur, In- und Outlinks, aber auch das visuelle Design einer Seite können eine Rolle spielen (Mandl 2006a).

Nutzer

Lewandowski (2005) listet „heterogene Informationsbedürfnisse“ und „geringe Kenntnisse über Suchfunktion“ als Merkmale von Nutzern im Web auf. Hinzu kommen zudem Eigenschaften aus dem Information Retrieval wie *Vagheit der Anfrage* oder die *Unsicherheit* im Retrievalprozess.

IR-System

IR-Systeme im Web sind im klassischen Sinne Websuchmaschinen, die nach Indexierung und Relevanzranking die Inhalte des Webs über einfache Suchformulare verfügbar machen. Ausgehend von einer Webseite können sich Nutzer durch die Nutzung von Links zu verwandten Informationen vorarbeiten. Im Sinne dieser Arbeit können auch Visualisierungen Teil eines IR-Systems oder eines übergreifenden IR-Prozesses sein.

Mandl (2006b) ordnet das klassische IR-Modell in den Kontext des Webs ein und zeigt die Herausforderungen, die sich dadurch ergeben. Den Kern bildet wieder der Ähnlichkeitsabgleich einer durch den Informationssuchenden gestellten Anfrage und deren Anfragerepräsentation mit einem Korpus von Webdokumenten und deren Objekteigenschaften. Einige Herausforderungen ergeben sich dabei durch den Webkorpus: Die Inhalte sind stark heterogen im Inhalt und ihrer Darstellung, der Gesamtumfang des Korpus ist unklar und schwer zu erfassen, im Korpus befindet sich durch Linkstrukturen Wissen über Beziehungen, Ähnlichkeiten und Verteilungen. Aufgrund der Größe des wachsenden Korpus erfordert die Indexierung und Repräsentation Heuristiken. Die Anfragen des Informationssuchenden sind in den meisten Fällen sehr kurz. In der Ähnlichkeitsberechnung fließen auch kommerzielle Interessen wie Werbung ein. Durch die Linkstruktur in den einzelnen Dokumenten der Ergebnismenge wird auch eine Fortsetzung des Informationsprozesses durch Browsing ermöglicht.

Eine Herausforderung liegt dabei in der großen Heterogenität der Qualität von im Internet angebotenen Wissensobjekten, deren Bewertung und der Nutzung dieser Bewertung als Faktor für das Ranking von Suchergebnissen im Web Information Retrieval. Dabei kann unterschieden werden zwischen den Begriffen Qualität und Relevanz (Mandl 2006b, 93ff). Als relevant wird ein Informationsobjekt angesehen, wenn es wichtig für ein akutes Informationsbedürfnis ist. Qualität von Informationsobjekten kann dagegen auch unabhängig von einem Informationsbedürfnis bewertet werden. Dabei kann die Relevanz auf Benutzer- oder Systemebene erfasst werden. Auf Benutzerebene wäre es der Abgleich eines kognitiv erfassten und natürlichsprachlich formulierbaren Informationsbedürfnisses, wobei pragmatische Kriterien einen Einfluss haben. Relevanz auf Systemebene wäre der Abgleich einer formulierten Anfrage mit Dokumenten im IR-System auf formaler Ebene, wie es für IR-Evaluationen genutzt wird. Mandl geht davon aus, dass bei großen Qualitätsunterschieden wie im Internet, Relevanzbeurteilungen durch Juroren auch von der Qualität mit beeinflusst werden.

Eingeteilt werden kann die Qualität von Informationsobjekten im Internet in die vier größeren Kategorien intrinsische Qualität, Kontext, Darstellung und Zugang (Mandl 2006b, 98ff). In diese Kategorien lassen sich Aspekte aus (1) umfangreichen Kriterienlisten als auch abstrakteren Definitionen wie (2) Autorität, (3) zeitliche Aspekte, (4) Gebrauchstauglichkeit, (5) wirtschaftliche Aspek-

te, (6) technische und Software-Qualität und (7) interkulturelle Unterschiede einordnen. Kriterienlisten enthalten eine Vielzahl von Kriterien unterschiedlicher Abstraktionsniveaus, können stark kontextabhängig sein und daher für die pragmatische Anwendung kaum anwendbar sein. Die Autorität stellt laut Mandl das zentrale Qualitätsmerkmal für Internetangebote dar. Prominent vertreten ist es durch den PageRank-Algorithmus von Google, der Webseiten anhand der Anzahl und der Qualität der eingehenden Links bewertet. Aber auch als Maß in der Szientometrie mit der Zentralität von Akteuren in einem Netzwerkgraphen spielt es eine wichtige Rolle. Zeitliche Aspekte können in Form der Aktualität der Webseite, des Zeitaufwands für die Lösung eines Informationsbedürfnisses oder der Erscheinung einer Internetquelle im Vergleich zur Printquelle eine Rolle spielen. Die Gebrauchstauglichkeit (Usability) lässt sich laut Mandl in der zentralen Frage zusammenfassen, ob der Nutzer die enthaltenen Informationen überhaupt aufnehmen und rezipieren kann (Mandl 2006a, 106). Die Information und ihre Darstellung müssen so aufbereitet sein, dass der Nutzer sie leicht rezipieren kann. Die Gebrauchstauglichkeit von Informationssystemen wird in den Feldern Mensch-Maschine-Interaktion und Softwareergonomie behandelt und wird durch verschiedene Bereiche wie (1) Formen der Mensch-Maschine-Interaktion, (2) Richtlinien für die Gestaltung, (3) Gestaltungsprinzipien und (4) Ästhetische Gestaltung beeinflusst (Mandl 2006b, 27ff). Verstößt man gegen diese prinzipiellen Erkenntnisse aus dem Bereich der Human-Computer-Interaktion, so kann die Information nur sehr viel schwerer aufgenommen werden. Wirtschaftliche Aspekte können der wirtschaftliche Erfolg eines Angebots basierend auf der Qualität und dem Vertrauen des Nutzers in das Angebot sein. Diese kann in Maßen wie der Benutzungshäufigkeit und dem Rückkehrverhalten gemessen werden. Der finanzielle Aufwand, der für die Erstellung von Internetangeboten investiert wird, kann auch ein Indikator für Qualität sein. Zudem existieren technische und Softwarequalitäts-Aspekte: Qualitätsmerkmale von Software im Allgemeinen, wie sie in DIN- und ISO-Normen ausgedrückt werden, technische Faktoren wie Downloadzeiten, Sicherheit und Zuverlässigkeit von Servern und die korrekte Programmierung von HTML-Seiten im Sinne von vollständiger Korrektheit und Barrierefreiheit. Qualitätskriterien, ihr Einsatz und ihre Häufigkeit scheinen stark kulturabhängig zu sein. So wird die Wichtigkeit der aufgeführten Kriterien in verschiedenen Kulturkreisen unterschiedlich bewertet.

Mandl untersucht das Qualitätsmodell für Websuchmaschinen anhand eines Prototyps. Dazu werden zuerst Ergebnisse von verschiedenen Suchmaschinen abgefragt. Dann werden die Top-Webseiten heruntergeladen, 113 verschiedene Features aus den Bereichen Grafik, visuelles Design und Inhalt extrahiert und daraus ein Qualitätsranking gebildet. Der Einfluss der verschiedenen Features

auf das Qualitätsranking wurde mit Machine-Learning-Algorithmen aus einem Korpus von menschlich qualitativ bewerteten Webseiten gelernt.¹

Für die Messung der Performanz des Qualitätsmodells wurde ein Nutzertest durchgeführt, der Teilnehmer Suchanfragen im Bereich Gesundheit durchführen und dann die Suchergebnisse nach Qualität und Relevanz beurteilen ließ. Es konnte gezeigt werden, dass die mit Qualität gerankten Ergebnisse für die Top-10 Webdokumente deutlich bessere Jurorenurteile für Qualität und Relevanz erhielten als Vergleichsrankings wie das originale Suchmaschinenranking, ein Zufallsranking oder das umgekehrte Qualitätsranking.

2.1.5 Semantic Web

Eine strukturierte Vernetzung von Informationsobjekten im Web bieten die Ansätze des *Semantic Webs* (Berners-Lee, Hendler und Lassila 2001) und *Linked Data* (Bizer, Heath und Berners-Lee 2009). Semantic Web wird als eine Erweiterung des heutigen Webs verstanden. Der *Linked Data*-Ansatz nutzt das Web für die Erstellung von typisierten Links zwischen Daten aus verschiedenen Quellen. Informationsobjekte aus verschiedenen Domänen wie geographische Orte, Personen, Unternehmen, Bücher, wissenschaftliche Publikationen, Filme, Musik-, Fernseh- und Radioprogramme, medizinische Informationen (bspw. Gene, Proteine, klinische Studien), Online-Communities, statistische Daten sollen dabei verbunden werden. Das finale Ziel von Linked Data ist, das Web als eine integrierte globale Datenbank nutzen zu können, die beliebige Informationstypen enthalten kann und eine große Anzahl an verteilten und heterogenen Datenquellen verbindet. Daten werden in Dokumenten in RDF (Resource Description Framework; Lassila und Swick 1999) ausgezeichnet. Das RDF-Modell enkodiert Daten in Form von Tripeln mit einem Subjekt, Prädikat und Objekt. Subjekt und Objekt repräsentieren dabei Entitäten und das Prädikat eine Beziehung zwischen diesen Entitäten. Alle drei Komponenten können dabei auch über einen URI (Uniform Resource Identifier) aufgelöst werden.

Für das Publizieren von Daten wurde ein Satz von Regeln aufgestellt (Berners-Lee 2006):

- 1) URIs stehen als Bezeichner für Dinge.
- 2) Der Einsatz von HTTP-URIs, so dass Nutzer diese Dinge nachschlagen können.
- 3) Beim Aufruf von URIs durch den Nutzer, sollten weitere nützliche Informationen auf der Basis von Standards (RDF, SPARQL) angeboten werden.

¹ Die Suchmaschine Google geht mit dem Panda-Update in 2011 einen ähnlichen Weg und nutzt Qualitätskriterien für das Ranking inklusive Faktoren wie Design, Vertrauenswürdigkeit, Geschwindigkeit und Rückkehrhäufigkeit. Dabei wurden tausende Webseiten durch menschliche Qualitätsprüfer bewertet und die Faktoren und Gewichtung durch Machine-Learning-Algorithmen gelernt <http://en.wikipedia.org/wiki/Google_Panda>.

- 4) Links zu anderen URIs sollten angeboten werden, so dass mehr Dinge entdeckt werden können.

Nach diesen Regeln kann jeder Anbieter Daten zur *Linked Data Cloud*² hinzufügen und mit bestehenden Datensätzen verlinken. Da somit alle Daten verlinkt sind, können Anwendungen wie Linked Data Browser genutzt werden, um zwischen den Datenquellen zu navigieren. Mit Linked Data Search Engines kann nach Tripeln analog zu Websuchmaschinen gesucht werden.

DBpedia (Bizer u.a. 2009) ist dabei ein zentraler Knotenpunkt in der Linked Open Data Cloud. DBpedia extrahiert Informationen von Wikipedia aus strukturierten Elementen wie Infoboxen, Templates oder der Kategorisierung. Sie beschreibt 3,64 Millionen Dinge³ (Stand Juni 2012) in bis zu 97 Sprachen. Dabei werden verschiedene Informationen zu Dingen wie Personen, Orte, Musikalben, Filme, Videospiele, Organisationen, biologische Arten oder Krankheiten extrahiert. Entitäten erhalten einen eindeutigen Identifier und können über eine URL (Uniform Resource Locator) aufgerufen werden und sind überdies über eine Ontologie verbunden, die über 320 Klassen und 1650 Eigenschaften⁴ verfügt. Die Entitäten sind vernetzt mit 2,72 Millionen Links zu Bildern, 6,3 Millionen Links zu externen Webseiten, 6,2 Millionen Links zu externen RDF-Datensätzen und zu anderen externen Quellen. Der Datensatz deckt dabei durch die enzyklopädische Struktur von Wikipedia viele Domänen ab und ist außerdem multilingual.

2.2 Methoden

Für die Informationssuche werden verschiedene Methoden angewandt. Der folgende Abschnitt zeigt die verwendeten Techniken und deren Definition im Information Retrieval, die sich teilweise von der in der Informationsvisualisierung unterscheiden. Die eingeführten Techniken werden später für die Anwendung in der Informationsvisualisierung genutzt und in das Modell aus Kapitel integriert.

2.2.1 Standard IR-Techniken

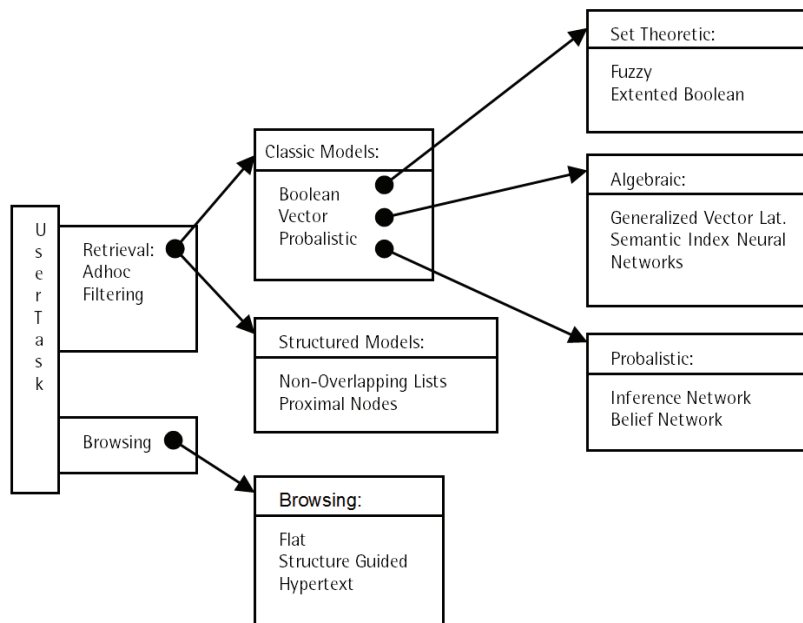
Basierend auf der Klassifikation von Retrievaltechniken von Belkin und Croft (1987, Abb. 2) stellen Baeza-Yates und Ribeiro-Neto (1999) eine Taxonomie der verschiedenen Retrievalmodelle auf. Auf erster Ebene wird zwischen Retrieval und Browsing unterschieden (vgl. Abbildung 4).

² <<http://linkeddata.org/>>.

³ <<http://wiki.dbpedia.org/Datasets>>.

⁴ <<http://wiki.dbpedia.org/Ontology>>.

Abb. 4: Taxonomie der Retrievalmodelle



Nach Baeza-Yates und Ribeiro-Neto (1999, Abb. 2.1).

Beim Retrieval wird zwischen Adhoc-Retrieval und Filtering unterschieden. Bleibt die Datenbasis während des Absetzens von Anfragen relativ stabil, spricht man von Adhoc-Retrieval. Verändert sich die Datenbasis laufend, aber die Anfragen bleiben gleich, so spricht man von Filtering. Klassisches Adhoc-Retrieval wird also beispielweise in der Literatursuche eingesetzt, wo der Anwender explizite Suchen nach Dokumenten ausführt. Sucht ein System anhand eines vom Anwender angelegten Suchprofils, bspw. in Wirtschaftsdaten oder Nachrichtenartikeln, spricht man von Filtering.

Browsing bezeichnet das Blättern in Dokumentenmengen. Wird nur eine Liste der Dokumente angeboten, so spricht man von *flat browsing*. Beim *structure guided browsing* wird der Dokumentenraum anhand inhaltlicher Kriterien strukturiert. Der Anwender kann dann über mehrere Stufen die Dokumentenmenge filtern und gleichzeitig darin zu blättern. *Hypertext-Browsing* ist das Verfahren, bei dem Dokumente durch eine Vielzahl von Links miteinander verbunden sind. Der Anwender kann sich so von Dokument zu Dokument seinen Weg wählen. Der bekannteste Hypertext-Zugang ist das World Wide Web.

Belkin und Croft (1992) stellen folgende charakteristische Merkmale für das Filtering auf:

- Ein Filteringsystem verarbeitet unstrukturierte oder semi-strukturierte Daten. Dies steht im Gegensatz zu einer typischen Datenbank-Anwendung, die strukturierte Daten enthält. Dabei bedeutet Struktur nicht nur, dass die Daten zum Beispiel durch Metaangaben weiter beschrieben sind, sondern auch, dass die einzelnen Datensätze aus einfachen Datentypen bestehen. Zum Beispiel kann eine Datenbank komplexe Dokumente wie Zeitschriftenartikel enthalten, deren Komponenten wie Text und Bild schlechter definiert sind als bspw. Metaangaben zu Autor und Titel. E-Mail-Nachrichten sind ein Beispiel für semi-strukturierte Daten, da sie definierte Header-Felder, aber unstrukturierten Text beinhalten.
- Filtersysteme befassen sich in erster Linie mit Text-Informationen. Unstrukturierte Daten werden häufig als Synonym für textuelle Daten verwendet. Filtersysteme sollten aber auch andere Datenarten wie Bilder, Sprache und Video als Teil von Multimediasystemen behandeln. Die Behandlung dieser Datenarten mit konventionellen Datenbank-Systemen stellt eine Herausforderung dar, da ihre Bedeutung und Beschreibung nur schwer zu extrahieren sind.
- Filtersysteme sind geeignet für große Datenmengen. Typische Anwendungen enthalten große Mengen an Textdaten oder große Datenmengen anderer Medien.
- Filter-Anwendungen behandeln typischerweise Ströme eingehender Daten, entweder von entfernten Quellen, wie News-Anbietern oder direkte Eingänge wie E-Mail. Filtering beschreibt auch den Abruf von Informationen aus Remote-Datenbanken, in denen der eingehende Datenstrom das Ergebnis einer Datenbankabfrage ist. Das beinhaltet auch Systeme, die mithilfe von intelligenten Agenten in entfernten heterogenen Datenbanken suchen.
- Filterung basiert auf der Beschreibung von Interessen von Einzelpersonen oder Gruppen. Die daraus entstehenden Profile vertreten meist langfristige Interessen.
- Filterung impliziert oft, dass nur die Daten, die aus dem Datenstrom extrahiert wurden, für den Nutzer interessant sind. Oft ist aber auch das Gegenteil der Fall und der Datenstrom selbst ist für den Nutzer interessant, wenn die Filterregel zum Beispiel die Beseitigung von Junk-Mail aus eingehenden Emails beschreibt. Nutzer können in ihrem Profil bestimmen, ob sie bestimmte Daten erhalten wollen oder nicht.

Tabelle 3 fasst die Unterschiede zwischen Information Filtering und Information Retrieval zusammen, die Belkin und Croft (1992) beschreiben.

Tab. 3: Abgrenzung von Information Filtering zu Information Retrieval

Information Filtering	Information Retrieval
Wiederholtes Benutzen des Systems von Anwendern mit langfristigen Zielen und Interessen.	Einmaliges Benutzen des Systems mit einmaliger Zielsetzung und einmaliger Anfrage.
Beim Filtering wird davon ausgegangen, dass Filterungsprofile korrekte Spezifikationen von Nutzerbedürfnissen sind.	IR geht von inhärenten Problemen in der Angemessenheit von Anfragen als Darstellungen von Informationsbedürfnissen aus.
Filterung befasst sich mit Verteilung von Texten an Gruppen oder Einzelpersonen.	IR befasst sich mit der Sammlung und der Organisation von Texten.
Filtering befasst sich vor allem mit der Auswahl oder dem Entfernen von Texten aus einem dynamischen Datenstrom.	IR befasst sich in der Regel mit der Auswahl von Texten aus einer relativ statischen Datenbank.
Filterung befasst sich mit langfristigen Veränderungen in einer Reihe von Informationssuchen.	IR befasst sich mit der Antwort auf Nutzerinteraktion innerhalb einer einmaligen Informationssuche.

Zusammengefasst aus Belkin und Croft (1992).

Einen umfangreichen Überblick über die methodischen und technologischen Aspekte von *Information Filtering* geben Hanani u.a. (2001).

2.2.2 Faceted Search und Faceted Navigation

Faceted Search oder *Faceted Navigation* (vgl. English u.a. 2002) ist ein populäres Konzept für die Navigation durch große Informationssammlungen auf der Basis von facettierten hierarchischen Metadaten. Es ist vornehmlich geeignet für große Datenbestände möglichst gleicher Informationselemente wie Webseiten, Produktkataloge und Bild- oder Dokumentensammlungen. Durch die Kombination von Freitextsuche und Browsing erlaubt es auch Nicht-Experten die intuitive Navigation durch sehr große Informationsräume. Nutzer können beispielsweise mit einer Freitextsuche beginnen, dann über Facetten die angezeigte Ergebnismenge einschränken und in der Folge weitere Suchterme zur Anfrage hinzufügen, ohne den Interaktionsprozess zu unterbrechen. Die Nutzung des Konzepts verhindert leere Ergebnislisten und durch die Anzeige der aktuellen Suchbegriffe und der ausgewählten Facetten wird das Gefühl beim Nutzer vermieden im Informationsraum verloren zu sein.

Voraussetzung für die Anwendung ist eine Informationssammlung, die auch Metadaten enthält. Diese Metadaten können facettiert sein, das heißt, es bestehen orthogonale Sets von Kategorien oder Eigenschaften. Zum Beispiel können Fotografien architektonischer Bauwerke mit *Facetten* wie Material (Beton, Backstein, Holz etc.), Stile (Barock, Gotisch, Ming, etc.), Personen (Architekten, Künstler, Entwickler), Ansicht, Orte oder Zeiten versehen sein. Die Facetten können *hierarchisch* („gelegen in Berkeley, Kalifornien, USA“) oder *flach* („von Ansel Adams“) sein. Facetten können *einwertig* („36 cm hoch“) oder *mehrwertig* (verwendet Ölfarbe, Tusche und Wasserfarbe) sein, das heißt, eine Facette kann nur eine oder mehrere Eigenschaften beinhalten (vgl. English u.a. 2002).

English et al. führen einen Nutzertest durch, der zwei verschiedene Designvarianten eines Suchsystems für Fotografien vergleicht und Komponenten eines erfolgreichen Suchsystems evaluiert. Das System enthält 40.000 Fotografien mit 16.000 Metadaten-Termen, die in die neun Facetten Personen, Orte, Strukturtypen, Materialien, Zeiten, Stile, Ansichten, Konzepte und Gebäudenamen unterteilt sind. Die Startseite der Varianten enthalten jeweils ein Suchfeld und eine Übersicht der hierarchischen Facetten auf erster Ebene. Die erste Designvariante *Matrix* zeigt auf der Ergebnisseite die Facetten auf erster Ebene auf der linken Seite an. Im oberen Bereich wird ein Navigationspfad mit den vom Nutzer ausgewählten Facetten angezeigt, mit der Möglichkeit, davon einzelne Facetten wieder zu löschen oder einen einzelnen Pfad auszuwählen. Darunter wird das Suchfeld angezeigt, mit der Möglichkeit auch in den Ergebnissen zu suchen und einzelne Suchterme wieder zu löschen. Unter dem Suchfeld erscheinen die Treffer der Suchanfrage gruppiert. Die Designvariante *Treeview* dagegen zeigt in den Facetten nur noch mögliche Unterkategorien an und wirkt dadurch übersichtlicher. Die Ergebnisse zeigen, dass die Nutzer die Matrixansicht für alle Aufgabentypen präferieren. Interessant ist auch die Verteilung des Zeitaufwands für die Nutzung einzelner Features der Oberfläche: Doppelt so viele Teilnehmer starten mit dem Browsing-Ansatz als mit einer Suche, Suchanfragen werden zudem sehr oft angepasst, um die Ergebnisse einzuschränken. Die Ergebnisse werden oft in der Matrixansicht per Browsing eingeschränkt, in der Treeansicht wird oft in den Ergebnissen gesucht und in beiden Ansichten wird oft die Suche komplett neu gestartet oder einen Prozessschritt zurückgegangen. Der Rest des Zeitaufwands verteilt sich auf Schritte für die Reduktion oder Expandierung des Suchergebnisses.

2.2.3 Information Foraging und Information Scent

Das Konzept *Information Scent* versucht auf der Basis der *Information Foraging*-Theorie (Pirolli und Card 1999) die subjektive Wahrnehmung der Kosten und Nutzen des Zugriffs auf eine Webseite auf der Basis von *Proximal Cues* wie zum Beispiel Web-Links, Icons und umgebenden Text vorherzusagen. Hierbei wird davon ausgegangen, dass Nutzer bei der Suche im Web oft von Webseite zu Webseite über Links navigieren und ihr Navigationsverhalten durch ihr Informationsbedürfnis bestimmt wird. In Chi u.a. (2001) werden zwei berechenbare Modelle für die Zusammenhänge zwischen dem Informationsbedürfnis eines Nutzers und seinen Aktionen vorgestellt: (1) Für ein gegebenes Navigationsmuster wird ein Informationsbedürfnis berechnet, (2) für ein Informationsbedürfnis wird ein wahrscheinliches Navigationsmuster berechnet.

Basierend auf dem Konzept *Information Scent* wurden *Scented Widgets* (Willett, Heer und Agrawala 2007) entwickelt. Scented Widgets sind Standard-Interaktionselemente in Webseiten und Nutzungsoberflächen, die mit Visualisierungen angereichert werden, welche die Navigation in Informationsräumen

erleichtern sollen. Ähnlich der Idee von Tufte's Sparklines (Tufte 1999) werden zu Navigationselementen wie Radio-Buttons, Slider oder Comboboxen *Visual Scents* in Form von Text, Icons, kleinen Säulen- oder Liniendiagrammen hinzugefügt, die Aufschluss über die Verteilung oder Nutzung der Information geben. Nutzer können also direkt beim Interaktionsschritt entdecken, wie sich zum Beispiel die Veränderung eines Sliders auswirkt, da darüber die Verteilung der Information angezeigt wird. Dies kann die Exploration von unbekannten Informationsräumen erleichtern.

2.3 Zusammenfassung

Modelle

Der Abschnitt Informationssuche stellt wichtige Modelle der Informationssuche vor: Das klassische IR-Modell, den Berrypicking-Prozess und Exploratory Search. Das klassische IR-Modell beinhaltet verschiedene Eigenschaften wie das Informationsbedürfnis des Nutzers, die Vagheit der Anfrage und die Unsicherheit der Ressourcen als auch des Nutzers im Retrieval-Prozess. Der Berrypicking-Prozess stellt heraus, dass Nutzer über verschiedene Dokumente und Informationen mit verschiedenen Suchtechniken iterieren und Suchabfragen immer weiter durch Denkprozesse auf Basis von Informationen anpassen, bis sie zum gewünschten Ziel gelangen. Exploratory Search erweitert diesen Ansatz und zeigt, dass neben der elementaren Informationssuche auch Lern- und Untersuchungsschritte eine wichtige Rolle spielen, um die Informationen zur Befriedigung des Informationsbedürfnisses in Wissen zu verwandeln.

Methoden

Nutzer wenden verschiedene Basis-Suchtechniken wie Adhoc-Retrieval, Browsing oder Filtering an, die in IR-Systemen und Websuchmaschinen ergänzt werden durch Techniken wie Faceted Search und Faceted Navigation. Auf Basis der Information Foraging-Theorie verfolgen Nutzer nur die Navigationswege und Links, die sie als wertvoll für die weitere Informationssuche erachten.

Dokumentenbasis

Die Dokumentenbasis wird in der Informationssuche immer komplexer: von Dokumenten über Fakten zu verschiedenartigen Informationsobjekten wie Bilder, Videos, Visualisierungen in verschiedenen Modalitäten und verschiedenen Medienformaten aus verteilten Datenquellen und mit unterschiedlicher Strukturiertheit. Die Granularität der Information reicht dabei von einfachen Rohdaten bis zu komplexen Informationstypen. Das Semantic Web bietet einen Ansatz, um die verteilten Informationsobjekte auf technischer Basis auszuzeichnen, mit URIs eindeutig identifizierbar zu machen und zu verbinden. Hier werden In-

formationsobjekte auf der Basis von Ontologien ausgezeichnet und modelliert und können miteinander verlinkt werden. Dadurch sind die Informationen auf technischer Basis verbunden, aber ein schlüssiges Konzept für die Informationssuche im Sinne des Nutzers, jenseits gängiger Suchmaschinenmodelle, fehlt.

3. Informationsvisualisierung

Die Informationsvisualisierung als eigenständiger Bereich hat zum Ziel große Datenmengen zu visualisieren und Erkenntnisse beim Nutzer zu generieren. Dabei hat sie im Gegensatz zur Informationssuche andere Ziele, Eigenschaften, Methoden und Ansätze, um Information darzustellen und durchsuchbar zu machen. Der folgende Abschnitt untersucht verschiedene Aspekte des Bereichs Informationsvisualisierung: (1) Welche Ziele und Vorteile hat die Informationsvisualisierung, wie können diese Ziele erreicht und gemessen werden und wie ist der Prozess modelliert, um diese Ziele zu erreichen? (2) Wie wird der grundlegende Prozess der interaktiven Informationsdarstellung modelliert? (3) Welche elementaren Informationsstrukturen werden in der Informationsvisualisierung als Hilfe zur Bildung eines kognitiven Modells angenommen und welche Forschungssysteme existieren in den jeweiligen Bereichen? (4) Welche Interaktionstechniken existieren in der Informationsvisualisierung? (5) Welche Modelle existieren für die koordinierte Anzeige mit mehreren Ansichten? (6) Welche Eigenschaften hat die Informationsvisualisierung im Web? (7) Wie werden Visualisierungen bereits in die Informationssuche eingebunden?

3.1 Ziele, Vorteile und Eigenschaften der Informationsvisualisierung

Dieser Abschnitt zeigt zuerst die Ziele, Vorteile und Eigenschaften, wie sie von der Informationsvisualisierung selbst beschrieben werden. Daraufhin erfolgt die Erläuterung allgemeiner Vorteile von Bild und Text. Nachfolgend werden die Eigenschaften subjektiver Wahrnehmung von Bildern, Grafiken und Visualisierungen vorgestellt, ausgedrückt im Modell der Nutzererfahrung.

Card u.a. (1999) definieren Information Visualization (IV) als „*The use of computer-supported, interactive, visual representations of abstract data to amplify cognition*“. Dabei fassen sie zusammen auf welchen Ebenen interaktive Visualisierungen den Kognitionsprozess verstärken können und schlagen sechs Aspekte vor, wie Informationsvisualisierung dies erreichen kann: (1) Die menschliche Speicher- und Verarbeitungsleistung kann durch interaktive Visualisierungen erhöht werden, (2) die Suche nach Informationen verringert sich, (3) die Mustererkennung wird durch visuelle Repräsentationen unterstützt, (4) Inferenzprozesse der Wahrnehmung werden ermöglicht, (5) Aufmerksamkeitsprozesse der Wahrnehmung können für die Überwachung genutzt werden und (6) Informa-

tion werden in ein manipulierbares Medium transferiert. Tabelle 4 gibt einen Überblick über diese Kategorien, ihre einzelnen Aspekte und Erläuterungen.

Tab. 4: Wie Informationsvisualisierung den Kognitionsprozess verstärken kann

Increased Resources	
High-bandwidth hierarchical Interaction	The human moving gaze system partitions limited channel capacity so that it combines high spatial resolution and wide aperture in sensing visual environments (Resnikoff 1987).
Parallel perceptual processing	Some attributes of visualizations can be processed in parallel compared to text, which is serial.
Offload work from cognitive to perceptual system	Some cognitive inferences done symbolically can be recorded into inferences done with simple perceptual operations (Larkin and Simon 1987).
Expanded working Memory	Visualization can expand the working memory available for solving a problem (Norman 1993).
Expanded Storage of Information	Visualization can be used to store massive amounts of information in a quickly accessible form (e.g., maps).
Reduced search	
Locality of Processing	Visualizations group information used together, reducing search (Larkin and Simon 1987).
<i>High data density</i>	Visualizations can often represent a large amount of data in a small space (Tufte 1983).
Spatially indexed addressing	By grouping data about an object, visualization can avoid symbolic labels (Larkin and Simon 1987).
Enhanced Recognition of Patterns	
Recognition instead of recall	Recognizing information generated by a visualization is easier than recalling that information by the user.
Abstraction and aggregation	Visualizations simplify and organize information, supplying higher centers with aggregated forms of information through abstraction and selective omission (Card, Robertson and Mackinlay 1991; Resnikoff 1987).
Visual schemata for organization	Visually organizing data by structural relationships (e.g. by time) enhances patterns.
Value, relationship, trend	Visualizations can be constructed to enhance patterns at all 3 levels (Bertin 1977/1981).
Perceptual inference	
Visual representation make some problems obvious	Visualization can support a large number of perceptual inferences that are extremely easy for humans (Larkin and Simon 1987).
Graphical Computations	Visualizations can enable complex specialized graphical computations (Hutchins 1996).
Perceptual Monitoring	Visualizations can allow for the monitoring of a large number of potential events if the display is organized so that these stand out by appearance or motion.
Manipulable Medium	Unlike static diagrams, visualization can allow exploration of a space of parameter values and can amplify user operations.

Aus Card et al. (1999, Tabelle 1.3).

North (2005) stellt besonders den Punkt des Erkenntnisgewinns heraus. Durch die Fähigkeit des menschlichen Gehirns visuelle Schlussfolgerungen zu ziehen, kann Wissen oder Erkenntnis generiert werden (Card, Mackinlay und Shneiderman 1999). Durch diese Fähigkeit können mentale Modelle der echten Phänomene, die in den Daten repräsentiert werden, gebildet werden (vgl. North 2005, 1222). Dabei können die Erkenntnisse folgendermaßen kategorisiert werden (vgl. North 2005, 1223):

Einfache Erkenntnisse:

- Zusammenfassung: Minimum, Maximum, Durchschnitt
- Finden: Suche nach bekannten Elementen

Komplexe Erkenntnisse:

- Muster: Verteilungen, Trends, Häufigkeiten, Strukturen
- Ausreißer: Ausnahmen
- Beziehungen: Übereinstimmungen, Wechselbeziehungen
- Abwägungen: Gleichgewicht, kombiniertes Maximum/Minimum
- Vergleich: Wahl (1:1), Kontext (1:M), Sets (M:N)
- Cluster: Gruppen, Ähnlichkeiten
- Pfade: Distanz, multiple Verbindungen, Zerlegungen
- Anomalien: Datenfehler

Schierl (2001) bietet eine ergänzende Auflistung von generellen Vorteilen von Bild und Text, die in der Werbung genutzt werden können (vgl. Tabelle 5).

Tab 5: Vorteile von Text und Bild

Vorteile des Bildes	Vorteile des Textes
Hohe Kommunikationsgeschwindigkeit	Eindeutiger als das Bild (kann sich selbst Zusammenhang schaffen)
Fast automatische Aufnahme ohne größere gedankliche Anstrengung	Kann den Leser ansprechen
Wird in der Regel zuerst fixiert	Kann argumentieren und somit wirklich verkaufen
Bildliche Informationsverarbeitung besonders effizient	Kann Schwerpunkte setzen und Einzelaspekte betonen
Einstellungen und Gefühle können subtiler übermittelt werden	Verfügt über Imperativ (Möglichkeit zum Auffordern)
Höhere Glaubwürdigkeit	Kann Argumente (u.U. in anderer Form) in einer Botschaft wiederholen
Höhere Anschaulichkeit (dadurch a. bessere Verstehbarkeit und b. einem Primärerlebnis ähnlicher)	Zeitliche Vorstellungen sind gut zu vermitteln
Platzsparende Information (viel spezifische Information auf wenig Raum)	
Allgemeine Verständlichkeit (auch für Lese- und Sprachunkundige)	
Räumliche Vorstellungen lassen sich gut vermitteln	

Aus Schierl (2001, Tabelle 9.1).

Die Vorteile von Bild aus Tabelle 5 lassen sich wie folgt erläutern:

- *Hohe Kommunikationsgeschwindigkeit:* Innerhalb von Hundertstelsekunden können vom Rezipienten die Grundinformationen bzw. das Thema wiedererkannt werden. Nach ungefähr zwei Sekunden kann ein Bild sicher wiedererkannt werden (Behrens und Hinrichs 1986, 85; zitiert in Schierl 2001). Abhängig von der Textschwierigkeit und Lesefähigkeit liegt die Lese- und Lesegeschwindigkeit bei 4-6 Worten pro Sekunde (Carpenter und Just 1983; zitiert in Schierl 2001) (vgl. Schierl 2001, 229).
- *Fast automatische Aufnahme:* Anzeigen, die überwiegend ihre Information über Bild vermitteln, verlangen weniger Denkvorgänge als über die textliche Vermittlung (Edell und Staelin 1983; zitiert in Schierl 2001). Bilder werden mehr auf holistische Weise aufgenommen und weniger gedanklich analysiert und kontrolliert als sprachliche Mitteilungen (Kroeber-Riel 1985, 124; zitiert in Schierl 2001). Nachteil ist, dass über Bilder Inhalte einfacher an der gedanklichen Kontrolle des Rezipienten „vorbeigemogelt“ werden können (vgl. Schierl 2001, 229).
- *Fixierung:* In Anzeigen- oder Plakatlayout wird das Bild zuerst unabhängig von der Platzierung fixiert. Dies gilt für Anzeigen, als auch für den redaktionellen Teil (Witt 1977; Barton-V. Keitz 1981, 711; Kroeber-Riehl 1985b, 123; Kroeber-Riehl 1993; zitiert in Schierl 2001).
- *Effizienz:* Bildliche Informationen werden effizienter verarbeitet als textliche Information und besonders gut behalten. Die Speicherung im Langzeitgedächtnis ist praktisch unbegrenzt (Schweiger 1985; zitiert in Schierl 2001). Bilder haben ein hohes Aktivierungspotential, können Beziehungen zu persönlichen Erlebnissen aktivieren, Emotionen auslösen und mit Produkten in Beziehung setzen (Behrens und Hinrichs 1986, 88; s.a. Weidemann 1988; Messaris 1997; zitiert in Schierl 2001) (vgl. Schierl 2001, 230).
- *Subtile Übermittlung von Gefühlen:* Mit Bildern können Einstellungen und Gefühle subtiler übermittelt werden als mit Text, z. B. kann die Zartheit und Weichheit von Kosmetik-Tüchern mit einem kleinen Kätzchen als Bildmetapher übermittelt werden (vgl. Schierl 2001, 230f mit Studien von Percy und Rossiter 1980, 166; Lyon 1968, 157ff und weitere).
- *Glaubwürdigkeit:* Bilder erscheinen besonders objektiv und manipulationsunverdächtig (Schifko 1981, 988; zitiert in Schierl 2001). Bilder werden selbstverständlich hingenommen und wirken somit glaubwürdiger als Text (Seyffert 1966, 674ff; zitiert in Schierl 2001). Bilddarstellungen werden heute jedoch wesentlich kritischer betrachtet und überprüft als noch vor wenigen Jahren (vgl. Schierl 2001, 232).
- *Anschaulichkeit:* Ein Bild kann das Primärerlebnis näherbringen, z. B. das Trinken einer Flasche Bier (Percy und Rossiter 1980, 166; Lülfig 1960, 565; zitiert in Schierl 2001).

- *Verständlichkeit*: Bilder haben eine hohe und schnelle Verstehbarkeit, können aber auch international unabhängig von Sprache und Kultur besser von jedermann verstanden werden (vgl. Schierl 2001, 234).
- *Räumliche Vorstellungen*: Viele räumliche Gegebenheiten, Zustände und Beziehungen lassen sich im Grunde nur analog bildhaft darstellen (vgl. Schierl 2001, 232). Zum Beispiel können technische Zeichnungen, Anleitungen, Illustration nur sehr schwer mit Text wiedergegeben werden (Koppelman 1981, 317; Krautmann 1981, 178; zitiert in Schierl 2001).

Dem gegenüber stehen die Vorteile von Text aus Tabelle 5:

- *Eindeutigkeit*: Text kann den Bedeutungszusammenhang, in dem er verstanden werden muss, selbst erklären und ist die einzige Möglichkeit, die funktionelle Darstellung einer Botschaft zu erläutern. Die analog-bildhafte Darstellung von Bild in der Werbung kann dagegen offen sein (vgl. Schierl 2001, 236).
- *Ansprechen des Lesers*: Mit Text kann der Konsument direkt angesprochen werden und damit auch auf bestimmte Zielgruppen und ihre Eigenschaften eingehen (Schierl 2001, 236f).
- *Argumentativ*: Text kann (komplex) argumentieren, Bilder dagegen nicht oder nur indirekt (Schierl 2001, 237f).
- *Schwerpunkt setzen*: Text kann Schwerpunkte setzen und Einzelaspekte weiter erläutern (Schierl 2001, 238f).
- *Imperativ*: Text kann aktiv den Konsumenten/Leser ansprechen und zu Handlungen auffordern (Schierl 2001, 237).
- *Argumente wiederholen*: Argumente in Textform können in variiert Form wiederholt werden und bleiben dadurch besser im Gedächtnis des Angesprochenen (Schierl 2001, 238).
- *Zeitliche Vorstellungen*: Zeitliche Verläufe, komplexe oder abstrakte Dinge, Tatbestände und ihre Beziehungen lassen sich besser in Text darstellen, worin er bildlicher Repräsentation überlegen ist (Schierl 2001, 238).

Eng verbunden mit der subjektiven Wahrnehmung von Bildern, Grafiken und Visualisierungen sind auch Untersuchungen, welche die Korrelation zwischen wahrgenommener Ästhetik und anderen Eigenschaften einer Person im Bereich der Sozialpsychologie (Dion, Berscheid und Walster 1972; Eagly u.a. 1991) oder der Benutzerfreundlichkeit von Nutzungsoberflächen im Bereich HCI (Chawda u.a. 2005; Kurosu und Kashimura 1995; Tractinsky, Katz und Ikar 2000; Tractinsky 1997) untersuchen. Dabei wurde evaluativ gemessen, dass die Wahrnehmung von Benutzerfreundlichkeit eng mit der Ästhetik zusammenhängt und wahrgenommene Benutzerfreundlichkeit anhand ästhetischer Wahrnehmung abgeleitet werden kann. Hassenzahl (2004) führt einige Experimente durch, bei denen sich hedonische Faktoren wie Identifikation und Stimulation mit der Ästhetik eines interaktiven Produkts verbinden lassen.

Einige Studien (Chawda u.a. 2005; Kurosu und Kashimura 1995; Tractinsky, Katz und Ikar 2000; Tractinsky 1997) fanden eine Korrelation zwischen Ästhetik und Benutzerfreundlichkeit vor und nach der Benutzung von Software. Allen Studien gemeinsam ist ein unpersönliches Studienobjekt: Kurosu und Kashimura (1995), Tractinsky u.a. (2000), Tractinsky (1997) beziehen sich auf Nutzungsoberflächen von Geldautomaten, Chawda u.a. (2005) nutzt Visualisierungen von Suchergebnissen. Hassenzahl (2004) dagegen findet keinen Zusammenhang von Ästhetik und Benutzerfreundlichkeit. In seiner Studie finden sich Zusammenhänge zwischen Ästhetik und subjektiven Eigenschaften wie Identifikation und Stimulation. Als Studienobjekte werden Oberflächen für MP3-Player Software genutzt. Diese sind sehr persönlich und bereits stark in verschiedene Richtungen grafisch gestaltet. Dass Hassenzahl in seiner Studie die zwei attraktivsten und die zwei unattraktivsten Oberflächen aus dem Pretest nutzt, verstärkt den Effekt der Identifikation. Ästhetik in den Studien von Kurosu, Kashimura (1995), Tractinsky (1997) und Tractinsky u.a. (2000) besteht dagegen nur in dem Arrangement der Oberflächenelemente der Nutzungsoberfläche eines Geldautomaten. So unterscheiden auch schon Lavie und Tractinsky (2004) zwischen den Begriffen klassischer und expressiver Ästhetik. Dabei bezeichnet klassische Ästhetik ein klares, strukturiertes und sauberes Design, das mit dem Begriff Usability verbunden ist. Gegenüber dem steht der Begriff der expressiven Ästhetik, die ein originelles und kreatives Design beschreibt. So verbindet auch Hassenzahl (2008) den Begriff der expressiven Ästhetik mit der Kategorie „hedonische Qualität – Stimulation“ seines Modells.

In allen Experimenten, die versuchen, Ästhetik mithilfe von vorgetesteten Studienobjekten zu beeinflussen, bleibt letztendlich offen, welche Attribute manipuliert wurden. Es kann unterschieden werden zwischen dem Bottom-up-Ansatz, in dem objektive, wahrnehmbare Faktoren die Ästhetik beeinflussen und Top-down-Ansätzen, wie die subjektive Meinung von Produkten mithilfe von Ästhetik beeinflusst werden kann (Hassenzahl 2004). Auch Tractinsky (2005) bemerkt, dass in einigen Studien (Kurosu und Kashimura 1995; Tractinsky 1997; Tractinsky u.a. 2000; Hassenzahl 2004) nur eine einzige Skala genutzt wurde, um Ästhetik zu beurteilen. Hier müsse feiner differenziert und abgefragt werden, welche Parameter die Wahrnehmung und das Verhalten des Nutzers beeinflussen. Erste Ansätze finden sich in Kim u.a. (2003), die untersuchen, welche spezifischen Elemente eines Designs mit menschlichen Emotionen verbunden sind. Tractinsky und Lowengard (2007) nehmen eine Differenzierung des Begriffs Ästhetik vor. Dabei setzen sie in ihrer Definition Ästhetik mit Schönheit, die auf der subjektiven Wahrnehmung von evaluativ nachgewiesenen Objekteigenschaften basiert, gleich.

Hassenzahl (2004) fordert ein gemeinsames Grundmodell der Nutzererfahrung als Prämisse für eine vergleichbare Forschung. Darauf aufbauend müssen Studien über möglichst differente Studienobjekte geführt werden, damit eine Kategorisierung der Ergebnisse möglich sei. Konsequenterweise entwerfen Hassenzahl und

Tractinsky (2006) ein grobes Modell der *user experience* (UX). Dabei besteht die Benutzererfahrung aus den Facetten *beyond the instrumental* (Usability, hedonische und ganzheitliche Aspekte), *emotion and affect* (1. Emotionen als Auswirkung von Produktnutzung und 2. die Bedeutung von Emotion als Vorgeschichte des Produkts und evaluativen Urteils) und *the experiential* (die räumlich-zeitliche Situation). Die Benutzererfahrung ist demnach ein Zusammenspiel aus dem inneren Zustand (Veranlagung, Erwartungen, Bedürfnisse, Motivation, Stimmung usw.) des Nutzers, den Charakteristiken des Systems (Komplexität, Zweck, Usability, Funktionalität usw.) und der Umgebung, in der die Interaktion auftritt (organisatorische/soziale Einstellung, Sinnhaftigkeit der Tätigkeit, Freiwilligkeit der Nutzung usw.). Dieses komplexe Arrangement führt zu unzähligen Design- und Erfahrungsvarianten, die auch bei der Evaluation von Informationsvisualisierungs-Systemen eine Rolle spielen und die Evaluation komplex gestalten.

3.2 Evaluation in der Informationsvisualisierung

Zusätzlich zu den Herausforderungen der Evaluation von Software oder interaktiven Online-Systemen (vgl. Hegner 2003)) spielt in der Informationsvisualisierung besonders der Aspekt eine Rolle, wie Vorteile, Ziele und Eigenschaften von Visualisierungen (vgl. vorheriger Abschnitt) wie z. B. Erkenntnisgewinn gemessen und damit Systeme für die Informationsvisualisierung angemessen evaluiert werden können. Dabei besteht besonders in der Informationsvisualisierung die Forderung nach komplexeren empirischen Untersuchungen, die nicht nur die Nutzerperformanz mit einfachen statistischen Messgrößen wie Zeit, Mausklicks oder Korrektheit der Aufgabenlösung messen. Vielmehr entsteht die Forderung in Evaluationen auch den Erkenntnisgewinn als höheres Ziel der Informationsvisualisierung zu messen. Aufgaben zur Messung einfacher Erkenntnisse sollen den Nutzer dazu veranlassen zu „vergleichen, assoziieren, unterscheiden, einteilen, gruppieren, korrelieren, kategorisieren“ (Plaisant 2004). Aufgaben auf höherer kognitiver Ebene können den Nutzer veranlassen, Verständnis für „Datentrends, Unsicherheiten, kausale Zusammenhänge, Vorhersage der Zukunft, oder das Erlernen einer Domäne“ (Amar und Stasko 2005 zitiert in (Carpendale 2008, 20)) zu entwickeln. North (2006) fordert dahingehend auch neue Evaluationsmethoden wie (1) das protokollierte, offene Experimentieren mit Visualisierungen mithilfe von initialen Fragen, (2) die Erfassung von Erkenntnis mit Think-aloud-Protokollen und verschiedenen Metriken wie Kategorie, Komplexität und Zeitverbrauch zum Erkenntnisgewinn und (3) die Fokussierung auf Nutzer aus der Domäne, um domänen-spezifische Ableitungen und Hypothesen bilden zu können. Analog fordert Carpendale (2008) den Fokus von empirischen Evaluationen auf „echte Nutzer, echte Aufgaben und große, komplexe Datensätze“.

Allerdings muss bei der Wahl einer passenden Evaluationsmethode zwischen den Konzepten *Generalisierbarkeit*, *Exaktheit* und *Realismus* unterschieden werden (McGrath 1995; zitiert in (Carpendale 2008)). Verschiedene Evaluationsmethoden können nur zu einem bestimmten Anteil einen oder zwei dieser Grundkonzepte abdecken. Alle drei Aspekte werden von keiner Methode abgedeckt. Carpendale (2008) gibt passend dazu eine Übersicht von etablierten Methoden der Sozialwissenschaften für die Informationsvisualisierung und wie sie diese Aspekte abdecken. Die grundlegende Einteilung erfolgt dabei anhand quantitativer Methoden, die feste Messgrößen wie Geschwindigkeit, Genauigkeit, Fehlerrate und Zufriedenheit messen und qualitativer Methoden, die einen holistischen Ansatz verfolgen und in realistischen Szenarien evaluieren, um ein besseres Verständnis für das Gesamtszenario zu erhalten.

Eine alternative Einteilung von Evaluationsmethoden und eine Übersicht, wie Evaluationen bisher in der Informationsvisualisierung eingesetzt werden, bietet Lam u.a. (2011). Sie schlagen eine szenariobasierte Einteilung von Evaluationsarten in der Informationsvisualisierung vor, um darauf basierend effektiv anhand von Zielen, Eigenschaften und Beispielen eine passende Art der Evaluation auswählen zu können. Sie führen ein Literaturreview von 800 Papieren von wichtigen Informationsvisualisierung-Konferenzen und -Journals durch. Davon haben 345 Papiere eine Evaluation durchgeführt, die sich in sieben verschiedene Kategorien einteilen lassen:

- *Evaluation von Arbeitsumgebungen und -weisen* (Welche Arbeitsumgebungen und -weisen existieren in einer bestimmten Domäne und wie können Informationsvisualisierungs-Systeme dazu beitragen diese Prozesse zu verbessern?)
- *Evaluation visueller Datenanalyse-Tools* (Wie werden bestehende visuelle Datenanalyse-Tools eingesetzt, um Daten zu explorieren, Hypothesen zu bilden und Wissen zu generieren?)
- *Evaluation von Kommunikation mit Visualisierungen* (Wie werden Visualisierungen eingesetzt, um Botschaften und Nachrichten zu transportieren oder Lerneffekte zu erzielen?)
- *Evaluation von kollaborativer Datenanalyse* (Wie werden Visualisierungen eingesetzt, um kollaborative Datenanalyse zu unterstützen?)
- *Evaluation der Benutzerleistung* (um z. B. verschiedene Visualisierungs- oder Interaktionstechniken gegeneinander zu vergleichen)
- *Evaluation der Nutzererfahrung* (im Sinne von Usability-Tests, um zu erforschen, was der Nutzer über eine Visualisierung oder bestimmte Interaktionstechnik denkt)
- *Automatische Evaluation* (z. B. von verschiedenen Layout-Algorithmen)

3.3 Knowledge Crystallization

Auf welche Weise Informationsvisualisierung den Kognitionsprozess verstärken kann, wird bei Card u.a. durch das Knowledge Crystallization-Modell beschrieben. Dabei wird Knowledge Crystallization als Prozess definiert, in dem der Nutzer Informationen zu einem bestimmten Zweck sammelt (Russell u.a. 1993, zitiert in Card u.a. 1999a), in einem Schema repräsentiert und als Grundlage für Kommunikation oder weitere Aktionen nutzt (Card u.a. 1999a, 10). Knowledge Crystallization-Prozesse sind charakterisiert durch die Verwendung von großen Datenbeständen heterogener Informationen (Card u.a. 1999a, 11). Das Modell lässt sich in folgende, typische Elemente einteilen:

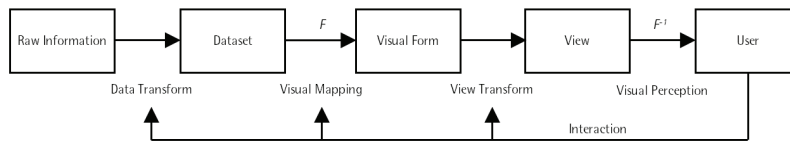
- *Information foraging*: Sammeln von heterogenen Informationen aus verschiedenen Datenquellen.
- *Search for schema (representation)*: Identifikation, welche Eigenschaften der Information wichtig für die Entscheidungsfindung sind.
- *Instantiate schema*: Instanziierung eines Schemas mit den Daten.
- *Problem solve to trade off features*: Nutzung des Schemas mit einfachen Operationen zur Entscheidungsfindung.
- *Search for a new schema that reduces problem to a simple trade-off*: Ist das Schema noch zu komplex, wird ein neues Schema gesucht, welches das Problem auf einfache Abwägungsprozesse reduziert.
- *Package the patterns found in some output product*: Entscheidungsfindung aufgrund der gefundenen Muster oder erneuter Zyklusstart.

Visualisierungen können für die meisten der Schritte im Prozess eingesetzt werden (vgl. Card u.a. 1999a, 12).

3.4 Visualisierungspipeline

Die Visualisierungspipeline (vgl. Abbildung 5) beschreibt den Prozess der Umsetzung von Rohdaten in eine Visualisierung, mit welcher der Nutzer interagieren kann. Zu Beginn stehen unbearbeitete Rohdaten, die mithilfe von Selektion, Filterung und Mapping in einen wohlgeformten Datensatz umgewandelt werden. Die einzelnen Entitäten dieses Datensatzes werden anhand des visuellen Mappings visuellen Glyphen zugeordnet und damit in eine visuelle Form gebracht. Der dritte Schritt stellt die Visualisierung auf dem Bildschirm dar, welche der Nutzer unter Verwendung des menschlichen, visuellen Systems wahrnehmen kann. Zwischen jedem Schritt hat der Anwender die Möglichkeit, durch Interaktion Einfluss auf den Prozess zu nehmen.

Abb 5: Visualisierungspipeline



Nach North (2005); adaptiert von Card, Mackinlay und Schneidermann (1999b).

Das visuelle Mapping ist der Kern der Visualisierungspipeline und besteht aus zwei Schritten. Zuerst wird jeder Datenentität eine visuelle Glyphe zugeordnet. Visuelle Glyphen können Punkte, Linien (Segmente, Kurven, Pfade), Regionen (Polygone, Volumen, Flächen) und Ikonen (Symbole, Bilder) sein. In einem zweiten Schritt werden die Attribute der Datenentität einer Eigenschaft der visuellen Glyphe zugeordnet. Diese unterscheiden sich in Position und der räumlichen Anordnung (x-, y-, z-Koordinaten), Größe (Länge, Fläche, Volumen), Farbe (Farbverlauf, Farbton, Sättigung), Orientierung im Raum (Winkel, Neigung, Einheitsvektor) und Gestalt. Andere visuelle Eigenschaften können Textur, Bewegung, Blinken, Dichte und Transparenz sein.

3.5 Informationsstrukturen

Die Struktur der Information ist ein Hilfsmittel für die initiale Visualisierung der Daten. Auch bildet sie in vielen Fällen die Grundlage für das mentale Modell des Nutzers (North 2005, 1226). Die Informationsstrukturen werden grob eingeteilt in tabellarische, räumliche/zeitliche, baum-/netzwerkförmige, textförmige sowie kombinierte Informationsstrukturen. Die Abbildung dieser Informationsstrukturen als interaktive Visualisierungen wird in diesem Abschnitt anhand der wichtigsten Forschungssysteme verdeutlicht. Die prinzipielle Auswahl der in diesem Abschnitt vorgestellten Systeme gründet dabei auf Übersichten wie (North 2005) oder (Andrews 2002), die eine Darstellung der wichtigsten Forschungssysteme anhand verschiedener Informationsstrukturen nutzten. Ergänzt werden sie durch Forschungssysteme, die ab diesem Zeitpunkt eine wesentliche Rolle für die Informationsvisualisierung spielen (vgl. auch Abschnitte 3.9 und 3.10). Prinzipiell lassen sich Forschungssysteme dabei in zwei Generationen einteilen: (1) die erste Generation, die durch die beginnenden grafischen Möglichkeiten von Computern in den 1990er Jahren entstanden ist. Hier wurden experimentell die neuen Visualisierungs- und Interaktionsmöglichkeiten in verschiedenen Systemen ausprobiert. Beispielsweise nutzten viele Systeme die Möglichkeiten der dreidimensionalen Darstellung für die Anzeige verschiedener Informationsstrukturen. Dabei entstanden zudem auch viele der prinzipiellen Visualisierung- und Interaktionstechniken, die heute noch genutzt werden (TableLens, Parallel Coordinates, Treemap, TopicMaps, ThemeRiver); (2) die zweite Generation, die auf der Entwicklung des Internets, seinen Austausch-

und Vernetzungsmöglichkeiten, bereits etablierten Visualisierungs- und Interaktionstechniken und enormen Datenmengen basiert, die in verschiedensten Domänen wie beispielsweise Lebenswissenschaften, Wetterdaten, Publikationsdaten, Nachrichten etc. entstehen.

3.5.1 Tabellarische Informationsstrukturen

Tabellarische oder mehrdimensionale Informationsstrukturen basieren auf Tabellen (vgl. North 2005, 1226). Tabellen bestehen aus Reihen (Entitäten), denen in Spalten Attribute zugeordnet sind. Über jedes Attribut kann ein Informationsraum aufgespannt werden, in dem jede Entität einen Punkt darstellt. Visualisierungen von Tabellen mit wenigen Attributen können einfach dargestellt werden. Beispiele für die Darstellung einfacher Tabellen sind Diagramme in verschiedenen Ausführungen, zum Beispiel Linien-, Säulen- oder Kreisdiagramme. Bei einer großen Anzahl von Attributen gestaltet sich die Visualisierung komplexer, da für jedes Attribut eine visuelle Eigenschaft von Glyphen genutzt werden muss. Die Problematik der Darstellung mehrerer Attribute greifen Forschungssysteme auf. Zwei wichtige Standardtechniken für die Anzeige multivariater Daten sind *TableLens* und *Parallel Coordinates*.

Das System *TableLens* (Pirolli und Rao 1996) stellt die gesamte Datenmenge auf einem Bildschirm dar und erlaubt so Muster und Ausreißer in den Daten zu erkennen. Dabei werden die verschiedenen Attribute in einer Tabelle abgebildet. Die Struktur der Daten bleibt erhalten und wird auf verschiedene Weise visuell kodiert. Numerische Daten werden als Balken dargestellt, Wochentage farbig angezeigt. Die einzelnen Spalten können auf- oder absteigend sortiert werden, woraus der Nutzer Abhängigkeiten zwischen den Spalten erkennen kann. Mithilfe der Maus können einzelne Spalten oder Bereiche markiert werden. Von der markierten Zeile werden die Attribute nun als Text angezeigt. Dies funktioniert im Sinne der Fokus&Kontext-Technik, welche die Details der markierten Zeile und den Kontext gleichzeitig anzeigen.

Ein anderer Ansatz ist die Methode *Parallel Coordinates* (Inselberg 1997). Hier werden mehrere Achsen parallel nebeneinander gesetzt, die jeweils ein Datenattribut repräsentieren. Jeder Datenpunkt wird nun als Polylinie mit dem Wert auf der jeweiligen Achse aufgezeichnet. Aus den Mustern, die mehrere Datenpunkte bilden, lassen sich Eigenschaften wie Clusterbildung oder Beziehungen zwischen Attributen ablesen. Wichtig ist dabei die Anordnung, Rotation und Skalierung der Achsen.

3.5.2 Räumliche und zeitliche Informationsstrukturen

Räumliche und zeitliche Informationsstrukturen haben einen ein-, zwei- oder dreidimensionalen Bestandteil in dem navigiert werden muss (vgl. North 2005, 1227). Beispiele für eindimensionale Informationsstrukturen sind Zeitleisten,

Musik, Videostreams, Listen, lineare Dokumente und Slideshows. Zweidimensionale Strukturen sind beispielsweise Karten, Satellitenbilder und Fotografien. Dreidimensionale Strukturen können beispielsweise medizinische MRT- oder CT-Bilder, CAD-Architektur-Planungen oder virtuelle Umgebungen sein.

Dreidimensionale Nutzungsoberflächen wurden neben natürlichsprachlichen Anfragen und Softwareagenten immer wieder als nächstes oder alternatives Paradigma zu zweidimensionalen, grafischen Benutzungsoberflächen gesehen. Der Vorteil wird in der Natürlichkeit der Räumlichkeit gesehen. Die Desktopmetapher könne sozusagen ausgeweitet werden auf eine RealWorld-Metapher. Das Ziel ist dabei die natürliche Orientierung im Raum des Menschen zu nutzen, so dass er in Benutzungsoberflächen ähnlich agieren kann, wie in der realen Welt. Hauptargument von Norman (1999) dagegen ist, dass das Bild einer dreidimensionalen Oberfläche mit wirklichem Raum verwechselt wird. Der Nutzer betritt nicht wirklich einen dreidimensionalen Raum, in dem er frei mit seinen Händen und Füßen interagieren kann, sondern er sieht nur ein dreidimensionales Bild auf einem Bildschirm, deren Interaktionsmedien auf Tastatur und Maus beschränkt sind. Die Interaktion wird eingeschränkt auf primitive Instrumente, welche die Vorteile der menschlichen Motorik nicht ansatzweise ausnutzen können. Lösungskonzept dafür sollen dreidimensionale Eingabegeräte wie Datenhandschuh oder 3D-Mäuse sein. Da Navigation aber zum Beispiel auf Handbewegungen abgebildet wird, bleibt die Interaktion durch die zusätzliche Dimension schwierig. Der Nutzer muss die Interaktionstechnik neu lernen und hat vermehrt Schwierigkeiten sich im Raum zu orientieren. Dieser Nachteil gilt auch für die Benutzung von Dreidimensionalität in visualisierenden Systemen. Zusätzlich zur Selektion muss der Nutzer auch in der Dreidimensionalität mit primitiven Mitteln navigieren, ohne wirklich immersiv in sie einzutauchen. Dreidimensionalität kann deshalb nur ansatzweise in wirklich einfachen Diagrammen genutzt werden, ob als besonderer visueller Effekt, oder um die Vergleichbarkeit möglich zu machen. Die Szenerie muss dabei aber immer vollständig durch den Nutzer überblickbar sein und eine Navigation ausgeschlossen werden.

Eine dreidimensionale Darstellung von Treemap wird in *Information Cube* (Rekimoto und Green 1993) verwendet. Hier werden geschachtelte, dreidimensionale Boxen genutzt, um komplette Hierarchien auf einer Bildschirmseite zu präsentieren. Die äußerste Box präsentiert die höchste Hierarchieebene, die darin liegende Box die zweithöchste usw. Auf der Oberfläche jeder Box wird der Titel angezeigt. Die Boxen sind semitransparent, wodurch die Struktur und die Inhalte gleichzeitig angezeigt werden können. Zudem können selektierte Boxen durch Veränderung der Transparenz hervorgehoben werden. Angezeigt werden die Boxen entweder auf einem Head-Mounted-Display oder auf einem Monitor. Interaktion findet über einen Datenhandschuh oder über eine 3D-Maus statt. Dabei können Boxen auf jeder Ebene mit einem Fokusstrahl aus-

gewählt und selektiert werden. Die Boxen können dann mit Drehungen oder Bewegungen der Hand gedreht und verschoben werden.

Eine dreidimensionale Visualisierung von großen Hierarchien auf der Basis von Cheops bietet Information Pyramids (Andrews, Wolte und Pichler 1997). Ein Plateau repräsentiert dabei das Wurzelement, Unterverzeichnisse werden als kleinere Plateaus auf dem Wurzelplateau dargestellt. Die Größe wird durch die Anzahl der Dateien und Unterverzeichnisse bestimmt. Dateien und Dokumente werden durch verschiedene Icons repräsentiert. Die Ordnung der Plateaus kann alphabetisch, chronologisch oder nach anderen Kriterien geschehen. Der Nutzer kann frei in der Pyramiden-Landschaft navigieren und zu beliebigen Details heranzoomen.

Ein Beispiel für die Darstellung hierarchischer Informationsstrukturen als dreidimensionale Baumdarstellung durch verbundene Kegel ist *Cone Tree* (Robertson, Mackinlay und Card 1991). Der Wurzelknoten wird mittig an der Raumdecke platziert, Knoten werden als Index-Karten abgebildet. Die Kanten werden mit einem semitransparenten Kegel umschlossen. Die Größe des gesamten Baums ist immer raumausfüllend, um den maximalen Platz auszunutzen. Dreidimensionalität wird genutzt, um den verfügbaren Platz für die Anzeige der gesamten Struktur möglichst gewinnbringend zu nutzen. Wird ein Knoten selektiert, wird er mit einer einsekündigen Animation in das Zentrum des Blickfelds verschoben und hervorgehoben. Diese interaktiven Animationen sollen die kognitive Last des Nutzers verringern und nach mehreren Selektionen sollen sich dem Nutzer Zusammenhänge innerhalb der Hierarchie erschließen. Zur Unterstützung der Wahrnehmung werden neben der dreidimensionalen Darstellung auch Schatten der Kegel und Knoten auf dem Boden angezeigt. Der Nutzer kann Teilbäume aus- und einblenden und an eine andere Stelle verschieben. Auch eine Suche ist möglich. Suchparameter werden durch Selektion oder Deskriptoreneingabe gesetzt. Die Suchergebnisse werden direkt im Baum angezeigt. Die entsprechenden Knoten werden mit einem roten Balken markiert, wobei die Größe die Relevanz kodiert. *Cam Tree* ist eine Variante, bei der der Wurzelknoten links mittig angeordnet wird.

Perspective Wall (Mackinlay, Robertson und Card 1991) ist ein System für die Anzeige von linearen Daten, die zeitlich oder alphabetisch geordnet sind. Es ist eine dreidimensionale Variante der Detail&Kontext-Technik, bei der 2D-Bildschirmseiten auf eine 3D-Tafel projiziert werden. Die mittlere Tafel zeigt die detaillierte Ansicht, die perspektivischen Tafeln zeigen den Kontext. Durch Interaktion kann ein Element der perspektivischen Seiten gewählt werden und mit einer Animation wird diese zum Zentrum.

3.5.3 Baum- und Netzwerkstrukturen

Baum- und Netzwerkstrukturen enthalten Verbindungen zwischen den einzelnen Entitäten (vgl. North 2005, 1229). In der Graphentheorie besteht ein Netz-

werk aus einer Anzahl von Knoten, die mit Kanten miteinander verbunden sind. Die Kanten können zwei Knoten direkt oder indirekt miteinander verbinden. Wie Knoten können auch Kanten Attribute enthalten. Netzwerkgraphen eignen sich beispielsweise für die Darstellung von Kommunikationsnetzwerken oder Hyperlinks zwischen Webseiten. Baumstrukturen sind hierarchische Strukturen, die Knoten mit einer Eltern-Kind-Beziehung verbinden. Dabei darf ein Kindknoten höchstens einen Elternknoten besitzen. Baumdarstellungen finden sich bei Darstellung von Dateiordnern von Computern, Menüs oder Organisationsprogrammen. Die Herausforderung bei der Visualisierung ist (1) das Layout des Netzwerkes oder Baumes, um die Struktur der Verbindungen aufzudecken und (2) die Visualisierung der Attribute von Entitäten und Verbindungen. Für die Visualisierung von Eltern-Kind-Beziehungen gibt es zwei Verfahren:

- Das Link-Verfahren nutzt Knoten-Link-Diagramme. Dabei werden Entitäten auf visuelle Knoten und Verbindungen als Linien zwischen den Knoten abgebildet.
- Das Containment-Verfahren, bei dem die zur Verfügung stehende Fläche maximal genutzt wird.

Treemap (Shneiderman 2009) bietet eine Baumdarstellung der gesamten Hierarchie mittels der Verschachtelung von Rechtecken. Ein rekursiver Algorithmus unterteilt Rechtecke abwechselnd horizontal und vertikal während der Baum nach unten traversiert wird. Werden auf einer Ebene zu viele Rechtecke auf einmal dargestellt, ist der Nutzer überlastet. Große Mengen an Rechtecken auf unterer Ebene können zu einem homogenen Feld zusammengefasst werden und erst beim Hineinzoomen wieder aufgelöst werden. Es existieren zahlreiche Varianten von Treemap, zum Beispiel mit alternativen Grundformen wie Polygonen und Kreisen. Benutzertests ergaben, dass neue Nutzer 10 bis 15 Minuten benötigen, um mit Treemap vertraut zu werden.

Auch die Anzeige der gesamten Hierarchie in der Darstellung einer Pyramide bietet Cheops (Beaudoin, Parent und Vroomen 1996). Visuelle Komponente für einen Knoten ist ein Dreieck. Diese werden ohne Kanten übereinander gestapelt, so dass sie die Form einer Pyramide erhalten. Unterknoten werden überlappend dargestellt, um Platz zu sparen. Durch die komprimierte Darstellung und fehlenden Kanten haben die inneren Knoten mehrere mögliche Elternknoten. Die Selektion des richtigen Kind-Knotens wird dabei durch den aktiven Elternknoten bestimmt. Somit ist nur der Wurzelknoten eindeutig bestimmt. Bei der Preselection-Technik wird der aktuelle Zweig beim Darüberfahren mit der Maus hervorgehoben. Bei der Selektion wird der aktuelle Zweig farbig markiert und der übrige Baum ausgegraut.

Eine Fokus&Kontext-Technik innerhalb einer hyperbolischen Baumansicht für die Anzeige großer hierarchischer Daten bieten Lamping u.a. (1995). Metapher ist die Abbildung in einer Spiegelkugel. Die Mitte stellt den Fokus dar und erscheint unverzerrt, zum Rand hin nimmt der Detailreichtum ab und die Ver-

zerrung zu. Durch Anklicken kann ein Knoten in den Mittelpunkt verschoben werden und der Kontext wird angepasst.

3.5.4 Textförmige Informationsstrukturen

Textförmige Informationsstrukturen bestehen aus Sammlungen von Dokumenten (vgl. North 2005, 1232), zum Beispiel Digitale Bibliotheken, Nachrichtenarchive oder Quellcodes von Software. Gegenüber den bereits vorgestellten Informationsstrukturen sind textförmige nur wenig strukturiert und lassen sich schwieriger darstellen. Text kann nicht einfach einem visuellen Mapping unterzogen werden. Eine Möglichkeit, Dokumentenräume zu visualisieren, ist durch Methoden von Semantic Maps. Dabei wird die semantische Nähe der Dokumente zueinander auf einer Karte eingezeichnet. Die thematische Ähnlichkeit kann zum Beispiel durch die Häufigkeit bestimmter Wörter im Text bestimmt werden.

WEBSOM (Honkela u.a. 1998) organisiert beliebige Sammlungen von Textdokumenten auf einer thematischen Karte. Inhaltlich ähnliche Dokumente befinden sich auf der Karte nahe beieinander. Einzelne Dokumente können durch Browsing oder durch eine Suche gefunden werden. In einem ersten Schritt werden die Dokumente vorbehandelt. Steuerungszeichen, häufige Wörter ohne Themenrelevanz und Wörter unterhalb einer gewissen Häufigkeit werden aus dem Textraum entfernt. Nun werden Wortkategorien zusammengefasst, die auf der nächsten Umgebung des Wortes basieren und auf einer Karte dargestellt werden können. Zur Erstellung der Dokumentenkarte werden die Dokumente mithilfe der Wortkategorien als Histogramm enkodiert. Diese werden mit einem Gaußschen Filter geglättet, um sie unempfindlicher gegen leichte Variationen im Text zu machen. Mithilfe der Histogramme werden die Dokumente dann auf einer Karte angeordnet, wobei ähnliche Dokumente nahe beieinander stehen. Häufigkeiten von Dokumenten werden durch hellere Farbtöne dargestellt. Der Nutzer kann stufenweise durch Anklicken einer Region in die Karte hineinzoomen, wobei zuerst die Übersicht, dann die vergrößerte Ansicht, eine Liste der dort angesiedelten Dokumente und dann der Volltext selbst erscheint. Mithilfe eines Interaktionselements kann innerhalb der Karte und zur höhergelegenen Ansicht navigiert werden. Alternativ kann der Nutzer mit Deskriptoren suchen, die gefundenen Regionen werden in der Karte eingekreist, wobei größere Kreise eine höhere Trefferanzahl visualisieren.

Ein System für die Darstellung von Ergebnissen von Suchanfragen auf Karten ist *VisIslands* (Andrews u.a. 2001). Dabei werden Dokumentencluster als Inseln auf einer zweidimensionalen Karte eingezeichnet. *VisIslands* ist eine Visualisierungskomponente von *xFIND*, einem System für das Sammeln und Indizieren von Datenquellen und der Ergebnisdarstellung von Suchanfragen. Suchresultate werden zuerst in Clustern vorsortiert. Die Mittelpunkte der Cluster werden zufällig auf einer zweidimensionalen Fläche angeordnet. Die dazugehörigen Dokumente werden kreisförmig darum eingezeichnet. Durch ein

iteratives Verfahren werden ähnliche Dokumente zueinander gezogen, bis sich das System nach einer bestimmten Zeit stabilisiert. Die einzelnen Dokumente tragen ihr Relevanzgewicht zu den Höhenlinien, auf denen sie liegen, bei. Die gesamte Darstellung erinnert an eine Reliefkarte von mehreren Inseln.

Eine dreidimensionale Darstellung von Dokumentenclustern als Höhenkarte findet sich in *IN-SPIRE Themeview* (Wong u.a. 2004). Das System nutzt Text-, HTML- oder XML-Dokumente als Datenbasis. Durch statistische Verfahren werden Schlüsselwörter und Themen extrahiert. *IN-SPIRE Themeview* baut auf *IN-SPIRE Galaxy* auf. Es extrahiert die relevanten Schlüsselwörter einer Galaxie und stellt diese als Themenberge in einer dreidimensionalen Landschaft dar. Die Höhe der Berge ist analog zu der Anzahl der Beiträge. Um den Effekt zu verstärken, wird die Höhe zusätzlich farbig kodiert. Als alternative Visualisierung zu *IN-SPIRE Themeview* nutzt *IN-SPIRE Galaxy* auch die Sternenhimmel-Metapher. Die Dokumente werden als Sterne dargestellt, semantisch nahe Dokumente bilden Galaxien. Durch Interaktion kann im Sternenhimmel gezoomt und navigiert werden. Mit diesen Visualisierungen kann einerseits die Struktur eines vollständigen Datenraums abgeschätzt werden als auch die Ergebnisse einer spezifischen Suche angezeigt werden, wobei die Relevanz mit einer Farbkodierung der Elemente arbeitet.

Beispiele für die Anzeige von Dokumenten mit der Metapher eines Sternenhimmels finden sich in *Infosky* (Kappe u.a. 2002) und *IN-SPIRE Galaxy* (Wong u.a. 2004). In ihnen wird die Visualisierung von großen hierarchisch angeordneten Ablagesystemen möglich. Dokumente werden als Sterne dargestellt, ähnliche Dokumente werden in geometrischer Nähe zueinander angeordnet. Durch diese thematische Anordnung entstehen Sternhaufen und Galaxien. Thematisch nahe Galaxien werden auch wieder nahe beieinander positioniert. In *Infosky* stehen Sterne für Dokumente, Galaxien und Cluster bilden hierarchische Strukturelemente wie Ordner ab. Diese werden als polygone Linienzüge dargestellt, die Sternhaufen umschreiben. Das Teleskop ist die Metapher für die Interaktion. Der Nutzer kann bestimmte Objekte fokussieren und zoomen (zooming) und über den Sternenhimmel schwenken (panning). *Infosky* besteht aus einer dreigeteilten Oberfläche: einer Toolbar für Optionen und Zugang zur Suche; einer Baumansicht auf der linken Seite für das Öffnen und Schließen von Ordnern und einer Teleskopdarstellung der aktuell selektierten Galaxie. Baumansicht und Teleskop-Darstellung sind dabei immer synchronisiert.

ThemeRiver (Havre, Hetzler und Nowell 2000) zeigt aggregierte Thematiken, beispielweise eines Dokumentenkörpus in einer Flussmetapher auf einer Zeitachse an. So kann die Entwicklung einer Thematik über die Zeit beurteilt werden und gleichzeitig mit Zeitereignissen, welche in das Diagramm eingezeichnet sind, in Verbindung gebracht werden.

Eine Anwendung der Treemap-Visualisierung wird auch für die Visualisierung von Nachrichten in der Online-Anwendung *Newsmap*⁵ genutzt. Hier wird die Anzahl ähnlicher Nachrichten und damit ihre momentane Wichtigkeit in Größe der Blöcke kodiert, verschiedenen Kategorien sind Farben zugeordnet und die Anzeige der Nachrichten ist dynamisch und ändert sich nach der Nachrichtenlage.

Die *Flip Zoom*-Technik (Holmquist 1997) ist eine Variante der Fokus& Kontext-Technik, um eine Überblicks- und Detailansicht gleichzeitig anzuzeigen. Das aktuelle Dokument wird vergrößert dargestellt, die Seiten davor und danach werden als Thumbnails angezeigt. Dabei erfolgt die Ordnung der Thumbnails um die aktuelle Seite analog ihrer Reihenfolge von links nach rechts und von oben nach unten. So kann vom Nutzer erkannt werden, in welchem Kontext sich die aktuelle Seite befindet. Um den Nachteil nicht lesbarer Schrift in den Thumbnails auszugleichen, werden die Überschriften aus den Seiten extrahiert und vergrößert in den Thumbnails dargestellt. Die Technik wurde im System Zoom Browser integriert.

Neben der zweidimensionalen Darstellung von textförmigen Informationsstrukturen existieren auch dreidimensionale Systeme. VR-VIBE (Snowdon und Jää-Aro 1997) ist eine dreidimensionale, virtuelle Umgebung, in der mehrere Nutzer gemeinsam an Daten arbeiten und miteinander kommunizieren können. Dokumente oder Referenzen werden als Symbole in einer 3D-Umgebung angezeigt. Der dreidimensionale Raum wird aufgespannt durch Point of Interests (POI), die Stichwörter der aktuellen Suchanfrage repräsentieren. Die Position der Dokumente zeigt den Bezug zu den einzelnen Stichwörtern an. Je näher Dokumente zueinander positioniert sind, umso näher sind sie sich thematisch. Die Relevanz eines Dokumentes wird durch die Größe und Helligkeit des Icons angezeigt. So kann zwischen Dokumenten dessen Bezug zu Stichwörtern der aktuellen Suchanfrage gleich, deren Relevanz aber unterschiedlich ist, unterschieden werden. Die Relevanz wird durch Algorithmen errechnet, welche die Häufigkeit der gesuchten Deskriptoren in Titel, Abstract und im Text ermitteln. POIs werden als grüne Oktaeder repräsentiert, eine weiße Kugel zeigt den aktuellen POI an. Dokumente werden als blaue Blöcke angezeigt, mehrere Nutzer mit verschiedenen Ikonen gekennzeichnet. Nutzer können im Raum navigieren, einzelne Dokumente auswählen und POIs verschieben. Durch das Verschieben der POIs wird der Dokumentenraum verändert. Neue Suchen können über das Anlegen eines neuen POI oder das Spezifizieren von Suchtermen geschehen. Verschiedene Nutzer können unterschiedliche Visualisierungen der Daten nutzen.

LyberWorld *Navigation Cone* (Hemmje 1995; Hemmje, Kunkel und Willett 1994) ist Teil des LyberWorld-Projekts für die Schlagwortfindung. Als visuelles

⁵ <<http://newsmap.jp/>>.

Mittel wird ein ConeTree eingesetzt. Ausgehend von einem Schlagwort wird eine ConeTree-Ebene erzeugt, die alle Texte anzeigt, die das Schlagwort enthalten. Öffnet man einen Textknoten, werden in einer weiteren Ebene alle Terme des Textes angezeigt. In einer dritten Ebene werden alle Texte angezeigt, die den jeweiligen Term enthalten. So können ausgehend von einem Schlagwort weitere Schlagwörter für die Suche gefunden werden, die dann in der RelevanceSphere verschieden gewichtet werden können.

RelevanceSphere (Hemmje u.a. 1994, 253ff) ist eine Entwicklung innerhalb des LyberWorld-Projekts für die Gewichtung von Suchbegriffen. Dokumente werden innerhalb einer Kugel abgebildet. Je höher die Relevanz eines Dokumentes, umso näher ist es an einem Suchbegriff positioniert. Dokumente, die alle Suchbegriffe enthalten, verbleiben im Zentrum, wenn die Deskriptoren gleichmäßig auf der Oberfläche verteilt sind. Diese werden als Trabanten in Form von Kugeln außerhalb der Sphäre dargestellt. Die Suchbegriff-Icons können verschoben werden, um so Suchbegriffe eines Kontextes räumlich nahe beieinander zu positionieren und damit relevante Dokumente zu finden. Durch das Anklicken eines Icons kann der Deskriptor höhergewichtet werden und damit werden relevante Dokumente mehr in diese Richtung angezogen. Durch Verkleinerung der Kugel treten Dokumente aus der Sphäre heraus, die nun im Volltext gelesen werden können. Krause (1996, 32ff) kritisiert das System, da die genutzte Metapher Mehrdeutigkeiten aufweist und die Dreidimensionalität eher negativ wirkt.

SpaVis (Keskin und Vogelmann 1997) und *City of news* (Sparacino, Davenport und Pentland 1996) ordnen Informationen als Stadtlandschaften an. *SpaVis* nutzt die Stadtbild-Metapher, um hierarchische Bäume dreidimensional darzustellen. Für strukturierte, relationale Informationen, wie hierarchische und Netzwerkinformation wird eine Generalisierung von Säulendiagrammen in 3D genutzt. Knoten werden als auf Ebenen geordnete Quader oder Zylinder angezeigt. Höhe und Position kodieren die Ebene in der Hierarchie, die Größe kodiert die Anzahl der Kindknoten. Dabei werden die Kindknoten kreisförmig um den Elternknoten angeordnet.

City of news nutzt die Stadtbild-Metapher, verfolgt aber einen philosophischen Ansatz. Ausgehend von einer Homepage werden Texte und Bilder von Webseiten auf Hochhäuser und Straßen projiziert. Die Stadt wird dabei weiter in Bezirke, zum Beispiel für Finanzen, Unterhaltung und Einkäufe unterteilt. Genutzt werden soll in diesem System die Fähigkeit der natürlichen Orientierung vom Menschen im Raum. Lynch (1960) identifiziert fünf Elemente, welche für die Bildung einer kognitiven Karte einer städtischen Umgebung wichtig sind: Orientierungspunkte (landmarks), Bezirke (districts), Wege (paths), Knotenpunkte (nodes) und Begrenzungslinien (edges). Mit dem Einsatz dieser Charakteristika soll die Bildung einer kognitiven Karte gefördert und die Orientierung im System erleichtert werden.

3.5.5 Kombinierte Informationsstrukturen

Häufig werden in realen Anwendungen verschiedene Informationsstrukturen verwendet, die in Beziehung zueinander stehen (vgl. North 2005, 1232). Eine Visualisierung verschiedener Informationsstrukturen in einer kombinierten Ansicht ist ausgesprochen schwierig. Da bereits eine Informationsstruktur den gesamten Platz beim visuellen Mapping ausfüllen kann, können kombinierte Informationsstrukturen beim Mapping in einem Konflikt enden. Eine Möglichkeit, dies zu umgehen, sind Mehrsichtensysteme, in denen die Informationsstrukturen in separierten Ansichten angezeigt werden. Dabei kann jede Informationsstruktur unabhängig voneinander das optimale visuelle Mapping nutzen. Die verschiedenen Sichten werden mit einem interaktiven Linking verbunden. Da interaktive Linkings nur wenige Assoziationen gleichzeitig anzeigen können, muss der Nutzer die Beziehungen zwischen den Strukturen im Laufe der Zeit mental zuordnen. So können vom Benutzer interessante Beziehungen übersehen werden. Werden mehrere Strukturen in einer Ansicht integriert, muss eine Informationsstruktur als räumliche Basis dienen. Die anderen Strukturen werden dann in diesen Raum eingefügt. Dies stellt die Beziehung der zweiten zur ersten Struktur heraus, aber die Klarheit über den Aufbau der zweiten Struktur kann verloren gehen.

Tab. 6: Mögliche Visualisierungstypen für Informationsstrukturen

Informationsstruktur	Mögliche Visualisierungstypen
Tabellarisch	Diagramme, TableLens, Parallel Coordinates, Aufteilung der Attribute in separate Ansichten
Textförmig	Makrolevel (Dokumentenübersichten): Topic-Maps, Baumstrukturen für TOC, tabellarische Strukturen für Metadaten, Netzwerkstrukturen für Zitations- oder Autorennetzwerke Mikrolevel (einzelne Dokumente): Tagclouds, Word Tree, Wordle
Räumlich & Zeitlich	1D: Zeitleisten, Musik, Video, Streams, Listen, lineare Dokumente, Slideshows 2D: Karten, Satellitenbilder, Fotografien, Blaupausen 3D: MRT- & CT-Bilder, CAD-Architekturpläne, Virtuelle Umgebungen
Baum- und Netzwerk	Netzwerkgraph, Baumdarstellung in verschiedenen Layouts: gegliedert-verschachtelt (Explorer-Ansicht), SpaceTree, HyperbolicTree, ConeTrees
Kombiniert	Kombinierte Ansicht oder Mehrsichtensysteme

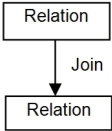
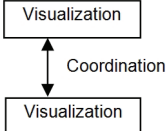
Wang-Baldonado et al. (2000) entwickeln ein Modell für koordinierte Mehrsichten (*Multiple Coordinated Views*) und geben Regeln vor, damit der positive Mehrwert nicht durch die erhöhte Komplexität gestört wird. Die Kernidee ist, dass Daten in verschiedenen Ansichten verbunden werden können. Werden Daten in einer Ansicht selektiert, werden sie auch in den anderen Ansichten hervorgehoben (Brushing-and-Linking). North und Shneiderman präsentieren ein alternatives Visualisierungs-Modell, welches auf dem relationalen Daten-

modell basiert (North und Shneiderman 2000). Das System *Snap* (North u.a. 2002) ist eine Implementierung dieses Modells. Es erlaubt dem Nutzer, Datenbasen auszuwählen und Visualisierungen zuzuweisen. In einem zweiten Schritt können verschiedene Visualisierungen verbunden und koordinierte Visualisierungen generiert werden. Hervorhebungen oder andere Aktionen werden zwischen den verschiedenen Sichten koordiniert. Werden Daten zum Beispiel in einer Visualisierung ausgewählt, werden sie auch in den anderen Ansichten ausgewählt.

Tabelle 6 fasst die grundlegenden Informationsstrukturen für die Bildung eines mentalen Modells zusammen und zeigt mögliche Visualisierungstypen.

3.5.5.1 Snap-Visualization-Modell

Tab. 7: Analogie zwischen dem relationalen Datenbank-Modell und dem Snap-Visualization-Modell

	<i>Relational Databases</i>	<i>Snap Visualization</i>
<i>Usage:</i>		
Design goal	Data design	Visualization design
Design method	Data schema	Visualization schema
Designer	Data owner	Data owner
Design change	Rapid, dynamic	Rapid, dynamic
Adaptability	Flexible	Flexible
<i>Perspectives:</i>		
Theory	Relational data model	Snap visualization model
User interface	Relational data schema	Snap visualization schema
Architecture	Relational DBMS	Snap visualization server
<i>Schema primitives:</i>		
Theory	Relation	Visualization component
	Tuple	Visual item
	Attribute	Visual property
	Selection	User interaction
	Join	Coordination
User interface		

Nach North u.a. (2002, Tabelle 2).

North u.a. (2002) stellen ein Modell für die koordinierte Anzeige von mehreren Sichten vor. Das Modell wurde analog zum relationalen Datenmodell entworfen, wobei sich folgende Beziehungen zwischen den Konzepten ergeben (vgl. North u.a. 2002, 215f; vgl. auch Tabelle 7):

- *Visualisierungskomponente* = *Datenrelation*: Eine Visualisierungskomponente ist eine Ansicht, die eine Datenrelation oder ein Anfrageergebnis abbildet. Ein Beispiel ist ein Streudiagramm, das eine binäre Relation abbildet.

Eine Visualisierungskomponente kann einen bestimmten Visualisierungstyp nutzen oder automatisierte Techniken anwenden, um dynamisch Visualisierungen zu erstellen.

- *Visuelles Element (item) = Datentupel*: Datentupel werden als visuelle Elemente in der Visualisierungskomponente angezeigt. So wird beispielsweise ein Tupel als Punkt in einem Streudiagramm angezeigt.
- *Visuelle Eigenschaft (property) = Datenattribut*: Datenattribute werden genutzt, um die Grafiken in der Visualisierungskomponente darzustellen. Der Nutzer kann Datenattributen komponentenspezifische visuelle Eigenschaften zuweisen. Zum Beispiel kann ein Datenattribut der X-Achse zugewiesen werden. So werden die Tupel anhand ihres Wertes für dieses Attribut visuell auf der X-Achse ausgerichtet.
- *Nutzerinteraktion = Selektion einer Untermenge von Tupeln*: Eine Interaktion des Nutzers in einer Visualisierungskomponente wählt eine Teilmenge von Tupeln analog zu einer Selektionsanfrage und verändert normalerweise die Anzeigeart der ausgewählten Tupel. So kann der Nutzer beispielsweise eine Menge von Tupeln in einem Streudiagramm direkt selektieren oder er zoomt auf ein einzelnes Tupel, um Details über diesen Datensatz zu erfahren. Interaktionen werden durch jede Visualisierungskomponente definiert und nur durch ihren Namen identifiziert (z. B. „select“, „zoom“). Jede durch eine Visualisierungskomponente definierte Interaktion hat eine Teilmenge an Tupeln, die sie kontrollieren kann. Diese Teilmenge besteht aus null oder mehr Tupeln, wobei null Tupel anzeigt, dass kein Tupel ausgewählt ist. Tupel einer Teilmenge können anhand ihres Primärschlüssels identifiziert werden. Jede Visualisierungskomponente hat inhärent die komplette Relation geladen, die in der Visualisierungskomponente angezeigt wird.
- *Koordinationen zwischen Visualisierungen = Verbund der Relationen*: Eine Koordination verknüpft eine Interaktion in einer Visualisierungskomponente zu einer Interaktion in einer zweiten Visualisierungskomponente, wobei die beiden Teilmengen basierend auf einer Join-Operation der beiden zugrundeliegenden Relationen abgeglichen werden. Interaktion des Nutzers in einer Visualisierungskomponente verursacht visuelle Aktionen in der per Join-Operation verbundenen zweiten Visualisierungskomponente. Die durch Interaktion ausgewählte Teilmenge in der ersten Visualisierungskomponente wird mit einer Inner-Join-Operation mit der zweiten Visualisierungskomponente verbunden, woraus eine neue Teilmenge entsteht, die als Basis für die Aktion in der zweiten Visualisierungskomponente dient. Zum Beispiel ist ein Brushing-and-Linking zwischen zwei Streudiagrammen, die auf derselben Relation basieren, eine Koordination der Select-Aktion in einer 1-zu-1 Join-Verknüpfung. Werden Elemente in einem Diagramm hervorgehoben, werden die verknüpften Elemente in dem zweiten Diagramm automatisch hervorgehoben. Eine Koordination kann jedes Paar von Aktionen zwischen

Visualisierungskomponenten verknüpfen. Koordinationen wie Joins sind dabei bidirektional.

Es wird im Snap-Visualization-Modell zwischen vier Arten des Joins unterschieden (vgl. North u.a. 2002, 216f):

- *Self Join*: Der Self Join stellt eine 1-zu-1 Beziehung zwischen zwei Visualisierungskomponenten her, die auf derselben Relation beruhen.
- *Single Join*: Eine Koordination kann zwischen zwei Visualisierungskomponenten hergestellt werden, deren Relationen eine direkte Assoziation haben.
- *Compound Join*: Eine Koordination kann zwischen zwei Visualisierungskomponenten hergestellt werden, die eine indirekte Assoziation über eine oder mehrere dazwischen liegende Relationen haben.
- *Multiple Alternative Join*: Der Multiple Alternative Join ist ähnlich dem Compound Join, es bestehen jedoch mehrere Möglichkeiten der Join-Assoziationen zwischen den Relationen zweier Visualisierungskomponenten.

3.6 Übersichtstechniken

Der Visualisierungsprozess von sehr großen Datenmengen ist problematisch, da nicht alle Informationen gleichzeitig auf einem Bildschirm visuell dargestellt werden können. Um den Visualisierungsprozess von sehr großen Datenräumen zu unterstützen, schlägt Shneiderman ein Design Mantra vor: „*overview first, zoom and filter, then details on demand*“ (Shneiderman 1996). Der Nutzer startet mit einer Übersicht und kann mithilfe von Interaktionstechniken zur gewünschten Information zoomen und Daten filtern. Zuletzt kann er sich durch gezielte Selektion Details zu den Daten ansehen. Die Übersicht zu Beginn hat mehrere Vorteile (North 2005, 1234):

- Mithilfe der Übersicht kann sich der Nutzer ein mentales Model des Informationsraumes machen.
- Der Nutzer kann erkennen, welche Informationen zur Verfügung stehen und welche nicht.
- Durch die Übersicht besteht ein direkter Zugang und durch Selektion kann direkt zu den gewünschten Daten navigiert werden.
- Eine Übersicht unterstützt die Erforschung der Daten.

Es existieren zwei Methoden, um Übersichten darzustellen, die auf einer großen Datenmenge basieren: (1) Aggregation und (2) Verringerung der Glyphen. Die erste Möglichkeit zur Verringerung der Datenmenge ist die Gruppierung von Datensätzen zu einem neuen Datensatz. Der erste Schritt dabei ist die Wahl der Datensätze, die zusammengefasst werden sollen. Datensätze können anhand gemeinsamer Attribute oder durch fortgeschrittene Techniken ausgewählt werden. Der entstehende Datensatz sollte die gruppierten Datensätze hinreichend repräsentieren. Statistische Methoden wie Mittel, Minimum, Maximum oder Anzahl werden oft eingesetzt (North 2005, 1235). Alternativ können die

Glyphen, die durch den visuellen Mapping-Prozess entstehen, minimiert werden. Dabei wird pro Glyphe so wenig wie möglich Fläche verbraucht, so dass eine maximale Anzahl von Daten auf dem Bildschirm dargestellt werden kann.

3.7 Navigationstechniken

Navigationstechniken werden für den Übergang von einer Überblicks- zu einer Detailansicht der Daten in der Visualisierung genutzt. Diese Interaktionsmöglichkeit befindet sich auf dritter Ebene der Visualisierungspipeline. Es werden hauptsächlich drei Navigationstechniken genutzt: Zoom & Panning, Overview & Detail, sowie Focus & Context (North 2005, 1237).

3.7.1 Zoom & Panning

Zoombare Visualisierungen beginnen mit einem Überblick und lassen den Anwender dynamisch in eine Region hineinzoomen, um Details über die Daten zu erfahren. Der Nutzer kann wieder zum Überblick hinaus- und in eine andere Region hineinzoomen. Durch Schwenken oder Scrollen kann auf jeder Zoom-Ebene im Raum navigiert werden (North 2005, 1237). Vorteile dieser Technik sind das effektive Ausnutzen des Bildschirmplatzes und die unendliche Skalierbarkeit der Ansicht. Nachteile sind die Risiken für den Nutzer, beim Hineinzoomen in die Visualisierung den Überblick zu verlieren und die langsame Navigation, da ständige Interaktion notwendig ist (vgl. North 2005, 1238).

3.7.2 Overview & Detail

Overview & Detail nutzt mehrere Sichten, um gleichzeitig eine Überblicks- und eine Detailansicht anzuzeigen. Ein Sichtfeld-Indikator in der Übersichtsansicht zeigt die aktuelle Position in der Detailansicht an. Die beiden Ansichten sind verbunden, so dass ein Verschieben des Sichtfeld-Indikators auch die Veränderung der Detailansicht zur Folge hat. Umgekehrt beeinflusst das Verändern der Detailansicht auch den Sichtfeld-Indikator (North 2005, 1237). Vorteile sind (vgl. North 2005, 1238) eine stabile Übersichtsansicht über die gesamten Daten und gleichzeitig eine skalierbare Detailansicht, die miteinander verkettet sind und unterschiedliche Fokusse auf die Daten bieten können. Nachteile sind ein optischer Bruch und der damit verbundene nötige Blickwechsel des Anwenders zwischen den verschiedenen Ansichten. Zusätzlich besteht eine Konkurrenz zwischen den Größen der beiden Ansichten, insofern als dass die Überblicksansicht in den meisten Fällen sehr klein ausfällt.

3.7.3 Focus & Context

Mit der Fokus und Kontext-Technik werden Übersichts- und Detailansicht verbunden. Der Fokus wird erweitert und vergrößert dargestellt, um detaillierte

Information für einen Bereich darzustellen. Der Anwender kann einfach in der Übersichtsansicht navigieren, um den Fokus zu verschieben. Um den nötigen Raum für die Fokusregion zu beschaffen, wird der Raum um den Fokus verzerrt dargestellt. Der Vergrößerungsgrad der Detailansicht ist im Brennpunkt am größten und nimmt zum Rand hin ab. Diese Technik funktioniert wie ein Vergrößerungsglas oder ein Fischauge, das die Details im Brennpunkt vergrößert darstellt, aber trotzdem die Umgebung anzeigt. Vorteil ist, dass die Details visuell verbunden zu ihrer Umgebung angezeigt werden. Die Skalierbarkeit ist aber auf einen Faktor von 1:10 limitiert und die Verzerrung führt zu einer instabilen Ansicht der Übersicht. Diese dynamische Verzerrung kann beim Nutzer zu Desorientierung führen. (vgl. North 2005, 1238).

3.8 Interaktionstechniken

Interaktionstechniken bieten die Möglichkeit, die Komplexität der Daten zu reduzieren. Die vorrangigsten Techniken sind: Selecting, Linking, Filtering, Rearranging and Remapping.

3.8.1 Selecting

Die grundlegendste Interaktionstechnik ist die interaktive Auswahl von Daten durch den Nutzer. Dies kann aus mehreren Gründen geschehen: Entweder, um sich detaillierte Informationen zu den Daten anzeigen zu lassen, um Daten hervorzuheben, um verwandte Daten zu gruppieren oder um sie für eine weitere Verarbeitung zu markieren. Die Auswahl der Daten kann dabei direkt oder indirekt erfolgen. Die direkte Auswahl kann über direktes Anklicken der Glyphen oder die Auswahl einer Gruppe von Glyphen mithilfe des Mauszeigers erfolgen. Eine indirekte Auswahl erfolgt zum Beispiel über Attributkriterien, die den Daten zugeordnet sind oder die Auswahl ganzer Pfade in Netzwerk- oder Baumdarstellungen. Selecting-Techniken sollten so gestaltet sein, dass der Anwender sehr einfach mehrere Glyphen markieren und weitere hinzufügen kann. Auch sollten Glyphen aus der aktuellen Auswahl entfernt oder die gesamte Auswahl gelöscht werden können (vgl. North 2005, 1240).

3.8.2 Linking

Mit Linking werden Daten interaktiv in mehreren Ansichten verbunden. Daten können in mehreren Ansichten unterschiedlich in Bezug auf verschiedene Perspektiven, unterschiedlichen Fokus oder Art der Darstellungsweise angezeigt werden. Die bekannteste Form des Linkings ist Brushing-and-Linking (Becker und Cleveland 1987), bei der das Markieren von Daten in einer Ansicht auch die Daten in den anderen Ansichten markiert und hervorhebt. Beim Linking kann der Nutzer die individuellen Stärken verschiedener Repräsentati-

onen nutzen. So können in einer Ansicht Daten markiert werden, um diese in einer anderen abzufragen (North 2005, 1240f).

3.8.3 Filtering

Filtering lässt den Nutzer dynamisch die Menge der Daten in der Visualisierung verkleinern und den Fokus auf die Daten lenken, die den Betrachter interessieren. Slider in Form grafischer Schieberegler können Filter-Parameter bestimmen, die der Nutzer in Echtzeit verändern kann. So kann das gefilterte Ergebnis direkt vom Nutzer gesehen werden. Zusätzlich kann die Verbindung von Filterparameter zu Datenattribut durch das direkte visuelle Feedback durch den Nutzer erkannt werden. Dynamic Queries können den Nutzer bei der Suche unterstützen, indem er die vorhandenen Daten erkunden und unerwünschte herausfiltern kann. Durch die Manipulation grafischer Bildelemente können die Parameter der Datenbankabfrage in Echtzeit verändert werden und die Ergebnismenge wird sofort angezeigt (North 2005, 1241).

Ein erstes System für die Anwendung von Dynamic Queries war der *Film Finder* (Ahlberg und Shneiderman 1994; Jog und Shneiderman 1995). In einem zweidimensionalen Koordinatensystem wird eine Filmdatenbank in der Sternenhimmel-Metapher visualisiert. Dabei ist die x-Achse mit dem Erscheinungsjahr und die y-Achse mit der Popularität von eins bis neun belegt. Jeder Film wird farblich kodiert nach Genre in das Diagramm als Rechteck eingezeichnet und je nach Zoom-Stufe mit Titel angezeigt. Klickt man einen Film an, werden Detailangaben angezeigt. Über die Achsen können die Popularität und das Erscheinungsjahr eingeschränkt werden. Auf der rechten Seite können die Filme nach Titel, Schauspieler und Regisseur alphabetisch gefiltert werden. Die Technik wird auch weiterhin in der kommerziellen Version *Spotfire* genutzt und wird auch immer mehr als Standard-Interaktionselement im Web eingesetzt (durch sie ist eine einfache und schnelle Einschränkung von Preisen oder Zeiten bspw. in Hotel-Portalen möglich).

3.8.4 Rearranging and Remapping

Die Visualisierung von Daten wird meist gemäß ihrer Informationsstruktur dargestellt. Die Darstellung nach Ordnungen kann aber zu Abbildungen führen, aus der sich für den Nutzer keine Erkenntnisse ziehen lassen. Die Wahlmöglichkeit zwischen mehreren Visualisierungsarten kann demnach beim Nutzer zu vermehrten Einsichten führen, wenn sich durch die andersartige Darstellung bestimmte Muster in den Daten erkennen lassen. Auch die Parameter der Visualisierung wie Belegung der Achsen, Auswahl der Glyphen usw. können durch den Nutzer bestimmt werden, um Beziehungen zwischen Daten oder Datenattributen zu erkennen. Generell können alle Mapping-Prozesse der Visualisierungspipeline durch den Nutzer angepasst werden (North 2005, 1241).

3.9 Visualisierungen im Web

McKeon (2009) wendet das klassische Visualisierungs-Referenzmodell für verteilte Visualisierungsanwendungen im Web an. Durch die Integration im Web steht die Visualisierung nicht mehr isoliert da. Auf den verschiedenen Stufen des Modells ergeben sich durch Standards für Daten und Medien und verteilte Webapplikationen neue Möglichkeiten, Visualisierungen interaktiv zu integrieren. Auf der ersten Stufe stehen verschiedene Datenquellen und-applikationen im Web zur Verfügung. Datengrundlagen können zum Beispiel HTML-Tabellen, Tabellen aus Online-Office-Applikationen, RSS-Feeds oder Online-Datenbanken sein. Diese können auf zweiter Stufe in Mashup-Services wie Yahoo! Pipes oder DabbleDB zu einem integrierten Datensatz arrangiert werden. Auf dritter Stufe stehen dann verschiedene Visualisierungs-Anbieter (IBM Many Eyes, Google Maps, Google Visualization API), die es ermöglichen, den Datensatz auf verschiedene Weise zu visualisieren. Diese Ansichten können nun auf unterschiedliche Arten aggregiert, arrangiert und weitergenutzt werden, z. B. in Dashboards und Wikis oder auf Social Media-Plattformen wie sozialen Netzwerken und Blogs. Das ermöglicht die Integration der Visualisierungen als Diskussionsgrundlage.

Heer u.a. (2009) verweisen dabei besonders auf den kollaborativen Aspekt von Visualisierungen im Web unter dem Mantra „Point, Talk, Publish“. Daten und Visualisierungen können in Onlinesystemen veröffentlicht werden und Interessierte können zeit- und ortsunabhängig interessante Aspekte und Muster in Visualisierungen analysieren und diskutieren. Analog zur kollaborativen Analyse von Daten und Visualisierungen außerhalb des Internets sollen auch online Teilnehmer in der Lage sein auf besondere Aspekte in Visualisierungen hinzuweisen („Point“), darüber zu sprechen („Talk“) und verschiedene Ansichten von Visualisierungen zu publizieren („Publish“).

Als Diskussionsgrundlage muss eine gemeinsame Ansicht der Daten vorliegen. In interaktiven Visualisierungen müssen dafür die Daten und Kontrollparameter wie Zoom-Stufe, Farbauswahl usw. reproduzierbar sein. In webbasierten Systemen können URLs genutzt werden, um die spezifischen Ansichten festzulegen, mit der Möglichkeit die URL über Email, Instant-Messages, Blogs usw. weiterzugeben. Standard für das Hinweisen auf Datenbereiche („Point“) ist die Brushing-Technik, wobei ein Teil der Daten selektiert und hervorgehoben wird. Bei visuellen Effekten wie Animation können sichtbare Spuren helfen, den Weg der ausgewählten Daten zu verfolgen. Zudem sind grafische Annotationen wie das Markieren von bestimmten Daten mit Freihandformen ein wünschenswerter Effekt.

Die asynchrone Diskussion („Talk“) über Beobachtungen, Fragen und Hypothesen der Daten wird über die Kopplung von bestimmten Ansichten und dazugehörigem Text erlaubt. Das kann dadurch geschehen, dass Text mithilfe

der URLs auf bestimmte Marker in der Visualisierung verweist oder passender Text auch bei der Exploration der Daten in der Visualisierung angezeigt wird.

Wesentlich für die vorherigen Punkte ist die Veröffentlichung („Publish“) von Daten, Visualisierungen, Kommentaren und Diskussionen im Web. Systeme wie *IBM Many Eyes* (Viegas u.a. 2007) erlauben die Veröffentlichung und Präsentation von Informationen mit verschiedenen Visualisierungen. Nutzer können ihre eigenen Daten hochladen und mit der Nutzercommunity teilen. Dabei können verschiedene Visualisierungstypen wie Diagramme, Treemaps oder Tagclouds für die Anzeige gewählt werden. Diskutiert werden kann sowohl direkt auf derselben Seite als auch durch die Integration von interaktiven Varianten in Blogs, so dass eine im Web verteilte Diskussion stattfinden kann.

Eine Reihe von Toolkits kann verwendet werden, um Dashboards mit Visualisierungen zu erstellen. *Tableau* (Mackinlay, Hanrahan und Stolte 2007) oder *Spotfire* (Ahlberg 1996) sind kommerzielle Lösungen, welche die Erstellung von Dashboards mit unterschiedlichen Visualisierungstypen ermöglichen. *Dashiki* (McKeon 2009) ist eine wiki-basierte kollaborative Plattform zur Erstellung von Visualisierungs-Dashboards. Benutzer können hier Visualisierungen integrieren, die Live-Verbindungen zu Datenquellen enthalten. *Exhibit* (Huynh, Karger und Miller 2007) ist ein leichtgewichtiges Framework zur einfachen Veröffentlichung von strukturierten Daten im Web. Benutzer können Daten importieren und in verschiedenen Ansichten wie Karten, Tabellen, Thumbnails und Zeitleisten präsentieren. Kombinierte Eigenschaften von Informationsvisualisierung wie das Erstellen thematischer Dashboards, Techniken wie Zooming und Filtering können in einer Präsentation kombiniert werden, um eine bestimmte Idee oder Geschichte zu präsentieren und nutzerfreundlich aufzubereiten (Gershon und Page 2001). Es existieren außerdem eine Reihe von Visualisierungs-Toolkits, die es ermöglichen, Visualisierungen in unterschiedlichen Programmiersprachen zu erstellen: das InfoVisToolkit (Fekete 2004), Prefuse/Flare (Heer, Card und Landay 2005), Processing (Reas und Fry 2003) oder Protovis (Bostock und Heer 2009).

3.10 Visualisierungen für die Informationssuche

Das vorherrschende Paradigma für die visuelle Exploration von Daten in Visualisierungen ist das Information-Seeking-Mantra von Shneiderman (1996): „*Overview first, zoom and filter, then details-on-demand*“. Er unterscheidet in seiner *Task by Data Type Taxonomy* sechs Datentypen (1-,2-,3-dimensionale Daten, zeitliche und multidimensionale, Baum- und Netzwerk-Daten), wobei alle Datentypen eine gewisse Anzahl an Attributen enthalten. Die initiale Suchanforderung besteht nun darin, Elemente zu filtern, deren Attribute bestimmte Werte annehmen. Dafür schlägt Shneiderman einen Prozess in sieben Schritten vor:

- *Overview*: Erhalt eines Überblicks über die gesamte Kollektion

- *Zoom*: Zoomen auf Elemente von Interesse
- *Filter*: Herausfiltern von uninteressanten Elementen; Dynamic Queries werden als Hauptmethodik für die interaktive und schnelle Filterung der Daten genannt
- *Details-on-demand*: Auswahl eines Elementes oder einer Gruppe; Details-on-demand Pop-Ups können HTML-Links zu weiteren Informationen enthalten
- *Relate*: Anzeige von Beziehungen zwischen den Elementen
- *History*: Aufzeichnen einer Historie von Aktionen, um Operationen wie Undo, Replay oder Refinement zu unterstützen
- *Extract*: Nutzerunterstützung für die Extraktion von Sub-Kollektionen und der Filter-Parameter

Keim (2002) hat einen sehr ähnlichen Ansatz und unterscheidet in seiner Klassifikation der Informationsvisualisierung und der Data-Mining-Techniken zwischen (i) dem Datentyp der visualisiert wird, (ii) den genutzten Visualisierungstechniken und (iii) den Interaktions- und Distortion-Techniken. Die Interaktionstechniken die unter Punkt (iii) aufgezählt werden, sind: (1) Interactive Projection (siehe Rearranging und Remapping Abschnitt 3.8.4), (2) Interactive Filtering, (3) Interactive Zooming, (4) Interactive Distortion (Fokus & Kontext-Techniken Abschnitt 3.7.3), (5) Interactive Brushing-and-Linking.

Zahlreiche Systeme mit unterschiedlichen Datentypen und verschiedenen Visualisierungstechniken werden genutzt, um Informationen darzustellen und nach dem Information-Seeking-Mantra zu durchsuchen. Systeme geordnet nach Art der Informationsstruktur wurden in Abschnitt 3.5 vorgestellt, weitere Übersichten bieten auch Shneiderman (1996), Keim (2002) und North (2005).

Ein Beispiel für das Filtern heterogener Informationen ist das System FacetMap (Smith u.a. 2006). Hier werden in einer Treemap-ähnlichen Anzeige Top-Level Facetten für die interaktive Filterung von persönlichen Informationen wie Webseiten, Emails, Fotos, Dokumente etc. genutzt. In einer Übersicht werden Facetten wie Informationstyp, Datum, Personen, Ort etc. angezeigt, wobei die Größe der Facetten die Anzahl der Informationselemente kodiert. Durch einen Mausklick auf die gewünschte Facette filtert der Nutzer bis zu den gewünschten Informationen und bekommt sie direkt als Thumbnail angezeigt. Als Alternative zu diesem Browsing-Ansatz kann der Nutzer auch direkt nach den Informationselementen suchen.

3.10.1 Facettierte Filterung von Information im Web

Das Information-Seeking-Mantra wird auch für die Repräsentation und Filterung von Information im Web genutzt. Dafür werden populäre Visualisierungen im Web wie Tagclouds, Karten, Diagramme, aber auch anderen Arten genutzt.

Wood et al. (2007) nutzen die kombinierte Anzeige von *Tag Clouds* und *Tag Maps* für die interaktive Anzeige von Log-Daten einer mobilen Applikation für

lokale Information in Google Earth. Der Service auf dem Mobiltelefon schlägt standortabhängig Geschäfte, Lokale oder Dienstleister in der Nähe vor. Die Log-Dateien beinhalten unter anderen Informationen zu Zeit, Ort, User-ID und Name des gewählten Services. Diese Daten können nun ortsabhängig in Google Earth angezeigt werden. Beispielsweise zeigt eine Tagcloud die wichtigsten Lokalitäten für einen Kartenausschnitt an (Größe der Tags kodiert Anzahl des ausgewählten Services durch den Nutzer). Klickt man auf eine Tag, so wird in Google Earth auf den Kartenausschnitt gezoomt, welches diese Tags enthält und zeigt nun weiterhin *auf* der Karte Tags an (Tag Maps), welche Services die Nutzer in dieser Region gewählt haben. Tagclouds werden also auch hier für die Filterung und den Übergang von einer Übersicht zu Details von Interesse genutzt. Tag Maps können auch animiert werden, um den Verlauf über eine Zeitperiode anzuzeigen.

Das System VisGets (Dörk u.a. 2008) kombiniert verschiedene Visualisierungen wie Karte, Tagcloud und Säulendiagramm zur Repräsentation und Filterung von abgerufenen Webressourcen wie RSS-Feeds, Wikipedia-Einträge oder Flickr-Fotos anhand von Facetten wie Zeit, Ort und Thematik. Basierend auf dem Konzept von Dynamic Queries können Resultate interaktiv durch Manipulation der Visualisierungen gefiltert werden. VisGets implementiert auch Coordinated Views. Fährt der Nutzer mit der Maus über ein visuelles Element, werden alle verbundenen Elemente in der Visualisierung in der Trefferliste hervorgehoben. Der neu eingeführte Ansatz Weighted Brushing wird genutzt, um stark verbundene Elemente stärker als schwach verbundene Elemente hervorzuheben.

Einen ähnlichen Ansatz bietet ein System für die dynamische Anzeige von Twitter-Nachrichten, die zu einem bestimmten Event oder Thema gepostet werden (Dörk u.a. 2010). In drei verschiedenen Visualisierungen werden Facetten der Nachrichten angezeigt, die auch wieder zur Filterung der Nachrichten genutzt werden können und untereinander mit Brushing-and-Linking verbunden sind. Die Hauptvisualisierung *Topic Streams* ist ein gestapelter Graph, der Nachrichten-Topics über die Zeit anzeigt und sich dynamisch anpasst. Die *People Spiral* zeigt die Teilnehmer der Diskussion und ihren Aktivitätsgrad, in der *Image Cloud* werden populäre Fotos zum Thema nach Popularität unterschiedlich groß angezeigt. Zusammen mit der Liste der Twitter-Nachrichten werden die Grafiken dynamisch aktualisiert und zeigen damit die aktuelle Diskussion, aber auch die vergangenen Thematiken.

3.10.2 Links innerhalb des Informationsraums

Informationen die mit dem Information-Seeking-Mantra gefiltert wurden, können auf Detail-Ebene Links zu verwandten Informationen enthalten, die im Sinne von Punkt 5 des Information-Seeking-Mantras: *Relate* oder der IR-Retrieval-Technik *Browsing* vom Nutzer verfolgt werden können, um im Navi-

gationsraum zu navigieren. Ein gutes Beispiel dafür ist das System BrainGazer (Bruckner u.a. 2009), das Volumendaten visualisiert, die durch konfokale Mikroskopie erworben wurden. Zum Beispiel kann das Gehirn einer Fruchtfliege mit annotierten anatomischen Strukturen visualisiert werden, um explorierbar zu machen, wie anatomische und physiologische Zusammenhänge im Nervensystem das Verhalten beeinflussen. Das System bietet drei grundlegende Arten von visuellen Abfragen: (1) Mit einem Mausklick auf ein Objekt werden Informationen zur Struktur und Links zu verwandten Objekten angeboten (semantic queries). Damit kann der Nutzer durch den Informationsraum navigieren. (2) Außerdem werden Links zu Objekten angezeigt, die sich in der räumlichen Nähe des Objektes befinden (spatial queries). (3) Der Nutzer kann einen Pfad frei Hand auf die Visualisierung zeichnen, wodurch Informationen über die ausgewählten Objekte angezeigt werden (path queries).

3.10.3 Graphisches Ergebnisretrieval

Einen ergänzenden Ansatz für die visuell gestützte Recherche bietet die Systemkomponente Wing-Graph (Wolff 1996) im Forschungsprojekt Wing-IIR, welches die Unterstützung der multimodalen Faktenrecherche nach Werkstoffdaten untersucht. Dabei verbindet das System verschiedene Zugangswege (multimodale Systemgestaltung) zu Datenbanken in einer Benutzerschnittstelle: (1) der natürlichsprachliche Zugang für die Recherche und (2) eine direktmanipulative graphische Benutzerschnittstelle als allgemeine Arbeitsumgebung und für die Recherche. Die Komponente Wing-Graph erlaubt es, Werkstoffparameter als Liniendiagramme darzustellen und graphische Operationen direkt im Liniendiagramm als visuelle Definition für den Anfrageaufbau zu nutzen. Dabei soll das visuelle Denken der Benutzer für das Retrieval in Faktendatenbanken genutzt werden, wobei die Recherchesituation, der bereits ein Retrievalprozess vorangegangen ist, im Mittelpunkt stehen soll (graphisches Ergebnisretrieval).

Der konkrete Anwendungsbereich ist dabei die Domäne der Werkstoffinformation. Hier lag die Beobachtung vor, dass die graphische Informationsdarstellung von Werkstoffparametern als Kurvendarstellung in der Literatur, bei der Informationsbeschaffung und -interpretation weit verbreitet ist und als sehr wichtig eingestuft wird. Fachexperten sind in der Lage, ein bestimmtes Kenngrößenverhalten eines Werkstoffes als Kurve auch ohne Vorlage zu skizzieren. Auch nach der Suchphase nutzen Werkstofffachleute visualisierte Informationen bei der Dateninterpretation und für die Recherche nach ähnlichen Werkstoffen. Dabei sind in der Recherche vergleichende, sich auf vorgegebene Werkstoffe und Daten beziehende Fragestellungen besonders relevant.

Bei Standardsystemen zur Werkstoffrecherche liegt für diese und ähnliche Fragestellungen eine kognitive Bruchstelle zwischen Ergebnisdarstellung/-interpretation und erneuter Anfragestellung vor. Die Werkstoffkenngrößen

werden grafisch als Liniendiagramme dargestellt. Fachexperten denken empirisch beobachtet visuell und möchten ähnliche Werkstoffe anhand dieser Vorlage finden. Dieses Informationsbedürfnis müssen sie nun wieder vom grafischen Modus in einen textuellen Modus (Modalitätswechsel) und in das Schema der Anfragesprache (z. B. Query-by-example-Formulare oder Boolesche Suchanfragen) umwandeln. Dabei entsteht eine erhöhte mentale Belastung für die verschiedenen Umwandlungsprozesse und Information kann verloren gehen, da sich visuelle Vorstellung nur mit erhöhtem Aufwand in exakte sprachliche Konstrukte durch den Nutzer umwandeln lässt.

Das System Wing-Graph bietet nun die Möglichkeit Anfragen direkt in der Ergebnisvisualisierung zu definieren. Durch unterschiedliche Ansätze wie (1) abstrakte Verfahren (Setzen mehrerer Suchpunkte oder eines Streubandes), relative Verfahren (Suche nach ähnlichen Werkstoffen durch Setzen eines Streubandes) oder produktive Verfahren (Skizzieren einer Werkstoffkurve als Suchhypothese) können Folgeanfragen direkt in der Ergebnisdarstellung definiert werden. Durch die Zusammenlegung von Ergebnis- und Anfragedisplay lassen sich die Probleme beim Übergang von Ergebnisinterpretation zur Anfragedefinition vermindern oder beseitigen. Dabei werden insbesondere Folge-recherchen unterstützt, in denen ein Ähnlichkeitscharakter zu bestehenden Ergebnissen vorliegt. Durch die Anfrageerstellung direkt in den Ergebnissen ergibt sich ein Kreislaufmodell, bei der die Visualisierung gleichzeitig für die Anfrage- und Ergebnisdarstellung genutzt wird.

Dabei setzt das graphische Ergebnisretrieval auf Pinkers Modell der Graphenwahrnehmung auf (vgl. Abschnitt 4.7) und wird um den Rechercheaspekt erweitert. Dazu werden zwei Komponenten hinzugefügt: (1) Eine externe Aufgabenstellung kann die konzeptuellen Fragen an die Informationsdarstellung beeinflussen und (2) die grafische Information kann ein neues Informationsbedürfnis beim Nutzer auslösen, das durch die visuelle Anfragedefinition definiert wird.

Die Leistungsfähigkeit des graphischen Ergebnisretrievals erklärt Wolff durch die duale Kodierung der Information. Visuelle Informationen werden im Gehirn doppelt als visuelle und als sprachliche Repräsentation kodiert (vgl. Abschnitt 4.3) und zu höherwertigen visuellen Einheiten aggregiert, mit denen sich im Kurzzeitgedächtnis visuelle Operationen durchführen lassen (vgl. Abschnitt 4.7). Das Informationsbedürfnis nach ähnlichen Werkstoffen innerhalb einer bestimmten Parameterspanne lässt sich mit diesen visuellen Operationen innerhalb der Graphendarstellung visuell vorstellen und in eine visuelle Anfragedefinition überführen. Dabei ist die Durchführung der visuellen Operation und Übertragung in die Diagrammdarstellung kognitiv weniger aufwendig und vom Benutzer selbst leistbar, als es durch die Dekodierung von transformierten visuellen Elementen zu komplexen sprachlichen Konstrukten und der erneuten Enkodierung in eine formale Anfragesprache machbar wäre.

Weitere Eigenschaften von IR-Retrieval wie Vagheit der Anfrage werden auch hier behandelt (Wolff und Womser-Hacker 1997), indem visuelle Anfragen nicht direkt in numerische Datenbankabfragen umgewandelt werden, sondern vorher mit Domänen-Modellen vager Konzepte und Regelsammlungen behandelt werden, um die Intention und Unschärfe der visuellen Anfrage besser zu berücksichtigen.

3.10.4 Visualisierungen im Document Retrieval

Eibl (2003) unterscheidet in seiner Kategorisierung von Visualisierungen im Document Retrieval zwischen (1) der Visualisierung der Anfrage, (2) der Visualisierung der Ergebnismenge und (3) der integralen Visualisierung von Anfrage und Ergebnismenge. Die Visualisierung der Anfrage wurde als Forschungsfokus in den 1980er Jahren von einigen Systemen umgesetzt (vgl. Eibl 2003, 58ff) und stützt sich meist auf das Boolesche Suchmodell. Der Anwendungsschwerpunkt der Visualisierung soll hier in der Unterstützung des Nutzers bei der Anfrageformulierung mit Booleschen Operatoren liegen. Da sich der Einsatz von Booleschen Operatoren bei komplexen Anfragen für Laien als auch Experten äußerst schwierig gestaltet und eine Diskrepanz zwischen natürlichsprachlicher und logischer Verwendung existiert, soll die Visualisierung den Nutzer bei der Anfrageformulierung unterstützen (vgl. Eibl 2002, 142).

Systeme für die Visualisierung von Ergebnismengen wurden seit Beginn der 1990er Jahre entwickelt (Eibl 2003, 57). Eibl unterscheidet grob zwischen (1) der Darstellung des Informations- oder Dokumentenraumes und (2) der Darstellung des Dokumenteinhaltes und wie sich Suchkriterien darin verhalten. Beispielsysteme für die Darstellung des Dokumentenraums wie WEBSOM, VisIslands, INSPIRE, Infosky, ThemeRiver, VR-VIBE oder LyberWorld wurden bereits in Abschnitt 3.5.4 aufgeführt (vgl. auch Eibl 2003, 62ff). Beispielsysteme für die Verteilung von Suchhäufigkeiten in Dokumenten stellt Eibl (2003, 79ff) vor. Eibl (2003) entwickelt das System DEVID für die integrierte Darstellung von Suchanfrage und Suchergebnis in einer Oberfläche mit dem Fokus auf Grafikdesign und Softwareergonomie. Dabei wurde mit Interviews und Beobachtung von professionellen Datenbankanwendern in der sozialwissenschaftlichen Domäne festgestellt, dass Nutzer am besten bei der Anfrageformulierung und mit direktem Feedback über die Größe des Suchergebnisses unterstützt werden können. Die Nutzer reformulierten die Boolesche Suchanfrage typischerweise so oft, bis sie ein Ergebnis von 20-30 Dokumenten erreicht hatten. Dabei traten Schwierigkeiten bei der Formulierung von komplexen Booleschen Suchanfragen auch für Experten auf. Diese waren begründet durch die Unterschiede in der natürlichsprachlichen und streng logischen Verwendung von Booleschen Operatoren, durch die Verwendung des „Not“-Operators und durch die komplexe Klammerung innerhalb von Suchanfragen. Als gestalterischer Ausgangspunkt wurde das Forschungssystem InfoCrystal (vgl. Eibl 2003, 103ff) genutzt. DEVID nutzt die

Grundidee der visuellen Kodierung von Suchbegriffen und deren Boolescher Ergebniskombination und nimmt eine starke Vereinfachung der visuellen Sprache vor. Deskriptoren können vom Nutzer unter einen farblich kodierten Winkel als Sinnbild für eine Karteikarte eingegeben werden und die Größe des Suchergebnisses wird direkt dort neben dem Eingabefeld angezeigt. Begriffe innerhalb eines Eingabefeldes werden mit „OR“ verknüpft, die Art des Eingabefeldes (Titel, Schlagwort, Ort, Datum usw.) lässt sich auswählen. Bei der Eingabe von Suchbegriffen in mehr als zwei Eingabefeldern wird die Boolesche „AND“-Verknüpfung gebildet und die Größe des Suchergebnisses unter der Kombination der beiden farblich kodierten Winkel in einer neuen Spalte angezeigt. So können die Endnutzer direkt erkennen, wie sich die Boolesche Kombination auf die Ergebnismenge auswirkt. Durch das Ausfüllen mehrerer Felder und der direkten Kombination der Deskriptoren lassen sich komplexe Anfrage bilden und direkt die Größe der Ergebnismenge ablesen. Durch Interaktion mit dem Mauszeiger über Felder oder Kombinationen werden nicht benötigte visuelle Elemente ausgeblendet. In einer erweiterten Version kann die Wichtigkeit von Deskriptoren für das Suchergebnis durch Verschiebung auf der X-Achse angezeigt werden oder der Nutzer kann die Wichtigkeit durch Verschieben selbst bestimmen. Vages Retrieval wird interaktiv in der Oberfläche unterstützt, indem durch ein Kontext-Icon basierend auf einer Similaritätsfunktion, weitere dem Suchergebnis ähnliche Dokumente gesucht werden, ohne die Suchanfrage zu ändern. In einer vergleichenden Evaluation mit neun Topics basierend auf dem GIRT-Korpus als Teil der TREC-Initiative wurde angegeben, dass DEVID den Vergleichssystemen Messenger und freeWais in Recall und Precision überlegen war; die Anzahl der durchschnittlichen Reformulierungen war vergleichsweise gering und die Ergebnisse der Fragebogenauswertung zu Einfachheit, Gestaltung, Präsentation und Nutzungsfreude des Systems waren sehr positiv. Kritisch lässt sich anmerken, dass der Einfluss auf Precision und Recall nur auf Ebene der Benutzungsoberfläche durch die Anfrageformulierung der Testpersonen bestimmt wurde. Weitere wichtige Einflussfaktoren wie der Rankingalgorithmus, die Visualisierung selber, Interaktion, ästhetische Gestaltung usw. wurden nicht näher untersucht.

3.11 Zusammenfassung

Die Vorteile und Ziele von Visualisierungen reichen von der Unterstützung einfacher Mustererkennung bis zum Einsatz als Werkzeug, um die Limitierungen des kognitiven Prozesses abzumildern. Im Knowledge Crystallization-Prozess wird abgebildet, auf welchen unterschiedlichen Stufen im Prozess der Wissensgenerierung Visualisierungen eingesetzt werden können, um den kognitiven Prozess zu verstärken. Herausgestellt werden kann, dass komplexe heterogene Informationen aus verschiedenen Datenquellen gesammelt, geordnet und in einer externen Repräsentation (Text, Bilder, Tabellen, Visualisierung

gen usw.) und einem internen Schema (mentalen Modell) repräsentiert werden, um daraus Informationen abzuleiten und Entscheidungen zu treffen. Verschiedene Informationsstrukturen wie tabellarisch, räumlich/zeitlich, baum-/netzwerkförmig und Text werden für die initiale Visualisierung und die Bildung eines mentalen Modells genutzt. Komplexere Informationsstrukturen können in kombinierten Ansichten angezeigt werden. Ein Modell dafür ist das Snap-Visualization-Modell, das eine Analogie zwischen dem relationalen Datenbank-Modell und koordinierten Ansichten bildet.

Der Transfer von Visualisierungen in das Internet verstärkt die Ausgangslage von verteilten Datenquellen und heterogenen Informationen, bietet aber die Möglichkeit, diese mit Online-Tools zu aggregieren und anzuzeigen. Die entstandenen Visualisierungen können als kollaborative Arbeitsgrundlage verwendet und als Komponenten in Blogs oder Dashboards weitergenutzt werden.

Die Methodik der Informationssuche richtet sich in der Informationsvisualisierung stark nach dem Mantra von Shneiderman, d.h. ausgehend von einer Übersicht der Daten wird in den Systemen immer weiter zu Punkten von Interesse gefiltert. Als Technik für die Anzeige von Relationen zwischen verschiedenen Daten in Visualisierungsansichten werden Multiple Coordinated Views und die Brushing-and-Linking-Technik eingesetzt.

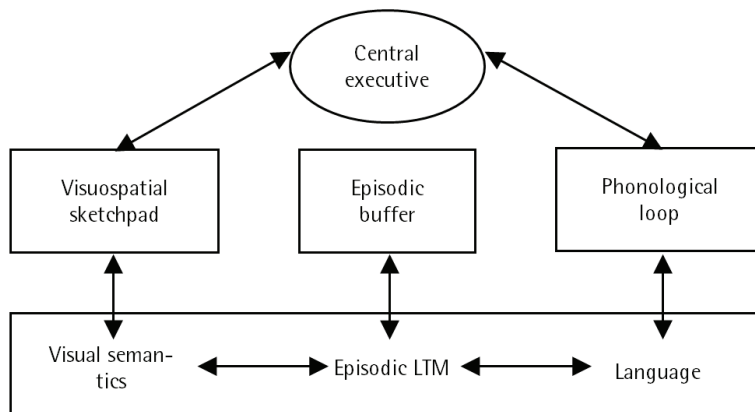
4. Informationsverarbeitung

Die beiden vorangegangenen Abschnitte haben die grundlegenden Modelle und Eigenschaften von Informationssuche und Informationsdarstellung vorgestellt. Beide Bereiche setzen sich zum Ziel, das Informationsbedürfnis des Nutzers mit dem Finden und der Darstellung von Information zu befriedigen. Um zu verstehen wie verschiedene Informationen wie Text, Bild und interaktive Visualisierungen im menschlichen Gehirn verarbeitet werden, muss der kognitive Prozess und die grundlegenden menschlichen Eigenschaften der Informationsverarbeitung untersucht werden. Dafür wird in diesem Abschnitt zuerst der prinzipielle Aufbau des Gedächtnisses beschrieben und die Prinzipien der dualen Kodierung und kognitiven Last eingeführt. Auf dieser Basis bilden sich zwei Theorien des multimedialen Lernens, die das Erfassen und Verständnis von Text und Bildern in Modellen kombinieren. Aus den Modellen lassen sich Prinzipien ableiten, welche die kognitive Last verringern sollen und damit den Lerneffekt erhöhen. Informationen verschiedener Modalitäten werden in einem mentalen Modell abgebildet. Darauf aufbauend wird gezeigt, wie Nutzer einfache bis komplexe Visualisierungen kognitiv verarbeiten und in einem qualitativen mentalen Modell integrieren, um daraus Informationen abzuleiten.

4.1 Aufbau des Gedächtnisses

Beim Aufbau des Gedächtnisses wird unterschieden zwischen dem Ultrakurzzeitgedächtnis, dem Kurzzeitgedächtnis (in anderen Modellen auch Arbeitsgedächtnis genannt) und dem Langzeitgedächtnis (vgl. Brand und Markowitsch 2004). Das Ultrakurzzeitgedächtnis hat eine Vorhaltezeit von Millisekunden und speichert Wahrnehmungserlebnisse. Visuell oder auditiv Wahrgenommenes wird durch verschiedene Kanäle als sinnlicher Eindruck kurzzeitig festgehalten. Das Kurzzeitgedächtnis hat eine Vorhaltezeit von Sekunden bis sehr wenige Minuten und unterliegt verschiedenen Restriktionen. Zum Beispiel kann es nur ungefähr sieben plus/minus zwei Informationseinheiten gleichzeitig behalten (Miller 1956). Eine Aufteilung des Kurzzeitgedächtnisses in ein komplexeres System mit operierenden Einheiten bietet das Arbeitsgedächtnismodell von Baddeley (1986, 2000). In diesem Modell wird unterschieden zwischen dem (1) räumlich-visuellen Notizblock, der visuelle Eindrücke speichert, (2) der phonologischen Schleife, die verbale Informationen speichert und (3) dem episodischen Buffer, der sowohl visuelle als auch verbale Informationen kurzfristig in Form von Episoden speichern kann. Die zentrale Exekutive verwaltet die verschiedenen Systeme und steuert die Verknüpfung mit Systemen im Langzeitgedächtnis (vgl. Abbildung 6).

Abb. 6: Gedächtnismodell



Nach Baddeley (2000, Abb. 1).

Das Langzeitgedächtnis hat eine sehr hohe Aufnahmekapazität und die Dauer der Speicherung kann für ein ganzes Leben reichen. Dabei wird zwischen folgenden inhaltlichen Einheiten des Langzeitgedächtnisses unterschieden (vgl. Aufstellungen von Brand und Markowitsch 2004; Meier 1999; und Gedächtnismodellen von Squire u.a. 1993; Squire 1992):

- *Prozedurales Gedächtnis*: Speicherung von motorischen Fähigkeiten und Routinehandlungen (Tulving 1985).
- *Primingsystem*: Das unbewusste Erkennen von Objekten oder Geräuschen aufgrund ihrer wahrgenommenen Charakteristika (Tulving und Schacter 1990).
- *Perzeptuelles Gedächtnis*: Das bewusste Erkennen von Objekten oder Geräuschen aufgrund ihrer wahrgenommenen Charakteristika (Tulving und Schacter 1990).
- *Semantisches Gedächtnis*: Das Wissenssystem, das Weltwissen und Fakten kontextfrei abspeichert (Tulving 1972).
- *Episodisches Gedächtnis*: Speicherung von persönlichen Erlebnissen mit Raum-/Zeitbezug und emotionalem Anteil (Tulving 1972).

Für die Aufnahme von Informationen in das Langzeitgedächtnis wird zwischen folgenden Prozessen unterschieden (Brand und Markowitsch 2004):

- Die Einspeicherung (Enkodierung) von Informationen,
- die Konsolidierung (Festigung) und Ablagerung von Informationen durch Einbettung in bereits bestehende Netzwerke von Gedächtniseinheiten, und
- den Abruf (Erinnerung) von Informationen.

4.2 Cognitive Load Theorie

Die Cognitive Load Theorie (CLT) wurde von Chandler und Sweller (1991) anhand empirischer Untersuchungen entwickelt. Dabei ist Lernen mit kognitiver Belastung verbunden. Je niedriger die kognitive Belastung ist, umso besser ist der Lerneffekt. Das Arbeitsgedächtnis ist für die Informationsverarbeitungsprozesse verantwortlich, das Langzeitgedächtnis hält die verarbeiteten Informationen in Schemata vor. Restriktiv für den Lernprozess sind die Limitierungen, die das Arbeitsgedächtnis und das Langzeitgedächtnis haben. Ist die kognitive Belastung beispielsweise größer als der Speicher im Arbeitsgedächtnis, entsteht eine Überlastung und der Lerneffekt wird negativ beeinflusst.

Die Kapazität und Vorhaltezeit des Arbeitsgedächtnisses ist beschränkt. Laut Miller (1956) kann das Kurzzeitgedächtnis sieben plus/minus zwei Informationen gleichzeitig speichern oder zwei bis vier Informationen gleichzeitig verarbeiten. Peterson und Peterson (1959) fanden heraus, dass ohne Wiederholung die Informationen im Kurzzeitgedächtnis innerhalb von zwanzig Sekunden verloren gehen. Lernen ist definiert als die Änderung von Schemata im Langzeitgedächtnis. Neue Informationen werden im Arbeitsgedächtnis verarbeitet und als Schemata im Langzeitgedächtnis gespeichert und mit anderen Schemata verknüpft. Schemata sind kognitive Konstrukte, die es erlauben, mehrere Informationen als eine Einheit zu kategorisieren.

Es existieren drei verschiedene Faktoren, die zusammen die kognitive Belastung bestimmen: die lernbezogene (germane), die extrinsische (extraneous) und die intrinsische (intrinsic) Belastung (loads). Die verschiedenen Faktoren

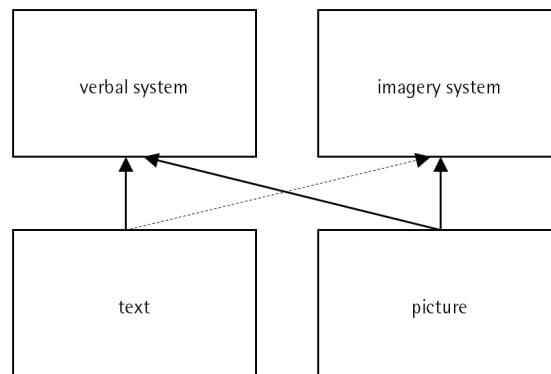
sind additiv und bestimmen gemeinsam die kognitive Belastung beim Lernprozess (vgl. Sweller 2005). Die lernbezogene Belastung ist direkt mit dem Lernprozess verbunden. Die kognitive Belastung besteht aus dem Verstehen des Lernmaterials und dem Aufbau von Schemata im Langzeitgedächtnis. Belastungsfaktoren, die das Lernen zum Positiven oder Negativen beeinflussen, sind die Aufbereitung des Lernmaterials, die Aufmerksamkeit, das Vorwissen und die Motivation des Lernenden.

Die extrinsische Belastung wird verursacht durch falsche Aufbereitung des Lernmaterials, wenn Limitierungen des Arbeitsgedächtnisses und der Vorgang des Aufbaus von Schemata im Langzeitgedächtnis nicht beachtet werden. Ist die Darstellung des Lernmaterials fehlerhaft, muss zu viel Zeit auf die Informationssuche aufgewendet werden. Auch zu viele unnötige Wiederholungen und Erklärungen erhöhen die kognitive Belastung. Durch die Anwendung der Multimediaprinzipien (Sweller 2005; vgl. Kapitel 0) kann die extrinsische Belastung gesenkt werden. Die intrinsische Belastung entsteht durch den Lerninhalt selber. Je schwerer der Inhalt, desto höher die Belastung. Je mehr die einzelnen Lerninhalte miteinander verknüpft sind und interagieren, umso höher ist die kognitive Belastung (vgl. Sweller 2003).

4.3 Duale Kodierungstheorie

Die duale Kodierungstheorie geht auf Paivio (1986) zurück. Er stellt ein Modell für die kognitive Verarbeitung von Texten und Bildern auf. Dabei geht er von der Hypothese aus, dass zwei unabhängige kognitive Systeme existieren: das verbale und das bildliche System (vgl. Abbildung 7).

Abb. 7: Duale Kodierung



Nach Paivio (1986).

Bei der Verarbeitung von Texten werden diese im verbalen System kodiert. Bilder hingegen werden bei nicht zu kurzen Präsentationszeiten doppelt kodiert

abgespeichert, einmal im bildlichen System und zusätzlich im verbalen System. Die duale Kodierung bei Bildern führt dabei zur besseren Memorierung als die einfache Kodierung von Text. Zudem ist der Größenvergleich von Objekten in grafisch kodierter Form schneller als bei Wörtern, da das bildliche System für den Vergleich verantwortlich ist.

4.4 Das CTML-Modell

Basierend auf der dualen Kodierungstheorie nach Paivio (1986), der Theorie der kognitiven Belastung nach Chandler und Sweller (1991) und der Theorie des aktiven Prozesses (Mayer 2001; Wittrock 1989) stellt Mayer (2005, 2001) eine Theorie des multimedialen Lernens auf: Aus psychologischer Sicht versteht man darunter das kombinierte Erfassen und Verstehen von Text und Bildern. Dabei wird Text als Synonym für gesprochene und geschriebene Sprache gebraucht. Webbasiertes multimediales Lernen verschiebt Texte und Bilder in den Kontext des Computers und des Internets. Beim multimedialen Lernen werden verschiedene externe Repräsentationen wie gesprochener oder geschriebener Text, Zeichnungen, Bilder, Fotos, Diagramme, Klang oder Ton als Informationsquelle für die Bildung eines internen mentalen Modells genutzt. Das mentale Modell wird zuerst im Arbeitsgedächtnis gebildet und dann in das Langzeitgedächtnis überführt.

Die kognitive Theorie des multimedialen Lernens (Mayer 2005) basiert auf drei Annahmen zum menschlichen Lernprozess:

- *Duale Kodierung*: Das menschliche Informationsverarbeitungssystem hat zwei Kanäle für die Verarbeitung von Text und Bild, den visuellen/bildlichen Kanal und den auditiven/verbalen Kanal.
- *Begrenzte Kapazität*: Jeder Kanal hat eine begrenzte Kapazität für die Verarbeitung und Speicherung von Informationen.
- *Aktiver Prozess*: Bedeutsames Lernen besteht aus einem Paket von fünf kognitiven Prozessen.

Die fünf kognitiven Prozesse aus Punkt 3 sind Folgende:

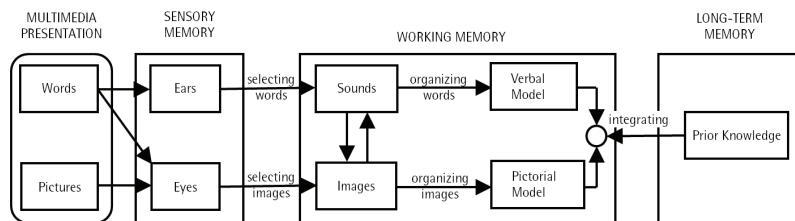
- Auswahl der relevanten Wörter aus dem Text.
- Auswahl der relevanten Bildinhalte aus dem Bild.
- Strukturierung der relevanten Wörter zu einer kohärenten verbalen Repräsentation.
- Strukturierung der Bildinhalte zu einer kohärenten bildlichen Repräsentation.
- Integration der verbalen und bildlichen Repräsentation zu einem Modell mit Hilfe von Vorwissen aus dem Langzeitgedächtnis.

Abbildung 8 zeigt das Modell der kognitiven Theorie des multimedialen Lernens nach Mayer (2005). Die einzelnen Boxen repräsentieren die verschiedenen Gedächtnisstufen: sensorisches Gedächtnis, Arbeitsgedächtnis und Langzeitge-

dächtnis. Text und Bilder werden im Rahmen einer Multimediapäsentation in das sensorische Gedächtnis überführt. Texte werden als geschriebene und gesprochene Sprache über die Wahrnehmung durch Augen und Ohren in das visuelle und auditive sensorische Arbeitsgedächtnis abgelegt. Bilder werden über die Augen in das visuelle sensorische Arbeitsgedächtnis abgelegt. Das sensorische Gedächtnis kann das genaue Abbild von Text- oder Bildfragmenten für sehr kurze Zeit speichern.

Auf der folgenden Stufe findet im Arbeitsgedächtnis der kognitive Prozess zur Auswahl von relevanten Wörtern oder Bildinhalten statt. Dabei wird aktiv ein Fokus auf bestimmte Wörter oder Bildinhalte gelenkt und in das Arbeitsgedächtnis für Ton oder Bild übernommen. Text kann dabei als gesprochene oder geschriebene Sprache über Ohren und Augen aufgenommen werden und wird in dem Arbeitsgedächtnis für Ton oder Bild verarbeitet. Durch kognitive Prozesse kann ein Austausch zwischen dem verbalen und auditiven Arbeitsgedächtnis stattfinden, wenn zu Wörtern mental Klangbilder konstruiert oder wenn Klangbildern Wörter zugeordnet werden. Der Selektionsprozess ist aufgrund der begrenzten Kapazität der Kanäle notwendig. Dabei ist der Selektionsprozess ein aktiver Prozess, da der Lernende entscheiden muss, welche Inhalte fokussiert werden.

Abb. 8: Kognitive Theorie des multimedialen Lernens



Nach Mayer (2005, Abb. 3-2).

Im nächsten Schritt werden die internen Repräsentationen aus dem Arbeitsgedächtnis für Ton und Bild in zwei getrennte mentale Modelle integriert: ein verbales und ein bildliches, mentales Modell. Diese werden dann mithilfe von Vorwissen aus dem Langzeitgedächtnis zu einem integrierten mentalen Modell zusammengeführt.

Basierend auf einer Vielzahl von empirischen Untersuchungen stellt Mayer in *Cognitive Theory of Multimedia Learning* von 2005 die grundlegenden Prinzipien des multimedialen Lernens vor:

- *Prinzip der dualen Kodierung (multimedia principle)*: Der Lerneffekt ist höher bei einer integrierten Darstellung von Text und Bild als nur bei Text.
- *Prinzip der räumlichen Nähe oder Kontiguitätsprinzip I (split-attention principle, spatial contiguity principle)*: Eine integrierte Präsentation von

Text und Bild ist besser als eine getrennte Darstellung. Text und Bild sollten in der Präsentation nahe beieinander stehen.

- *Prinzip der simultanen Darstellung oder Kontiguitätsprinzip II (temporal contiguity principle)*: Die gleichzeitige Darstellung von Text und Bild ist besser als eine sukzessive Präsentation von Text und Bild.
- *Multimodalitäts-Prinzip oder Modalitätsprinzip (modality principle)*: Die audiovisuelle Darstellung von Bild und Text ist besser als geschriebener Text mit einem Bild.
- *Kohärenz-Prinzip (coherence principle)*: Irrelevante, visuelle oder akustische Informationen beeinträchtigen die Lernleistung.
- *Redundanz-Prinzip (redundancy principle)*: Die audiovisuelle Darstellung von Bild ist besser als eine redundante Darstellung von Bild, Ton und Text. Die gleichzeitige Darstellung von geschriebenem und gesprochenem Text beeinträchtigt den Lernerfolg.

4.5 Das ITPC-Modell

Schnotz (2005) vereint verschiedene Modelle und Konzepte aus den Bereichen duale Kodierung, multimediales Lernen und einer kognitiven Architektur mit mehreren Gedächtnisstufen in einem integrierten Modell des Text- und Bildverständnisses. Basis ist ein Zweikanalmodell (Schnotz und Bannert 2003) ähnlich dem Paivios (1986), bei dem zwischen externen und internen Repräsentationen unterschieden wird.

Bei der externen Repräsentation wird differenziert zwischen Deskriptionen und Depiktionen. Deskriptionen bestehen aus Symbolen, die auf Konventionen beruhen und keine Ähnlichkeit mit dem Bezeichneten haben. Beispiele sind mathematische Formeln oder Text als das gängigste Zeichensystem. In Texten werden Substantive als Symbole für Objekte, Verben als Symbol für Relationen, und Adjektive als Symbol für Attribute genutzt. Depiktionen bestehen aus Ikonen und sind Repräsentationen mit Ähnlichkeiten zu dem Bezeichneten; bspw. Fotografien, Zeichnungen, Bilder oder Karten besitzen eine Analogie zu dem Abgebildeten.

Interne Repräsentationen entstehen durch die Perzeption und kognitive Verarbeitung. Beim Lesen von Text entstehen drei mentale Repräsentationsformen:

- 1) *Oberflächenstruktur (text surface representation)*: Der Text kann wiedergegeben werden, wurde aber noch nicht verstanden.
- 2) *Propositionale Repräsentation (propositional representation)*: Die Ideen und Konzepte des Textes können unabhängig von der Wortwahl und der Syntax des Textes wiedergegeben werden.
- 3) *Mentales Modell (mental model)*: Einzelne Ideen und Konzepte werden zu einem mentalen Modell zusammengefügt.

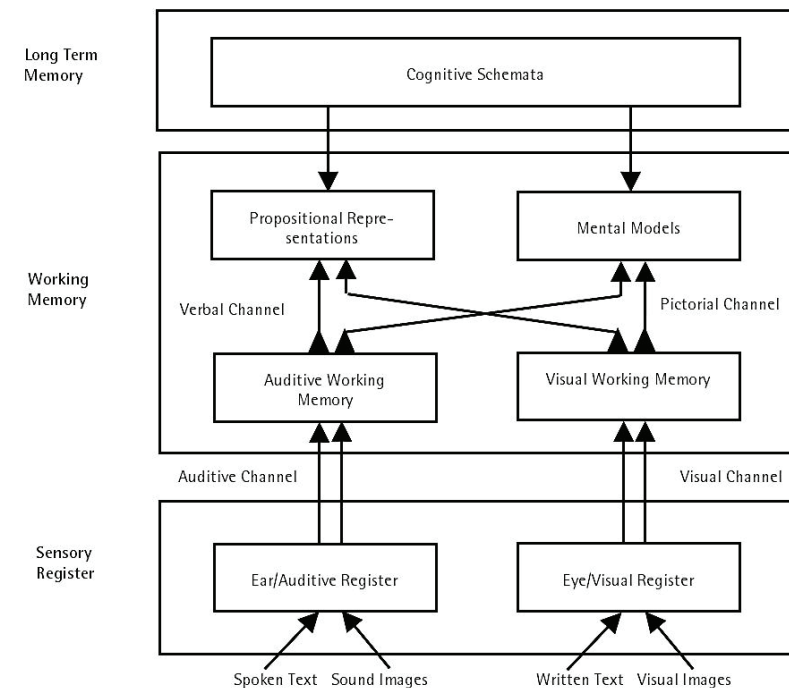
Beim Erfassen von Bildern entstehen zwei bzw. drei mentale Repräsentationsformen:

- 1) Repräsentation der Wahrnehmung (perceptual representation): Es wird eine interne Abbildung des Bildes angelegt.
- 2) Mentales Modell (mental model): Aus dem internen Bild wird ein mentales Modell gebildet, das die Informationen aus dem Bild zu einem Modell zusammenfügt, aus dem mentalen Modell können dann spezifische Informationen abgeleitet werden, die als Propositionen enkodiert werden.

Auch in der internen Repräsentation wird unterschieden zwischen deskriptiven und depiktiven Repräsentationen. Deskriptiv sind dabei die Oberflächenstruktur und die propositionale Repräsentation, da sie Symbole zur Beschreibung nutzen. Depiktiv ist das mentale Modell, da es eine ikonische Repräsentation ist, die eine analoge Struktur zum Sachverhalt hat.

Schnotz verbindet nun die Modelle und ordnet die Prozesse auf den Ebenen der sinnlichen Wahrnehmung, des Arbeitsgedächtnisses und des Langzeitgedächtnisses an (vgl. Abbildung 9). Das Arbeitsgedächtnis hat dabei eine begrenzte Aufnahmefähigkeit.

Abb. 9: Integriertes Modell des Text- und Bildverständnisses



Nach Schnotz (2005, Abb. 4.3).

Auf der Ebene der sinnlichen Wahrnehmung gibt es zwei Wahrnehmungskanäle. Gesprochener Text und Geräusche gelangen durch das Ohr über den auditiven Kanal in das auditive Arbeitsgedächtnis. Geschriebener Text und Bilder gelangen durch die Augen über den visuellen Kanal in das visuelle Arbeitsgedächtnis. Die Wahrnehmungskanäle haben eine begrenzte Aufnahme- und Verarbeitungskapazität.

Im Arbeitsgedächtnis findet die Verarbeitung in zwei weiteren Kanälen statt: dem verbalen und dem bildlichen Kanal. Gesprochene oder geschriebene Sprache wird im verbalen Kanal verarbeitet. Visuelle Bilder und Ton werden im bildlichen Kanal verarbeitet und in propositionale Strukturen oder mentale Modelle verarbeitet. Auch diese Kanäle haben eine beschränkte Aufnahme- und Verarbeitungskapazität.

Das Langzeitgedächtnis unterstützt den Prozess der Verarbeitung in propositionale Strukturen und mentale Modelle durch Vorwissen, das für die Verarbeitung benötigt wird. So sind für das Verstehen von geschriebenem Text die Kodierung der Symbole und die Syntaxstruktur notwendig. Das Verstehen von Bildern kann durch Vorwissen vereinfacht werden.

In dem Modell kann zwischen Perzeption und Kognition unterschieden werden. Die Perzeption findet beim Transfer von Information von der Außenwelt durch den auditiven und visuellen Wahrnehmungskanal in das Arbeitsgedächtnis statt. Die kognitive Informationsverarbeitung findet im Arbeitsgedächtnis mithilfe von Langzeitgedächtnis in den verbalen und bildlichen Kanälen statt.

Beim Lesen von Text wird die Information über das Auge aufgenommen und durch den visuellen Kanal in das visuelle Arbeitsgedächtnis weitergeleitet. Im visuellen Arbeitsgedächtnis befindet sich nun eine Repräsentation der Wahrnehmung. Die verbale Information wird nun herausgefiltert (schwarzes Dreieck in Abbildung 9) und über den verbalen Kanal in das Arbeitsgedächtnis für propositionale Repräsentationen geleitet, wo die Bildung oder Erweiterung von Propositionen und mentalen Modellen mithilfe des Langzeitgedächtnisses angestoßen wird.

Beim Hören eines Textes wird die Information über das Ohr aufgenommen und durch den auditiven Kanal in das auditive Arbeitsgedächtnis weitergeleitet. Im visuellen Arbeitsgedächtnis befindet sich nun eine Repräsentation der Wahrnehmung. Die verbale Information wird wieder herausgefiltert und weiter verarbeitet, analog zum Vorgang beim Lesen von Texten.

Beim Erfassen von Bildern wird die Information über das Auge aufgenommen und durch den visuellen Kanal in das visuelle Arbeitsgedächtnis weitergeleitet und dort repräsentiert. Die bildliche Information wird herausgefiltert und durch den bildlichen Kanal in das Arbeitsgedächtnis für mentale Modelle geleitet, wo sich ein mentales Modell mithilfe des Langzeitgedächtnisses bildet oder wo dieses erweitert wird. Aus dem mentalen Modell können dann wieder Propositionen abgeleitet werden.

Beim Erfassen von Ton wird die Information über das Ohr aufgenommen und durch den auditiven Kanal in das auditive Arbeitsgedächtnis weitergeleitet. Aus der nun hier befindlichen Repräsentation der Wahrnehmung wird die hör-bildliche Information herausgefiltert und durch den bildlichen Kanal in das Arbeitsgedächtnis für mentale Modelle geleitet, wo sich ein mentales Modell mithilfe des Langzeitgedächtnisses bildet oder wo dieses erweitert wird. Aus dem mentalen Modell können auch hier wieder Propositionen abgeleitet werden.

Sinnvolles Lernen aus Texten und Bildern ist folglich laut Schnotz die Gesamtleistung aus kognitiven Prozessen wie Selektieren und Organisieren der Information, Aktivierung von Vorwissen und die Integration von Wissen aus verschiedenen Quellen, die zu einer propositionalen Repräsentation und zu einem mentalen Modell führen. Beim Verständnis von visuellen Bildern beispielsweise werden relevante bildliche Informationen aus einer Zeichnung, Karte oder einem Diagramm als externe Quelle selektiert und organisiert. Vorwissen wird als interne Quelle genutzt, was zur Bildung eines mentalen Modells führt, mit dessen Hilfe neues Wissen in propositionale Repräsentation überführt werden kann.

Basierend auf dem ITPC-Modell trifft Schnotz mehrere Voraussagen, die sich mit den vorliegenden empirischen Studien decken sollen (vgl. Schnotz 2005, 62ff) und die Prinzipien des multimedialen Lernens entweder erweitern oder einschränken:

- Bilder verbunden mit gesprochenem Text führen zu einem besseren Lerneffekt, da der gesprochene Text über den auditiv-verbalen Kanal und das Bild über den visuellen-bildlichen Kanal gleichzeitig verarbeitet werden können (multimedia principle).
- Werden Bild und Text kombiniert, so sollten sie räumlich nah beieinander stehen (spatial contiguity principle).
- Es ist besser ein passendes Bild vor der Textlektüre als nach der Textlektüre zu präsentieren.
- Ungeübte Leser profitieren mehr von Illustrationen in geschriebenen Texten als geübte Leser.
- Das Konzept der dualen Kodierung geht davon aus, dass die Kombination von Bildern und Text immer zu besseren Lerneffekten führt. Nach dem ITPC-Modell ergeben sich aber ggf. Synchronisations- und Ablenkungsstörungen.
- Bilder verbunden mit gesprochenem und geschriebenem Text führen zu einem schlechteren Lerneffekt, da Leser den Text oft schneller lesen können als er gesprochen wird und sich dadurch Synchronisationsprobleme zwischen gelesenen und gesprochenem Text ergeben (specific redundancy effect).
- Redundanz durch mehrere Quellen wie Text und Bild führt bei Lernenden mit großem Vorwissen zu schlechten Lerneffekten, da beide Informationen über den visuellen Kanal aufgenommen werden müssen und die Aufmerksamkeit des Lernenden zwischen den beiden Quellen geteilt werden muss (general redundancy effect).

- Visualisierungen müssen für die Aufgabe geeignet sein. Nicht geeignete Visualisierungen führen bei Lernenden mit geringem Vorwissen zu Schwierigkeiten bei der Bildung eines mentalen Modells, da die Visualisierung nicht direkt in ein mentales Modell umgewandelt werden kann. Bei Lernenden mit Vorwissen und existentem mentalen Modell kann es zu Schwierigkeiten bei der Einordnung der Visualisierung in das existente Modell kommen.
- Da das ITPC-Modell unterschiedliche Wege für Text- und Bildverständnis hat, besteht die Möglichkeit, Text- durch Bildverständnis oder umgekehrt zu ersetzen. So können zur Bildung eines mentalen Modells nur Text oder nur Bild herangezogen werden.
- Bei der Integration von mehreren Lernquellen müssen Vorteile und Kosten immer gegeneinander abgewogen werden. Mehrere externe Quellen bedeuten eine erhöhte kognitive Last, welche die Vorteile mehrerer Quellen nicht immer ausgleichen können.

4.6 Gemeinsamkeiten und Unterschiede zwischen dem ITPC-Modell und dem CTML-Modell

Die Modelle von Schnotz (ITPC-Modell) und Mayer (CTML-Modell) sind fast deckungsgleich und unterscheiden sich nur in wenigen Aspekten. Beide Modelle basieren auf einer kognitiven Architektur, die aus mehreren Gedächtnisformen mit limitierten Kapazitäten und Kanälen für die Verarbeitung und Speicherung von Information nach dem Konzept der dualen Kodierung besteht. Die Modelle unterscheiden sich jedoch in folgenden zwei Aspekten. Im CTML-Modell wird ausgegangen von einem kombinierten auditiv-verbalen Kanal und einem visuell-bildlichen Kanal, der die Ebenen der Sinneswahrnehmung und Repräsentation im Arbeitsgedächtnis zusammenführt. Im ITPC-Modell dagegen sind die beiden Ebenen und die darin enthaltenen Kanäle getrennt. Auf der Ebene der sinnlichen Wahrnehmung existieren der auditive und der visuelle Kanal und auf Ebene des Arbeitsgedächtnisses der verbale und bildliche Kanal. Information kann beim Lesen von geschriebenem Text durch das Auge über den visuellen Kanal in das Arbeitsgedächtnis gelangen und dann weiter im verbalen Kanal verarbeitet werden. So sind alle Kombinationen von Kanälen möglich, was die Verarbeitung von Information als gesprochener und geschriebener Text, von Bildern und Ton möglich macht. Zudem geht Mayer von der Bildung getrennter mentaler Modellen für verbale und bildliche Information aus („verbal model“ und „pictorial model“), die mithilfe des Langzeitgedächtnisses im Arbeitsgedächtnis integriert werden. Das ITPC-Modell unterscheidet nicht explizit zwischen zwei getrennten mentalen Modellen.

Das Prinzip der dualen Kodierung nach Mayer macht die Annahme, dass der Lerneffekt höher bei einer integrierten Darstellung von Text und Bild ist als nur bei Text. Visualisierungen unterstützen also den Lerneffekt, da Text und Bild über zwei getrennte Kanäle verarbeitet werden. Beide Modelle sagen voraus,

dass Visualisierungen dazu genutzt werden, um ein mentales Modell beim Nutzer zu bilden. Basierend auf der Annahme, dass Text- durch Bildverständnis ersetzt und zur Bildung eines mentalen Modells nur Text oder nur Bild herangezogen werden kann, können Visualisierungen für die Bildung eines mentalen Modells herangezogen werden.

Schnotz grenzt im ITPC-Modell die Umstände, unter denen diese Vorteile der dualen Kodierung wirken, ein. Ungeübte Leser profitieren mehr von Illustrationen in geschriebenen Texten als geübte Leser. Weiterhin müssen Visualisierungen für die Aufgabe geeignet sein. Es kann zu Schwierigkeiten bei der Bildung eines mentalen Modells oder bei der Einordnung in ein bestehendes mentales Modell kommen, wenn Visualisierungen nicht aufgaben- und nutzeradäquat eingesetzt werden.

4.7 Kognitive Verarbeitung von Visualisierungen

In den bisherigen Modellen wurde die kognitive Informationsverarbeitung von verschiedenen Modalitäten wie Text und Bild unter Berücksichtigung von Gedächtnismodellen und -limitierungen besprochen und erläutert, welche Effekte sich auf den Lerneffekt ergeben. In den folgenden Abschnitten wird tiefer darauf eingegangen, wie Visualisierungen von einfachen Diagrammen bis zu komplexen Visualisierungen kognitiv verarbeitet werden und welche Effekte sich daraus ergeben.

Bertin (1983) unterscheidet drei Hauptprozesse bei der Verarbeitung von Diagrammen und Visualisierungen:

- 1) Die Identifikation und Enkodierung der visuellen Elemente einer Grafik und die Zuweisung zu konzeptionellen oder realen Referenten über die der Graph Informationen übermittelt (external identification).
- 2) Die Identifikation der möglichen Varianten, welche die visuellen Elemente einnehmen können und zu welcher konzeptuellen Variable oder Skalierung sie gehören (internal identification).
- 3) Aus der Anordnung der visuellen Elemente kann nun interpretiert werden, welche Bedeutung sie für die konzeptuellen Variablen haben (perception of correspondence).

Cleveland und McGill (1984, 1986) untersuchen die psycho-physikalischen Aspekte der grafischen Wahrnehmung. Sie untersuchen die Frage, wie Informationen, die in Graphen enkodiert angezeigt werden, im menschlichen Gehirn wieder dekodiert werden. Dabei unterscheiden sie zwischen grundlegenden Operationen für die Extraktion von quantitativen Informationen aus Graphen und messen ihre Genauigkeit. Die elementaren Wahrnehmungsleistungen („elementary perceptual tasks“), geordnet nach Genauigkeit der menschlichen Wahrnehmung sind Folgende (Cleveland und McGill 1984, 536):

- 1) Position auf einer gemeinsamen Skala
- 2) Position auf mehreren identischen Skalen bei unterschiedlicher Anordnung

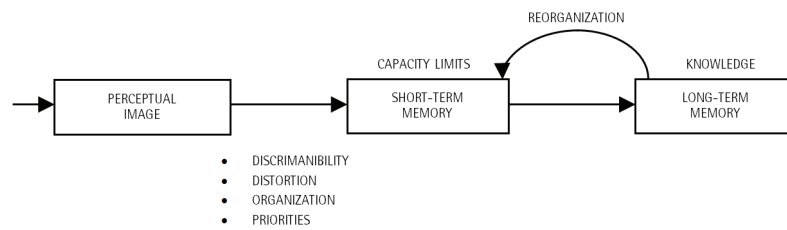
- 3) Länge, Richtung, Winkel
- 4) Flächengröße
- 5) Volumen, Krümmung
- 6) Schattierung, Farbton

Dabei werden die Wahrnehmungsleistungen kombiniert für das Lesen von verschiedenen Diagrammtypen angewandt. In Linien- oder Streudiagrammen wird zum Beispiel der erste Task „Position auf einer gemeinsamen Skala“ genutzt. Kombinierte Faktoren wie Winkel, Flächengröße oder Bogenlängen werden in Kreisdiagrammen eingesetzt, Farbschattierungen zum Beispiel in einer statistischen Karte. Im ersten Experiment (Cleveland und McGill 1984) zeigen sie, dass die Positionsbeurteilung genauer ist als die Längenbeurteilung bei Faktoren zwischen 1,4 und 2,5 und dass die Positionsbeurteilung um den Faktor zwei genauer ist als die Winkelbeurteilung. In einem weiteren Experiment (Cleveland und McGill 1986) mit einfachen Diagrammen ordnen sie die Genauigkeit der Wahrnehmungsleistungen nach 1. Position, 2. Länge, 3. Winkel und Neigung und 4. Flächengröße.

Kosslyn (1989) präsentiert das Standardmodell der visuellen Informationsverarbeitung (vgl. Spoehr und Lehmkuhle 1982 zitiert in Kosslyn 1989). Hier werden drei Stufen der visuellen Informationsverarbeitung unterschieden (vgl. Abbildung 10). Auf erster Stufe steht das wahrgenommene Bild (*perceptual image*), das sensorische Aspekte wie Kanten und Regionen enthält, solange man auf das Display schaut. Im Kurzzeitgedächtnis (*short-term memory*) erfolgt die Zwischenspeicherung des wahrgenommenen Bildes. Die Kapazität ist limitiert in Dauer und Anzahl der Informationen, die vorgehalten werden können. Zur Interpretation des Bildes wird vorher gelerntes Wissen aus dem Langzeitgedächtnis (*long-term memory*) abgerufen und im Kurzzeitgedächtnis findet die Reorganisation und Reinterpretation der Information statt. Im Langzeitgedächtnis wird eine große Anzahl an Informationen für eine lange Zeit gespeichert. Verschiedene Eigenschaften beeinflussen das Ablesen von Graphen, Diagrammen und anderer visueller Reize: (1) *Unterscheidbarkeit* (*Discriminability*): Ist der Stimulus zu klein oder hebt sich nicht genug ab, so wird er übersehen. (2) *Verzerrung* (*Distortion*): Verzerrung bei der Wahrnehmung von Eigenschaften von Objekten wie Größe und andere. (3) *Anordnung* (*Organization*): Stimuli werden bei der Wahrnehmung automatisch in Gruppen und Einheiten organisiert (*Gestalt-Gesetze*). (4) *Priorität* (*Priorities*): Manchen Eigenschaften von Stimuli wird mehr Beachtung geschenkt, zum Beispiel: Hervorhebung, starke Farbgebung etc.

Die Kapazitätslimitierung des Kurzzeitgedächtnisses limitiert das Vorhalten und die Verarbeitung von Informationen. Information im Kurzzeitgedächtnis wird im Langzeitgedächtnis enkodiert und mit bestehendem Wissen abgeglichen und interpretiert. Auf diese Weise können zum Beispiel auch abfallende Linien in einem Liniendiagramm als Abwärtstrend identifiziert werden.

Abb. 10: Drei Stufen des visuellen Informationsverarbeitungsprozesses



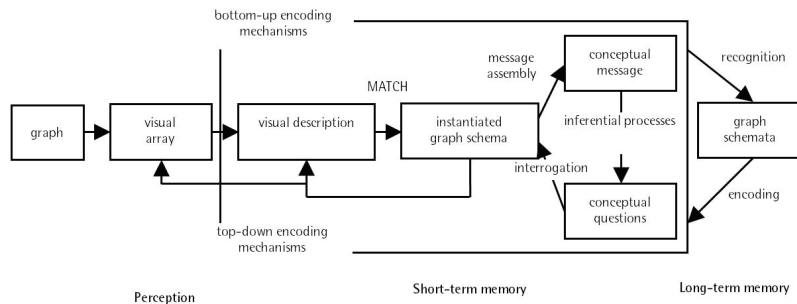
Nach Kosslyn (1989, Abb. 2).

Darauf baut Pinkers Modell der Graphenwahrnehmung (Pinker 1990) auf (vgl. Abbildung 11). Der Graph wird im visuellen Array (visual array) als weitgehend unverarbeitete, visuell-bildliche Repräsentation gespeichert. Mithilfe von visuellen Enkodierungsprozessen wird die Repräsentation in eine symbolische Repräsentation (symbolic representation) oder strukturierte Beschreibung (structural description) der Szene umgewandelt. In der Beschreibung werden die einzelnen Elemente der Szene (z. B. Legende, Achsen, Linien usw.), ihre Eigenschaften (Größe, Position, Form, Farbe, Textur usw.) und ihre räumliche Zuordnung untereinander in separate Symbole kodiert. Übergeordnete Kognitionsprozesse müssen dann nur auf die Symbole zurückgreifen, die für den aktuellen Verarbeitungsschritt benötigt werden. Bei der kognitiven Verarbeitung existieren verschiedene Prozesse:

- Der *Match-Prozess* erkennt individuelle Graphen als zugehörig zu einem bestimmten Typ, welche im Langzeitgedächtnis gespeichert sind und instanziiert ein Graphenschema im Kurzzeitgedächtnis. Dabei werden Parameter im Schema durch Konstanten in der visuellen Beschreibung ersetzt. Das instanziierte Graphenschema ist dabei durch die Kapazitätsgrenzen des Kurzzeitgedächtnisses beschränkt.
- Beim Prozess der Informationsenkodierung (message assembly) wird grafische Information in konzeptuelle Information übersetzt. Diese kann (je nach Komplexität des Graphen) als Liste von n-Tupeln gespeichert werden (z. B. $\{V_1 \text{ absolute-value} = \text{January}, V_2 \text{ absolute-value} = \$20/\text{oz}\}$ für die Information, dass in einem Balkendiagramm der Wert \$20/oz für Januar abzulesen ist).
- Interrogation: Konzeptuelle Fragen werden wiederum top-down aufgelöst. Dafür wird ein Parameter als Variable gesetzt: $\{V_1 \text{ absolute-value} = \text{March}, V_2 \text{ absolute-value} = ?\}$; V_1 in das entsprechende visuelle Element enkodiert und überprüft, ob es in der visuellen Beschreibung existiert.
- Konzeptuelle Information: Informationen, die aus dem Graphen abgeleitet werden.
- Konzeptuelle Fragen: Informationen, die aus dem Graphen abgeleitet werden *sollen*.

- Inferenzprozesse: Inferenzprozesse können aus den konzeptuellen Nachrichten mithilfe von mathematischen und logischen Inferenzregeln weitere konzeptuelle Fragen bilden, die nicht direkt aus dem Graph abgelesen werden können.

Abb. 11: Pinkers Modell der Graphenwahrnehmung



Nach Lohse (1991a, Abb. 3.1); adaptiert von Pinker (1990, Abb. 4.19).

Lohse (1993) entwickelt ein Computerprogramm namens UCIE (Understanding Cognitive Information Engineering), das die grafische Informationsverarbeitung simuliert. Einzelne Komponenten des Modells sind dabei (1) Sequenzen von Augenfixierungen, (2) Limitierungen in der Kapazität und Vorhaltezeit im Kurzzeitgedächtnis, sowie (3) der Schwierigkeitsgrad der Informationsverarbeitung in jedem Schritt. Dabei wird auch Pinkers Modell des instanziierten Graph-Schemas genutzt. UCIE (a) bildet den Graph, (b) analysiert die Anfrage, (c) modelliert den menschlichen Informationsverarbeitungsprozess und (d) weist jedem Schritt Zeitspannen zur Verarbeitung zu. Durch die Bildung der Summe der Zeitspannen, die für jeden Teilschritt benötigt wird, prognostiziert UCIE die Gesamtzeit des Verarbeitungsprozesses.

In einer Nutzerstudie stellen Carpenter und Shah (1998) fest, dass Nutzer mit der Augenfixierung zwischen verschiedenen Bereichen im Graphen wie visuellen Elementen, Achsen, der Legende, Titel usw. wechseln. Für jedes visuelle *chunk* bewegt sich die Augenfixierung zwischen diesen Bereichen, was folgern lässt, dass Nutzer zwischen verschiedenen Stufen der Verarbeitung iterieren. Wenn die Komplexität des Graphs erhöht wird (Erhöhung der chunks), erhöht sich auch die Bewegung zwischen den Regionen des Graphs, d.h. für jeden chunk wird eine Verarbeitungsiteration benötigt. Mehrere Zyklen der Verarbeitung werden benötigt, um Informationen zu integrieren.

Ratwani et al. (2008) führen mehrere Experimente durch, um zu zeigen, dass die Integration während der kognitiven Verarbeitung von Graphen zwei Komponenten hat: (1) Die visuelle Integration basiert auf wahrgenommenen Merkmalen (Farbe, Nähe, semantische Nähe usw.). Die individuellen Datenpunkte werden visuell integriert, um Cluster höherer Ordnung zu bilden. (2)

Die kognitive Integration macht es möglich, visuelle Cluster direkt zu vergleichen, was einen neuen Prozess darstellt. Wenn die Komplexität des Graphs erhöht wird, werden mehr Zyklen von visueller Clusterbildung und deren Vergleichen benötigt. Je komplexer ein Graph ist, umso mehr wird integriert, was einen zyklischen Prozess darstellt.

Auch Ware (2005) unterstützt die Annahme, dass nicht ein komplettes Modell der visuellen Beschreibung im Gedächtnis vorgehalten wird, sondern dass jede konzeptuelle Frage im Sinne Pinkers wieder in der Visualisierung abgelesen wird. Dies beruht auf Untersuchungen, dass auch das visuelle Arbeitsgedächtnis starken Limitierungen unterworfen ist und beispielsweise nur drei farbige Objekte gleichzeitig speichern kann (Vogel, Woodman und Luck 2001, zitiert in Ware 2005).

In einer Nutzerstudie zeigen Trafton et al. (2000), wie Wetterexperten Informationen aus komplexen Visualisierungen extrahieren, Informationen in ein qualitatives Modell integrieren und das mentale Modell nutzen, um quantitative Vorhersagen zu machen. Die bisherigen Modelle beschreiben, wie einfache Visualisierungen mit wenigen Variablen direkt abgelesen werden können. In realen komplexen Visualisierungen dagegen wird (1) oft Domänenwissen benötigt, (2) muss entschieden werden, welche Information angezeigt werden, (3) welche Information extrahiert werden und (4) wie das Wissen weiterverarbeitet wird. Die extrahierten Informationen werden dann von den Wissenschaftlern für das Ableiten, Generalisieren, Extrapolieren und Vorhersagen genutzt. Die komplexen wissenschaftlichen Visualisierungen enthalten viele Variablen mit unterschiedlichen Skalierungen und Verläufen über die Zeit. Das Informationsverhalten wurde mit einer Kognitiven Task Analyse (CTA) untersucht: Die Wissenschaftler haben (1) sich zuerst einen groben Überblick über Wetterbedingungen gemacht, (2) ein qualitatives mentales Modell (QMM) gebildet, (3) das Modell verifiziert und angepasst und (4) abschließend eine Zusammenfassung der Vorhersage geschrieben. Trafton et al. schlagen vor, dass im konsistenten QMM Informationen aus verschiedenen Quellen und Visualisierungen integriert werden. Dabei werden Informationen verschiedener Modalitäten und Arten integriert. Wissenschaftler haben Informationen aus Diagrammen, Graphen, Bildern und Text extrahiert, in ein mentales Modell integriert und daraus Vorhersagen abgeleitet. Das QMM kann genutzt werden, um quantitative Informationen, qualitative Informationen und What-if-Szenarien zu bilden und Vorhersagen zu machen. Im Gesamtprozess verbraucht die Bildung des QMM die meiste Zeit.

In weiteren Studien (2002, 2001) zeigen Trafton et al., dass Wissenschaftler Information nicht nur ablesen, sondern auch die mentale Vorstellungskraft und räumliche Transformation („spatial transformation“) nutzen, um Information, die nicht in der Visualisierung zu finden ist, zu extrapolieren. Räumliche Transformationen sind kognitive Operationen, die ein Wissenschaftler auf einer Visualisierung anwenden kann. Beispiele sind: mentale Rotation, Erzeugung

eines mentalen Bildes, Modifizieren des mentalen Bildes durch das Hinzufügen oder Löschen von Merkmalen, Animation eines Aspekts der Visualisierung, die Vorhersage des Fortgangs einer Zeitreihe, das mentale Verschieben eines Objekts, die Transformation einer 2D- in eine 3D-Ansicht oder umgekehrt und alles andere, was mental durchgeführt werden kann, um ein Problem besser zu verstehen oder zu vereinfachen. Diese Operationen können dabei entweder intern mental ausgeführt werden, aber auch extern (z. B. in einem IV-System). In dem Experiment mit Wissenschaftlern der Astronomie und Physik, die Daten mit Visualisierungen auswerten, wird gezeigt, dass fast genauso viel Information direkt abgelesen wird, wie Information extrapoliert wird, die nicht explizit in der Visualisierung vorhanden ist.

Liu et al. (2008) schlagen *Distributed Cognition* (DC) (vgl. Hutchins 1980, zitiert in Liu u.a. 2008) als theoretisches Framework für die Informationsvisualisierung, insbesondere für die Komponenten *Repräsentation* und *Interaktion* vor. Die traditionelle Ansicht ist, dass Kognition die Informationsverarbeitung im menschlichen Gehirn ist. Der Fokus der Forschung liegt deshalb darauf, wie Informationen im Gehirn aufgenommen und verarbeitet werden. Dabei wird der Kontext, in dem der Nutzer interagiert, außer Acht gelassen. In DC ist Kognition mehr eine Eigenschaft von Interaktion als eine Eigenschaft des menschlichen Verstands. Tools (zum Beispiel in der Informationsvisualisierung) erlauben es, schwierige Aufgaben in solche zu verwandeln, die mit einfacher, menschlicher (kognitiver) Mustererkennung gelöst werden können. In DC besteht das kognitive System aus Individuen und Artefakten, die sie benutzen: In der Informationsvisualisierung aus Personen, die zusammenarbeiten, den Informationssystemen, Hilfsmitteln wie Maus, Computer, Stift und Block usw. Die beobachtbaren Repräsentationen von Information befinden sich damit auch außerhalb des menschlichen Gehirns, aber innerhalb des kognitiven Systems. Auch bei individueller Analyse mit IV-Systemen existieren enge Kopplungen zwischen internen Repräsentationen wie Schemata, mentalen Modellen, Propositionen, lexikalen Modellen und mentalen Bildern sowie externen Repräsentationen in der Visualisierung. Wie interne und externe Repräsentationen genau zusammenhängen, bleibt dabei ein offenes Problem in DC. Externe Repräsentationen sind mehr als nur Stimuli, mehr als nur Inputs. In DC ist Interaktion kein Einweg-Prozess von Systemen zu Menschen, sondern Koordination zwischen interner und externer Repräsentationen. Dabei werden Visualisierungen nicht einfach von einem Anfangszustand in einen Endzustand mit Interaktionstechniken gebracht. Sondern es besteht ein andauernder Austausch zwischen intern und extern, wobei externe Repräsentationen auch als Denkhilfen gebraucht werden (zum Beispiel wie bei dem Computerspiel „Tetris“, wo Bausteine als Denkhilfe gedreht werden). Interaktion ist die Propagierung von Repräsentations-Zuständen („representation states“) in einem kognitiven System durch Koordination. Interne und externe Repräsentationen interagieren reziprok und kreieren auch Erkenntnis und konzeptionelle Veränderung.

4.8 Zusammenfassung

Dieser Abschnitt stellt die grundlegenden menschlichen Mechanismen der Informationsverarbeitung, insbesondere die kognitive Verarbeitung von Visualisierungen vor. Die menschliche Kognition basiert auf dem Aufbau des Gedächtnisses, auf den Theorien der dualen Kodierung und der kognitiven Last. Auf dieser Basis bilden sich zwei leicht unterschiedliche Modelle des multimedialen Lernens, welche die kognitive Verarbeitung von unterschiedlichen Modalitäten wie Text und Bild modellieren und beschreiben, wie die Informationen in einem mentalen Modell enkodiert werden.

Darauf aufbauend wird näher auf kognitive Prozesse für die Verarbeitung von Visualisierungen eingegangen. Die Modelle zeigen die Verarbeitung von Visualisierungen über die verschiedenen Gedächtnisstufen mit der Abrufung von bereits gelernten Schemata aus dem Langzeitgedächtnis. In den Experimenten von Trafton et al. wird gezeigt, dass Wissenschaftler Informationen verschiedener Modalitäten und verschiedener Darstellungsarten wie Diagramme, Graphen, Bilder und Text in einem qualitativen mentalen Modell integrieren.

Die Bildung des mentalen Modells benötigt dabei den längsten Zeitraum im Gesamtprozess. Das mentale Modell kann dann genutzt werden, um Informationen abzuleiten, zu generalisieren, zu extrapolieren und um Vorhersagen zu machen. Räumliche Transformation wird genutzt, um Information, die nicht explizit in der Visualisierung abzulesen ist, zu extrapolieren (Trafton u.a. 2002). Liu et al. betont die enge Kopplung zwischen externer Repräsentation der Information im IV-System und interner Repräsentation im mentalen Modell. Es besteht ein andauernder Austausch, wobei externe Repräsentationen auch als Denkhilfen gebraucht werden.

5. Diskussion und Schlussfolgerungen

Eines der Hauptziele dieses Artikels ist es, zu untersuchen, ob interaktive Visualisierungen in einen allgemeinen Suchprozess im Web integriert werden können. Grundlegende Eigenschaften des Informationsprozesses wie das Informationsbedürfnis, die Vagheit der Anfrage, die Unsicherheit und Iterativität und Nutzung verschiedener Suchtechniken im Retrieval-Prozess werden dabei im Bereich Information Retrieval (IR) definiert. Die Nutzung von interaktiven Visualisierungen zur Verstärkung des Kognitionsprozesses wird im Bereich Information Visualization (IV) propagiert. Die folgenden Definitionen zeigen noch einmal Gemeinsamkeiten und Unterschiede auf abstraktem Niveau.

Eine aktuelle Zielsetzung von Information Retrieval bieten Ingwersen und Järvelin (2007): „that the ultimate goal of information retrieval is to facilitate human access to and *interaction* with information that probably may entail cognition.“ Card, Mackinlay und Shneiderman (1999) definieren Information

Visualization als: „The use of computer-supported, interactive, visual representations of abstract data to amplify cognition“.

Sicherlich geht es in beiden Definitionen um den Informationsverarbeitungsprozess mit dem Ziel, Informationen oder Daten zu finden und daraus Erkenntnisse zu ziehen. In beiden Definitionen geht es um Interaktion und den Kognitionsprozess. Unterschiedlich ist, dass IR allgemein von Informationen spricht und IV von „visuellen Repräsentationen von abstrakten Daten“. Auf diesen Unterschied wird später in diesem Abschnitt noch eingegangen. IV nutzt computerunterstützte, interaktive Visualisierungssysteme, um den Kognitionsprozess zu verstärken. Die Definition von IR spricht nur allgemein von Interaktion mit Informationen.

Die Ausgangslage für die Informationssuche ist komplex. Die Informationsmenge im Web nimmt jedes Jahr kontinuierlich zu (Hilbert und López 2011) und die Informationen unterscheiden sich stark voneinander. Abschnitt 2.1.4 zeigt die Heterogenität von Informationen im Web, die sich in Eigenschaften wie unterschiedlicher Modalität, Informationsart, Strukturiertheit, Granularität, Qualität und ihrer Verteiltheit zeigt. Zum Beispiel liegen numerische Faktendaten wie historische Aktienkurse in Datenbanken, während textuelle Informationen zu passenden historischen Ereignissen in der Wikipedia enthalten sind. Diese Informationsarten sind heterogen, können aber verbunden und explorierbar gemacht werden, so dass der Nutzer die Verbindung nutzen kann, um Erkenntnisse zu generieren.

Abschnitt 2.1 zeigt Modelle auf verschiedenen Abstraktionsstufen, die für die Modellierung des Suchprozesses im Bereich der *Informationssuche* genutzt werden. Dabei werden die Modelle immer abstrakter und ihre Messbarkeit im Sinne einer quantitativen Analyse der Güte eines Suchergebnisses oder Suchsystems nimmt immer weiter ab.

Das klassische IR-Modell ist auch heute noch das Standardmodell für Suchsysteme wie Fachinformationssysteme oder Google als klassische Websuchmaschine. Die Güte dieser Suchsysteme wird mit IR-Evaluationen mit standardisierten Messgrößen wie Precision und Recall gemessen. Diese quantitative Evaluation einheitlicher Messgrößen erlaubt die Vergleichbarkeit der Güte verschiedener Suchsysteme im Sinne der Qualität des Suchergebnisses; andere Merkmale eines Suchprozesses wie zum Beispiel die Eigenschaften des Nutzerinteraktionsprozess, der Iterativität des Suchprozesses oder die Vagheit im Suchprozess werden mit anderen Evaluationsarten untersucht.

Für die Messung von IR-Modellen werden Testkollektionen wie TREC (Text Retrieval Conference), CLEF (Cross-Language Evaluation Forum) oder NTCIR (NII-NACSIS Test Collection for IR Systems) genutzt. Die TREC-Initiative (Voorhees und Harman 2005) beispielsweise wurde Anfang der 1990er Jahre gegründet und ist ein jährliches Experiment, an dem mehrere Forschungsgruppen teilnehmen. Dabei werden von TREC Testkollektionen gebildet, die für die Evaluationen verschiedener Retrievalsysteme eingesetzt

werden können. TREC-Testkollektionen bestehen aus (1) Dokumentenkollektionen, (2) einer Anzahl von Topics als Ausdruck von Informationsbedürfnissen und (3) Relevanz-Einschätzung, die eine Liste von Dokumenten beinhaltet, die für ein Topic gefunden werden sollten. Dabei werden mehrere simplifizierte Annahmen getätigt (vgl. Voorhees 2002): (a) Relevanz wird approximiert durch thematische Ähnlichkeit zwischen Topic und Dokument. Laut Voorhees beinhaltet das die Annahmen, dass alle Dokumente gleich erstrebenswert sind, dass die Relevanz von Dokumenten untereinander unabhängig ist und dass das Informationsbedürfnis statisch ist, (b) dass ein einzelner Satz von Relevanz-Einschätzungen für ein Topic repräsentativ sei und (c) dass die Liste von relevanten Dokumenten für ein Topic komplett sei. Dabei werden in den meisten Testkollektionen Relevanz-Entscheidungen als binäre Entscheidungen getroffen. Für TREC werden 25 Topics als Minimum und 50 Topics als Norm angenommen. Als finale Kenngröße wird die Mean Average Precision (MAP) genutzt, die auf relativ stabile Weise Werte als Abbildung von Precision und Recall für Testkollektionen liefert. Dabei betont Voorhees (2002), dass nur die MAP für Systeme verglichen werden kann, die exakt dieselbe Testkollektion nutzen. Verändert sich beispielsweise die Auswahl der Topics, sind die Ergebnisse nicht mehr vergleichbar. Der Prozess des Poolings bildet ein Subset aus der sehr großen Kollektion von TREC (durchschnittlich 800.000 Dokumente), damit es überhaupt möglich, wird einen Pool in überschaubarer Größe durch Juroren beurteilen zu lassen. Voorhees (2002) zeigt, dass die Annahmen des Poolings und der Relevanzbeurteilungen aus statistischer Sicht stabil sind, bis auf Änderungen in den Relevanz-Beurteilungen. Werden verschiedene Retrieval-Systeme mit exakt derselben TREC-Kollektion getestet, lässt sich die Performance der Systeme damit direkt vergleichen.

Allerdings stellen Armstrong et al. (2009) fest, dass sich die Adhoc-Retrieval-Ergebnisse in den Jahren von 1998 bis 2008 nicht messbar verbessert haben. In einer Längsschnittanalyse verglichen sie die Ergebnisse, die mit TREC Adhoc-, Web-, Terabyte- und Robust-Kollektionen erreicht und in Konferenz-Papieren von SIGIR (1998-2008) und CIKM (2004-2008) veröffentlicht wurden. Die Autoren stellen fest, dass sich die MAP der Retrieval-Systeme in den Folgejahren nicht messbar zu den Original-TREC-Systemen von 1999 unterscheiden und dass es zudem keinen Aufwärtstrend der MAP gab. In den Forschungspapieren berichtete signifikante Verbesserungen der MAP beruhen dabei meist auf zu schwachen ausgewählten Baselines. Zudem ergibt sich eine Diskrepanz zwischen der Optimierung des Gesamtsystems zur Erhöhung der MAP und der Optimierung von Einzelkomponenten. Da die Retrievalsysteme von den Forschungsgruppen selbst entwickelt werden, auf verschiedenen Suchmodellen und -funktionen beruhen, lässt sich aus den Gesamtergebnissen nicht erkennen, inwiefern sie zum Gesamtergebnis beitragen. Die Einzelkomponenten können additiv, aber auch negativ in Verbindung mit anderen Einzelkomponenten agieren. Oft ist es eine Sammlung von Optimierungsmaßnahmen,

die zum Gesamtergebnis führt. Da die Systeme und Komponenten nicht veröffentlicht werden, können die Komponenten nicht in Systemen anderer Gruppen weitergenutzt werden. Als einen Lösungsansatz schlagen die Autoren vor, Ergebnisse von Retrieval-Experimenten in einer öffentlichen Datenbank zu speichern und damit jederzeit einen nachvollziehbaren Stand der Forschung und nachvollziehbare Baselines zur Verfügung zu haben.

Analog zu fehlender Innovation auf messbarer quantitativer Evaluations-Ebene zeigen Ingwersen und Järvelin (2007) die Grenzen des klassischen Cranfield-Paradigmas auf Modell-Ebene auf. Das klassische IR-Modell und die dazugehörigen Retrieval-Experimente im Labor sind möglichst kontextfrei, echte Nutzer und Aufgaben werden ausgeblendet. Das Modell wird nur genutzt, um Algorithmen für das Retrieval und die Relevanz von Dokumenten zu vergleichen. In ihrer Modellerweiterung wird das klassische Modell um mehrere Ebenen erweitert: (1) den Suchkontext, (2) den Arbeitskontext und (3) den sozio-organisatorischen und kulturellen Kontext. Für diese Ebenen gelten andere Evaluationsmaße als für den klassischen IR-Kontext. Im Suchkontext muss die Usability und die Qualität der Information/Prozesse gemessen werden. In den anderen Ebenen muss die Qualität der Arbeits- und Informationsprozesse, Resultate und die sozio-kognitive Relevanz gemessen werden. Hierfür werden natürlich ganz andere Evaluationsmethoden benötigt. Wie diese verschiedenen Aspekte evaluiert werden können, bleibt letztendlich offen. Doch auch ohne diese Präzisierung bleibt das Modell beschränkt auf die IR-Suche (vs. Browsing) und Matchingprozesse. Alternative Suchmethoden und Informationsarten werden ausgeblendet.

Eine Stufe abstrakter ist das Berrypicking-Modell, das versucht, sich vom klassischen IR-Modell abzusetzen und aufzuzeigen, dass über verschiedene heterogene Quellen iteriert wird, die Anfrage auf Basis der Informationen immer weiter angepasst wird und die Suchtechniken wechseln. Man kann es als Erweiterung des klassischen IR-Modells sehen. Aufgrund der Abstraktionsstufe des Modells greifen Standard-IR-Evaluationen hier nicht mehr. Suchsysteme oder Prozesse, die durch das System abgebildet werden, müssen für jeden Task und für jede Domäne einzeln evaluiert werden und können nur Hinweise geben, dass das Verfahren für diesen Task und für diese Domäne funktioniert. Die Vergleichbarkeit der Systeme ist damit nicht mehr gegeben.

Exploratory Search (ES) geht noch eine Stufe weiter und zeigt die starke Heterogenität von Informationen im Web und die Lern- und Untersuchungsschritte, die nötig sind, um diese Informationen zu verarbeiten. Der Suchprozess muss erweitert werden, um den Suchprozess realitätsnah abbilden zu können. Dabei besteht die Heterogenität nicht nur in den Informationen selbst, sondern setzt sich auch noch im Kognitionsprozess des Nutzers fort, der beispielsweise Text- und Bildinformationen in unterschiedlichen kognitiven Prozessen verarbeitet (siehe Abschnitt 4.3). Der Nutzer muss sehr aufwendige, kognitive Prozesse durchführen, um diese Informationen zu verarbeiten und in

einem qualitativen mentalen Modell zu integrieren (siehe Abschnitt Informationsverarbeitung 4). Marchionini (2006) führt hier Prozesse auf wie „scannen/betrachten, vergleichen und qualitativ beurteilen“. Diese Prozesse sind zeitaufwendig, was explizit auch von Trafton u. a. (2000) festgestellt wird. Auch wenn die Prozesse zeitaufwendig sind, so ergeben sich doch Mehrwerte für den Nutzer, nämlich dass laut Marchionini der, „des Erwerbs von Wissen, Verstehen von Konzepten oder Fähigkeiten, der Auslegung von Ideen und Vergleichen oder Zusammenfassungen von Daten und Konzepten“. Verschiedene, heterogene Informationen werden also in verschiedenen, kognitiven Prozessen untersucht und verglichen, um daraus einen Mehrwert wie einen Lernprozess oder Wissen zu generieren und ein Informationsbedürfnis zu stillen. Die Brücke zur Informationsvisualisierung wird durch Marchionini selbst geschlagen, indem er hochinteraktive Systeme als Lösung des Problems und Standard-IV-Techniken wie Dynamic Queries oder Brushing vorschlägt. Ein weiterer Vorschlag ist der Hyperlinking-Mechanismus, der als neuer glyphenbasierter Prozess für IV in dieser Arbeit eingeführt wird. Exploratory Search ist damit als IR-Modell auch im gebildeten Modell in diesem HSR-Focus (Hienert 2014) integriert. Auf unterster Ebene stehen stark heterogene Informationen, welche als eine Möglichkeit in interaktiven (koordinierten) Visualisierungen abgebildet werden können. Auf Interaktionsebene stehen die geforderten Interaktionsmöglichkeiten, nicht nur auf individueller Ebene, sondern auch für mehrere Visualisierungen. Die Iterativität des Suchprozesses wird gegeben durch die Eingliederung in den Hyperlinking-Mechanismus über glyphenbasierte IR-Techniken oder Textlinks. Zudem besteht noch weiterhin die Möglichkeit, innerhalb des in der Visualisierung angezeigten Datensatzes zu suchen, zu filtern und zu browsen. Es existieren hier also immer zwei Wege: interaktive Methoden für die Suche/Filterung nach Daten innerhalb der Visualisierung und der Weg nach außen in das Web zu verwandten Informationen.

Ergänzend zeigt Mandl mit einem Qualitäts-Modell für Webinformation, dass für das Ranking und die Relevanzbewertung von Suchergebnissen im Web-Retrieval nicht nur Autoritätsmerkmale wie PageRank eine Rolle spielen, sondern auch weitere Komponenten wie zeitliche Aspekte, Gebrauchstauglichkeit, wirtschaftliche Aspekte, technische und Software-Qualität sowie interkulturelle Unterschiede. Dabei konnte gezeigt werden, dass die Qualität von Informationen einen messbaren Einfluss auf die Relevanz für den Nutzer hat. Dabei stellt die Gebrauchstauglichkeit ein Qualitätskriterium dar, das sich auf die Inhalte als auch auf die Informationssysteme bezieht („Die Gebrauchstauglichkeit stellt ein entscheidendes Qualitätskriterium dar, das sich im Internet kaum von den Inhalten trennen lässt.“ (Mandl 2006b, 15). Unterstützend zeigt Eibl, dass eine ästhetische und interaktive Gestaltung der Benutzungsoberfläche einen messbaren Einfluss auf Precision und Recall haben kann (vgl. Abschnitt 3.10.4).

Visualisierungen im Sinne dieser Arbeit sind Informationssysteme für die Darstellung und Präsentation von Informationen, aber auch für die hochinteraktive Suche nach Informationen mit verschiedenen Techniken wie Suche, Browsing oder Filterung. Mandl dagegen ordnet Visualisierungen noch als Mehrwertkomponenten für die Suche ein (Mandl 2006b, 62). Visualisierungen können für die Darstellung des Suchergebnisses genutzt werden, um die Beziehungen zwischen Dokumenten zum Beispiel anhand thematischer/semantischer Nähe oder Ko-Autorenschaften darzustellen. Dabei geht Mandl davon aus, dass Mehrwertkomponenten wie Personalisierung oder geographische Einschränkung von Suchergebnissen im Web häufig angeboten werden, sich jedoch nicht etablieren (Mandl 2006b, 65): „Der Benutzer widersetzte sich allem, was über eine Eingabezeile und eine Ergebnisliste hinausgeht“. In diesem Sinne sind Visualisierungen noch beschränkt auf die Darstellung von Information oder für die facetthierte Filterung. Mandl ordnet Visualisierungen folglich noch im Sinne einer entweder Ergebnisvisualisierung oder Anfragevisualisierung nach der Einordnung von Eibl ein (vgl. Abschnitt 3.10.4). Die Erweiterung ist der iterative Kreislauf von Wolff, der Visualisierungen sowohl für die Darstellung von Ergebnissen nutzt, als auch die Möglichkeit bietet, direkt in der Ergebnisdarstellung visuell und interaktiv erneut Anfragen zu stellen (vgl. Abschnitt 3.10.3). Dabei bleibt der Nutzer in der Modalität der grafischen Darstellung für Ergebnis und Anfrage und muss nicht zwischen grafischer Ergebnisdarstellung und formaler textueller Anfragesprache wechseln. Durch den fehlenden Modalitätswechsel verringert sich der kognitive Aufwand und die Komplexität von Folgeanfragen mit Ähnlichkeitsbezug wird in das System verlagert. Das ermöglicht es, Anfragen zu festzulegen, die vom Nutzer nur schwer definierbar wären. Auch bei Eibl (vgl. Abschnitt 3.10.4) wird gezeigt, dass die visuelle interaktive Gestaltung der Benutzungsoberfläche den Nutzer bei der komplexen Anfrageerstellung unterstützen kann. Dabei abstrahiert die Modellierung in dieser Arbeit Wolffs Modell um verschiedene heterogene Informationen, Informationsstrukturen, Visualisierungstechniken, Interaktionstechniken und die Einordnung in das Web- und IR-Retrieval.

Der Knowledge Crystallization-Prozess (KC) als Modell für die Informationssuche und Verarbeitung in der *Informationsvisualisierung* ist ähnlich angelegt wie der Exploratory Search-Ansatz (ES) im Bereich Informationssuche. Bei beiden Ansätzen handelt es sich um iterative Prozesse, beide Ansätze werden genutzt, um Lernprozesse, Wissen, Entscheidungen, Kommunikation oder Aktionen usw. zu generieren. Als Grundlage für beide Prozesse werden heterogene Informationen aus verschiedenen Datenquellen gesammelt („1. Information foreaging“ (KC) vs. „Lookup“ (ES)). Was im Exploratory Search-Ansatz als „Learn“ mit Prozessen wie „scannen/betrachten, vergleichen und qualitativ beurteilen“ kodiert ist, wird im Knowledge Crystallization-Prozess feiner durch den Prozess der Schema-Repräsentation als Konzept des „Sense-making“ aufgelöst.

Im Original-Papier (Russell u.a. 1993) wird „Sensemaking“ als der „Prozess für das Enkodieren von gesammelten Informationen für die Beantwortung von aufgabenspezifischen Fragestellungen“ definiert. Dabei werden Schemata als externe Repräsentation von Information gesucht, erstellt und angepasst, um die Kostenstrukturen für Operationen im Informationsverarbeitungsprozess zu reduzieren. Als Beispiel werden Eigenschaften von Druckern analysiert und in einem sich entwickelnden Schema festgehalten. Anhand des Schemas können Drucker verglichen und Cluster von verwandten Elementen wie Subsystemen oder Funktionen manuell, aber auch computer-basiert ermittelt werden. Durch die Enkodierung der Information in Repräsentationen und die darauf basierende Möglichkeit, computerbasiert Cluster von Eigenschaften zu berechnen, kann Zeit eingespart werden, die durch die Informationsverarbeiter an anderen Stellen genutzt werden kann.

Der „Learning Loop Complex“ des „Sensemaking“ ist als zyklischer, 4-stufiger Prozess angelegt:

- 1) *Search for representations*: Suche nach passenden Repräsentationen, die wichtige Eigenschaften von Informationen repräsentieren.
- 2) *Instantiate representations*: Suche nach passenden Informationen und Enkodierung im Schema.
- 3) *Shift representations*: Ständige Anpassung der Repräsentationen, um auf die Informationen abgestimmt zu werden und um die Kosten für den späteren Vergleich zu reduzieren.
- 4) *Consume encodings*: Die Repräsentationen können dann genutzt werden, um aufgabenspezifisch einfacher bestimmte Eigenschaften von Informationen zu vergleichen.

In den von den Autoren untersuchten Anwendungsfällen nimmt der Prozessschritt der Suche nach Informationen, der Extraktion von wichtigen Eigenschaften dieser Informationen und der Enkodierung in der Repräsentation 75% der Zeit im Gesamtprozess ein (Schritt 2 im „Learning Loop Complex“).

Analog dazu werden im Knowledge Crystallization-Prozess die Prozesse „2. Search for schema“, „3. Instantiate schema with data“ und „4. Problem-solve to trade off features“ aufgeführt. Der Nutzer sucht also nach Eigenschaften in den Informationen, welche für die Entscheidungsfindung wichtig sind (Schritt 2), instanziiert dieses Schema in einer externen Repräsentation (zum Beispiel in einem IV-System) (Schritt 3) und kann dann dort gesuchte Eigenschaften der Daten leichter ablesen (Schritt 4) oder neu arrangieren, um das Ablesen zu erleichtern (Schritt 5: „Search for a new schema that reduces the problem to a simple trade-off“).

Untersucht man KC als Grundlage für die Informationsvisualisierung, zeigen sich einige Limitierungen des Modells. Diese werden dadurch verursacht, dass sich das Modell auf den Prozess des „Sensemaking“ und damit auf die Reduktion des Modells auf einen Eigenschaftsvergleich von Informationen gründet.

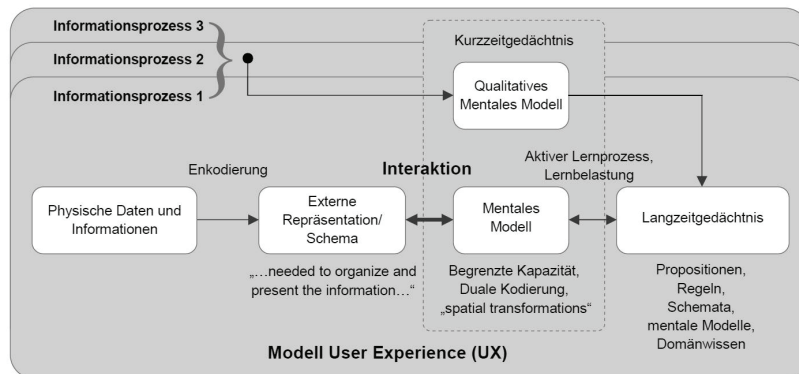
Im KC-Prozess können im ersten Schritt („forage for data“) als Informationsgrundlage große Mengen unterschiedlicher Daten („large amounts of heterogeneous data“) genutzt werden (Card, Mackinlay und Shneiderman 1999, 11). Diese heterogenen Daten müssen aber gemeinsame Attribute besitzen, damit sie überhaupt in einem externen Schema kodiert werden können. Heterogene Informationen im Sinne des Exploratory Search- und „Informationen im Web“-Ansatzes haben oft keine gemeinsamen Attribute, sondern sind vielleicht nur durch Verlinkungen oder einen Suchprozess miteinander verbunden. Folglich ist es schwer, diese in ein gemeinsames Schema und eine passende Repräsentation zu integrieren. Dies ist nur möglich, wenn die Informationen Attribute und außerdem auch noch weitere Eigenschaften wie zum Beispiel Modalität, Struktur oder Granularität teilen. Haben heterogene Informationen unterschiedliche Informationsstrukturen (tabellarisch, zeitlich/räumlich), können sie in koordinierten Ansichten repräsentiert werden, aber bei unterschiedlichen Modalitäten wird es ungleich schwieriger. Werden zum Beispiel heterogene Informationen aus verteilten Datenquellen wie Kurstrends und Finanznachrichten mit Abwägungsprozessen mental verglichen und abgewogen, funktioniert das KC-Modell bereits nicht mehr. Das einzige Informationsattribut, das die Informationen a priori teilen, ist hier das Zeitattribut oder ein bestimmter Zeitraum, in dem beide Informationen liegen könnten (ein Kurstrend liegt bspw. im gleichen Zeitraum wie eine Finanznachricht). Erst nach dem Erkenntnisprozess haben die Informationen eine Verlinkung oder Relation. Somit kann kein Schema festgelegt werden, das alle Informationen inkludiert, sondern der Suchprozess funktioniert nur, weil Visualisierungen interaktive Verlinkungen nach außen bieten (vgl. Modellbildung im vorherigen Artikel, Hienert 2014). Dagegen ist das von Card et al. angeführte Beispiel für den KCProzess recht simpel gehalten (Card, Mackinlay und Shneiderman 1999, 10): Hier werden einfache Eigenschaften von Laptop-Computern wie CPU-Geschwindigkeit, Gewicht, Dicke und Kosten in einer Tabelle verglichen, nicht komplexe heterogene Informationen. Auch die Schema-Repräsentation in einer einfachen Tabelle ist nicht passend für die Informationsvisualisierung, da eine Tabelle mehr ein visueller Formalismus (vgl. Nardi und Zamer 1993) als eine echte interaktive Visualisierung ist (auch wenn der Übergang zu Systemen wie TableLens fließend ist). Natürlich wurde das Beispiel sehr einfach gewählt, damit man den Prozess nachvollziehen und verstehen kann, aber die geforderten Eigenschaften des Modells werden dadurch nicht wirklich abgebildet. Die innerhalb des ersten Prozessschritts „forage for data“ aufgeführten, möglichen Interaktionsschritte sind neben dem Information-Seeking-Mantra („Overview, Zoom, Filter, Details-on-demand“) auch noch die Suchtechniken „Browse“ und „Search query“. Das Information-Seeking-Mantra wird innerhalb einer Visualisierung genutzt, Browsing und Search werden außerhalb der Visualisierung angewandt, wenn Visualisierungen in den nächsten Stufen des KC-Prozesses genutzt werden. Die Studie von Trafton u.a. (Trafton u.a. 2002) mit Wissenschaftlern zeigt,

dass heterogene Informationen nicht nur aus verschiedenen Quellen einer Repräsentationsschicht (zum Beispiel dem Web oder einem IV-System) entnommen werden, sondern dass Informationen auch aus verschiedenen hochkomplexen Expertensystemen und Visualisierungen (multiplen heterogenen externen Repräsentationen) entnommen werden, die sich so nicht einfach in eine externe Repräsentation einfügen lassen, sondern nur im mentalen Modell des Nutzers verarbeitet sind.

Dabei stellt sich die Frage, wie ein externes Schema, das zentral im KC-Prozess angelegt ist und ein qualitatives mentales Modell des Nutzers aufeinander aufbauen (vgl. Abbildung 12). Der KC-Prozess beschränkt sich, wie oben bereits aufgeführt, auf den Vergleich von Eigenschaften verschiedener Informationen, die in einem externen Schema repräsentiert („needed to organize and present the information“, (Nardi und Zarmer 1993, 10)) werden können und dessen visuelle Anordnung sich durch Interaktionsprozesse manipulieren lassen. So können im Sinne des „Sensemaking“ im KC-Prozess nur Ableitungen gemacht werden, wie sie in Abschnitt 3.1 Ziele und Vorteile der Informationsvisualisierung genannt werden. Dazu gehören Ableseprozesse wie Clusterbildung, Minimum, Maximum etc. Auf dieser Ebene wäre auch Pinkers Modell der Graphenwahrnehmung anzusiedeln (vgl. Abschnitt 4.7). Hier werden gelernte Schemata von Graphendarstellungen aus dem Langzeitgedächtnis abgerufen und angewandt. Externe Repräsentationen wurden als Schemata gelernt und sind als mentales Modell im Langzeitgedächtnis verfügbar. Sie können abgerufen und im Kurzzeitgedächtnis mit Daten instanziiert für die Ableitungen von konzeptuellen Fragen genutzt werden (einfache Ableseprozesse in der Grafik). Darüber hinaus ist aber auch in Pinkers Modell schon angelegt, dass durch Inferenzprozesse auch Informationen abgelesen werden, die nicht direkt in der Grafik abgebildet sind. Dies wird zum Beispiel in der Studie von Trafton u.a. (2002; 2001) gezeigt, in der Wissenschaftler Information nicht nur ablesen, sondern auch die mentale Vorstellungskraft und räumliche Transformation („spatial transformation“) nutzen, um Information, die nicht in der Visualisierung zu finden ist, zu extrapolieren. Diese Transformationen können sowohl mental als auch im IV-System durchgeführt werden. Auf dieser Ebene interagieren externe Repräsentation und das interne mentale Modell also sehr stark. Der Abgleich der beiden Instanzen geschieht durch Interaktion. Der Nutzer kann versuchen, Ableitungen aus dem mentalen Modell durch Interaktion wieder in der externen Repräsentation herzustellen oder aus der externen Repräsentation können sich Erkenntnisse bilden, die wieder im mentalen Modell repräsentiert werden. In diesem Sinne spielen die Interaktionsprozesse eine sehr wesentliche Rolle, wie auch durch Liu et al. (2008) mit Distributed Cognition als theoretisches Framework für die Informationsvisualisierung betont wird. Der andauernde Austausch zwischen externer und interner Repräsentation geschieht durch Interaktion und kann die Problemlösung fördern. Dabei lässt sich auch schon auf dieser Ebene die externe Repräsentation nicht mit

dem mentalen Modell gleichsetzen, da auf das mentale Modell weitere Faktoren wie der allgemeine Kontext (ausgedrückt durch das Modell der User Experience (Abschnitt 3.1)) und Wissen aus dem Langzeitgedächtnis (Propositionen, Regeln, Schemata, mentale Modelle) einen Einfluss haben. Die Konsolidierung der Erkenntnisse im Langzeitgedächtnis unterliegt einem aktiven Lernprozess mit Lernbelastungen (vgl. Abschnitte 4.4 bis 4.6).

Abb. 12: Zusammenhang zwischen physischen Informationen, externen Repräsentationen und dem mentalen Modell



Auf einer höheren Stufe unterstützen die externen Repräsentationen aber auch die Bildung eines qualitativen mentalen Modells, aus denen sich Ableitungen mental bilden lassen, wie in der Studie von Trafton et al. (2000) gezeigt wird. Hier werden aus stark heterogenen Informationsquellen und Darstellungen Informationen und Ableitungen integriert und daraus weitere Ableitungen gezogen. Genau diese mentalen Ableitungen machen interaktive Visualisierungen so wertvoll, da sie die Bildung eines qualitativen mentalen Modells unterstützen, indem komplexe mentale Prozesse wie Ableitungen, Abwägungen, qualitative Beurteilungen und Lernprozesse zwischen heterogenen Informationen, wie im ES-Modell gefordert, ermöglicht werden über einfache Vergleiche von Informationsattributen homogener Informationen, die sich in ein Schema enkodieren lassen. Gerade in diesen höherwertigen Vergleichen und Abwägungen fließen somit auch Domänenwissen, Erfahrung, Weltwissen usw. ein. Insofern muss beachtet werden, dass Visualisierungen nicht nur visuelle Repräsentationen abstrakter Daten darstellen, die isoliert in einer Visualisierung angezeigt werden, sondern auch Informationen repräsentieren, die zu anderen Informationen in Verbindung stehen.

Russel u.a. (1993) geben an, dass in den von ihnen untersuchten Anwendungsfällen die Aufgabe der Datenextraktion und -enkodierung 75% der Zeit im Gesamtprozess einnimmt. Das entspricht den Prozessschritten „forage for data“, „search for schema“ und „instantiate schema“ im KC-Prozess. Aber

gerade hier unterstützen die meisten Visualisierungen nicht, sondern bei den meisten interaktiven Visualisierungen besteht das Schema bereits, wird abgebildet, und der Nutzer sucht innerhalb dieser Visualisierung nach Mustern. Ergänzend geben Trafton u.a. (2002) an, dass die Bildung eines qualitativen mentalen Modells beim Nutzer die meiste Zeit im Gesamtprozess der Informationsintegration benötigt. Damit zeigt sich, dass sowohl die Aufarbeitung verschiedener Informationen in eine externe Repräsentation ein zeitintensiver Prozess ist als auch die Extraktion heterogener Informationen aus diesen externen Repräsentationen in ein qualitatives mentales Modell.

Card et al. (1999, 12) erklären, dass Visualisierungen auf den meisten Stufen des KC-Prozesses eingesetzt werden können und gibt Beispielsysteme für die einzelnen Prozessschritte. Moderne kommerzielle Visualisierungssysteme wie Tableau oder Spotfire sind in der Lage alle Prozessschritte des KC-Prozesses abzubilden. Aber der erste Schritt („forage for data“) wird nur insofern unterstützt, als dass Daten- und Informationsbestände in aufbereiteter tabellarischer Form geladen oder an Datenbanken angeschlossen werden können. Das steht diametral dem gegenüber, wie Informationen im Web durch Nutzer gesucht und gefunden werden. Die meisten Visualisierungssysteme beginnen bei Schritt 3 („instantiate schema“) im KC-Prozess: Informationen werden in einer Visualisierung dargestellt und der Nutzer kann durch die initiale Darstellung oder mit interaktiven Methoden Erkenntnisse aus der Darstellung generieren.

Verbindungen zwischen den einzelnen Schritten des Knowledge Crystallization-Prozesses bestehen damit bisher nur durch die Interaktion und Kognition des Nutzers. Der Nutzer ist das verbindende Glied der Prozesskette. Möchte der Nutzer neue Daten oder Informationen in der Visualisierung anzeigen, muss er sie erst im Web suchen, aufbereiten, ein passendes Schema kreieren und in ein passendes IV-System laden und anzeigen. Der Ansatz, Informationen auch auf Visualisierungsebene zu verbinden und mit Interaktionselementen browsbar zu machen (vgl. Abschnitt 3.4.4 in Hienert 2014), verbindet in diesem Modell zumindest die verschiedenen Schritte wie „forage for data“ und die Prozesse „search for schema“ und „instantiate schema“. Verschiedene heterogene Informationen sind bereits optimal in passenden Visualisierungsformen angezeigt und über browsbare Links auf Interaktionsebene verbunden.

IV beschränkt sich in seiner Rolle oft auf die interaktive Visualisierung von *Daten* (Definition zu Beginn dieses Abschnitts) versus *Informationen*.⁶ Dies stellt besonders den Anspruch der Informationsvisualisierung heraus *große Datenmengen* zu visualisieren (Card, Mackinlay und Shneiderman 1999, 11),

⁶ Definition zur Unterscheidung Daten vs. Information: Daten sind „zum Zweck der Verarbeitung zusammengefasste Zeichen, die aufgrund bekannter oder unterstellter Abmachungen Informationen (d.h. Angaben über Sachverhalte und Vorgänge) darstellen“ (Verlag, Hg. Gabler Wirtschaftslexikon, Stichwort: Daten, online im Internet. Abgerufen am 21. Sept. 2012.)

die sonst nicht mehr vom Nutzer kognitiv verarbeitbar wären. Da diese Daten aber entweder Eigenschaften von Informationen sind (wie es im KC-Prozess oder im Konzept des Sensemaking erklärt wird) oder bereits komplette Informationen abgebildet werden, repräsentieren sie auch immer Informationen. Diese Informationen haben im Gegensatz zu abstrakten Daten oft auch Verbindungen zu anderen Informationen im Web. Erweitert man also die Definition von IV von der Visualisierung von Daten auch auf die von Informationen ist nicht nur die Verbindung und Verlinkung innerhalb der Visualisierung und damit innerhalb eines Datensatzes wichtig, sondern auch die Verbindung nach außen zu anderen Informationen im Web.

Das Information Retrieval nutzt verschiedene Methoden für die Informationssuche wie Suche, Browsing, Filterung oder Faceted Search. In der Informationsvisualisierung herrscht das Paradigma von Shneiderman vor: Informationen werden hauptsächlich gefiltert und nur auf Detailebene soll zu verwandten Daten innerhalb des Datensatzes (und innerhalb der Visualisierung) gebrowst werden können. Dabei stellt sich die Frage, warum die Verlinkung von Informationen in Visualisierungen zu Informationen nach außen auf Detailebene, besonders im Kontext des Webs noch kein Standardkonzept ist, so dass die IR-Technik Browsing also auf allen Ebenen des Datensatzes genutzt wird. Gerade dadurch, dass der Nutzer im Web nur geringe Kenntnis über die angebotene Suchfunktionen und Recherchemöglichkeiten hat („Ill-formed queries“, (Lewandowski 2005)), sollten die Informationen verlinkt und durch Browsing schnell erreichbar sein. Voraussetzung für ein Browsing zwischen Information in und außerhalb von Visualisierungen ist eine Verlinkung auf Datenebene. Diese kann bereits auf verschiedenen Ebenen bestehen. Zum Beispiel werden Informationen im Semantic Web-Ansatz in heterogenen Informationsräumen vernetzt, um eine „integrierte globale Datenbank“ zu erreichen. Innerhalb von Informationsräumen bestehen Ontologien, die die Zusammenhänge und Verlinkungen innerhalb einer Domäne beschreiben. Besteht noch keine Verlinkung zwischen Informationen, so können Visualisierungen auch dazu genutzt werden, um diese Verlinkung zu erstellen. Dabei werden die Vorteile von interaktiven Visualisierungen genutzt, anhand von optimierten Schemata Informationen mittels Eigenschaften zu identifizieren. Geschieht dies für verschiedene Informationen in verschiedenen Ansichten, können diese Ansichten direkt genutzt werden, um die Informationen auszuwählen und mit einer einfachen Interaktionsmetapher zu verbinden. Dabei muss der Nutzer die grafische Modalität nicht verlassen und kann die Informationen dort verbinden, wo er sie gefunden hat (vgl. Abschnitt 3.10.3). Diese Verlinkung kann dann wiederum genutzt werden, um zwischen diesen Informationen zu browsen. Die Verlinkung ist auch eine wichtige Basis, um verschiedene Ansichten zu koordinieren. So können verlinkte Informationen oder verschiedene Eigenschaften von Information in verschiedenen Ansichten angezeigt werden. Die Ansichten kön-

nen so verbunden werden, dass bei der Auswahl einer Information in einer Ansicht verlinkte Informationen in anderen Ansichten hervorgehoben werden.

Auch die zweite Methodik aus dem IR, die Suche, wird nur selten als Standardtechnik im IV eingesetzt (vgl. einige Systeme in Abschnitt 3.10). Auch hier lässt sich argumentieren, dass Glyphen als Repräsentationen für die Suche eingesetzt werden können.

Filtering als dritte Technik kommt sowohl im IR als auch im IV zum Einsatz. Filtering im IR ist stark geprägt durch die Definition durch Belkin und Croft (1992), die auch von aktuellen Quellen wie Hanani u.a. (2001) als Basis genutzt wird. Dabei stellen sie folgende Punkte noch einmal heraus (Hanani, Shapira und Shoval 2001, 203f):

IF systems:

- are applicable for unstructured or semi-structured data (e.g. documents, e-mail, messages);
- handle large amounts of data;
- deal primarily with textual data;
- are based on user profiles; and
- their objective is to remove irrelevant data from incoming streams of data items.

Dabei wird Information Filtering in der Definition für IR stark auf Eigenschaften wie „semi-/unstrukturierte Daten“, „textuelle Daten“ und „Datenströme“ begrenzt. Diese Eingrenzung steht analog zur klassischen Definition von IR (vgl. Abschnitt 2.1.1), in der IR auf die Modalität Text eingegrenzt ist: „the means for identifying, retrieving, and/or ranking texts (or text surrogates or portions of texts) in a collections of texts, that might be relevant to a given query“ (Belkin und Croft 1987, 109). Auch wenn spätere IR-Modelle die Öffnung auf verschiedene Informationsarten (z.B. Fuhr, 1996 oder Exploratory Search, Marchionini, 2006) bereits vornehmen, wird hier *Information Filtering* als Konzept stark eingrenzt. Die Eigenschaften „are based on user profiles“ und „to remove data from incoming streams of data items“ schränken die Technik weiter auf ganz bestimmte Anwendungsbereiche des Konzepts ein. In der Informationsvisualisierung wird Filtering genutzt, um den Datenraum anhand bestimmter Parameter einzuschränken. Das kann mit interaktiven Methoden wie „Dynamic Queries“ durch den Nutzer durchgeführt werden (vgl. Abschnitt 3.8.3). Fokussiert man sich auf das eigentliche Konzept, so wird mit Filtering der Informationsraum im einfachsten Sinn anhand bestimmter Informationsattribute eingeschränkt. Diese Definition gilt sowohl für IR als auch für IV. Im IR könnte man sich auf die Einschränkung von Textdokumenten anhand bestimmter Informationsattribute (Titel, Jahr, Autoren etc.) fokussieren, aber im Sinne erweiterter Definitionen muss das Filtering für alle Informationsarten gelten. Im IV werden bereits unterschiedliche Informationsstrukturen anhand beliebiger Informationsattribute gefiltert, die für jede Instanz ausgewählt werden können. Dabei kann das Konzept Filtering in der Modellbildung zweifach

genutzt werden: (1) einmal ausgehend von Glyphen, die den Parameter für die Filterung bereitstellen (vgl. (Abschnitt 3.4.5 in Hienert 2014) oder (2) als globale Variante, die alle Ansichten gleichzeitig anhand extrahierter Parameter filtert (vgl. Abschnitt Modellbildung 3.4.2 in (Hienert 2014)).

Letztlich verschwimmt in einem Teilbereich die Grenze zwischen Informationsvisualisierung und Trefferlistendarstellung im Sinne des IR immer mehr, insofern als dass Visualisierungen hochinteraktive Systeme für Informationsdarstellungen sind. Dabei unterscheiden nur noch wenige Konzepte IR und IV: (1) Unterscheidung von Daten vs. Information, (2) Suchmethodiken, (3) textuelle vs. grafische Darstellung, (4) Interaktionskomponenten und (5) Ranking. Auf der Seite des IR werden Websuchmaschinen immer interaktiver. Interaktionstechniken wie Faceted Browsing oder Instant Search werden genutzt, um die Trefferlistendarstellungen analog zum Information-Seeking-Mantra sehr schnell zu filtern. Dabei sind die Trefferdarstellungen noch teilweise auf textuelle Strukturen beschränkt, aber andere Informationsarten und -darstellungen werden bereits in die Trefferlistendarstellung integriert (Bilder, Karten, Echtzeitnachrichten usw.), die auch einen sehr visuellen Charakter haben. Analog hängt Informationsdarstellung in IV auch sehr von textuellen Artefakten ab, ohne die die Erläuterung von visuellen Elementen nicht möglich ist. So sind auch in der Informationsvisualisierung meist Text und visuelle Elemente kombiniert. Abstrahiert man das Konzept in der Informationsvisualisierung von Datendarstellung zu Informationsdarstellung und führt weitere Techniken wie Suche, Filterung und Browsing ein, so lassen sich Visualisierungen auch für den Suchprozess im Sinne des IR nutzen.

Insgesamt lassen sich folgende Schlussfolgerungen für die Modellbildung ziehen:

- 1) Unterscheidet man nicht mehr explizit zwischen der Darstellung von Daten im IV und der Darstellung von Informationen im IR, werden Visualisierungen bereits für die interaktive Darstellung von Informationen genutzt.
- 2) Informationen haben Verbindungen zu anderen Informationen, die auch für den Suchprozess genutzt werden können.
- 3) Gerade dieser Prozess des Untersuchens, Vergleichens und Abwägens von heterogenen Informationen wird im Exploratory Search-Ansatz dargestellt.
- 4) Dagegen wird der Prozessschritt („forage for data“) im KC-Modell noch nicht wirklich durch interaktive Visualisierungen unterstützt. Der Schritt kann aber unterstützt werden, wenn das IV-Modell durch Interaktionstechniken wie Browsing, Filtering und Suche auf Informationsattribut-/ Glyphen-Ebene in der Visualisierung erweitert wird.
- 5) Diese Suchtechniken basieren teilweise auf einer Verlinkung von heterogenen Informationen, die aufgrund der Vorteile von Visualisierungen direkt auf visueller Ebene durchgeführt werden können.

- 6) Durch diese Erweiterung lassen sich interaktive Visualisierungen in einen allgemeinen Suchprozess im Web einbinden. Visualisierungen stehen nicht mehr isoliert da, sondern sind in den Suchprozess eingebunden.

Durch die Möglichkeit von IV, sehr unterschiedliche Informationen und Informationsstrukturen vorteilhaft abbilden zu können (vs. vorrangig Text im IR), können somit mehr Informationsarten intuitiv und interaktiv in den Suchprozess eingebunden werden.

Die hier beschriebenen Grundlagen der Informationssuche, Informationsvisualisierung und Informationsverarbeitung bilden die Voraussetzung für den Entwurf eines Modells für die Integration von interaktiven Visualisierungen in den Suchprozess des Webs, die im vorherigen Artikel dieses HSR-Focus (Hienert 2014) besprochen wurde.

References

- Ahlberg, Christopher. 1996. Spotfire: an information exploration environment. *SIGMOD Record* 25 (4): 25-9. doi: doi.acm.org/10.1145/245882.245893.
- Ahlberg, Christopher, und Ben Shneiderman. 1994. Visual information seeking: tight coupling of dynamic query filters with starfield displays. In *Proceedings of the SIGCHI conference on Human factors in computing systems: celebrating interdependence*, 313-7. CHI '94. New York, NY, USA: ACM. doi: 10.1145/191666.191775.
- Andrews, Keith. 2002. *Visualising Information Structures – Aspects of Information Visualisation*. Habilitation, Graz: Graz University of Technology.
- Andrews, Keith, Vedran Sabol, Wilfried Lackner, Christian Gütl, und Josef Moser. 2001. Search Result Visualisation with xFIND. In *Proceedings of the Second International Workshop on User Interfaces to Data Intensive Systems*, 50-8. UIDIS '01. Washington, DC, USA: IEEE Computer Society <<http://dl.acm.org/citation.cfm?id=572716.884776>>.
- Andrews, Keith, Josef Wollte, und Michael Pichler. 1997. Information Pyramids: A New Approach to Visualising Large Hierarchies. In *Proceedings of IEEE Visualization '97*, 49-52. VIS '97. Phoenix, Arizona: IEEE Computer Society <<http://citeseerx.ist.psu.edu/viewdoc/download?doi=10.1.1.87.3929&rep=rep1&type=pdf>>.
- Armstrong, Timothy G., Alistair Moffat, William Webber, und Justin Zobel. 2009. Improvements that don't add up: ad-hoc retrieval results since 1998. *Proceedings of the 18th ACM conference on Information and knowledge management*, 601-10. CIKM '09. New York, NY, USA: ACM. doi: 10.1145/1645953.1646031.
- Baddeley, Alan David. 1986. *Working Memory*. Oxford, UK; New York: Oxford University Press.
- Baddeley, Alan David. 2000. The Episodic Buffer: A New Component of Working Memory? *Trends in Cognitive Sciences* 4 (11): 417-23. doi: 10.1016/S1364-6613(00)01538-2.
- Baeza-Yates, Ricardo A., und Berthier Ribeiro-Neto. 1999. *Modern Information Retrieval*. Boston, MA, USA: Addison-Wesley Longman Publishing Co., Inc.

- Bates, Marcia J. 1989. The Design of Browsing and Berrypicking Techniques for the Online Search Interface. *Online Review* 13 (5): 407-24.
- Beaudoin, Luc, Marc-Antoine Parent, und Louis C. Vroomen. 1996. Cheops: A Compact Explorer for Complex Hierarchies. In *Proceedings of the 7th conference on Visualization '96*, 87-92. Los Alamitos, CA, USA: IEEE Computer Society Press <<http://dl.acm.org/citation.cfm?id=244979.245014>>.
- Becker, Richard A., und William S. Cleveland. 1987. Brushing scatterplots. *Technometrics* 29 (2): 127-42. doi: 10.2307/1269768.
- Belkin, Nicholas J., und W. Bruce Croft. 1987. Annual review of information science and technology. In *TITEL?*, hg. v. Martha E. Williams, 109-45. New York, NY, USA: Elsevier Science Inc. <<http://dl.acm.org/citation.cfm?id=42502.42506>>.
- Belkkin, Nicolas J., und W. Bruce Croft. 1992. Information filtering and information retrieval: two sides of the same coin? *Communications of the ACM* 35 (12): 29-38. doi: 10.1145/138859.138861.
- Berners-Lee, Tim. 2006. Linked Data – Design Issues. *W3C* <<http://www.w3.org/DesignIssues/LinkedData.html>> (Zugegriffen am 12. September 2012).
- Berners-Lee, Tim, James Hendler, und Ora Lassila. 2001. The Semantic Web. *Scientific American* 284 (5): 34-43.
- Bertin, Jacques. 1983. *Semiology of graphics*. University of Wisconsin Press.
- Bizer, Christian, Tom Heath, und Tim Berners-Lee. 2009. Linked data – the story so far. *International Journal on Semantic Web and Information Systems* 5 (3): 1-22.
- Bizer, Christian, Jens Lehmann, Georgi Kobilarov, Sören Auer, Christian Becker, Richard Cyganiak, und Sebastian Hellmann. 2009. DBpedia – A crystallization point for the Web of Data. *Web Semantics: Science, Services and Agents on the World Wide Web* 7 (3): 154-65. doi: 10.1016/j.websem.2009.07.002.
- Bostock, Michael, und Jeffrey Heer. 2009. Protovis: A Graphical Toolkit for Visualization. *IEEE Transactions on Visualization and Computer Graphics* 15 (6): 1121-8. doi: 10.1109/TVCG.2009.174.
- Brand, Matthias, und Hans J. Markowitsch. 2004. Lernen und Gedächtnis. *Praxis der Naturwissenschaften* 53 (7): 1-7.
- Bruckner, Stefan, Veronika Šoltészová, Eduard Groller, Jiří Hladřuvka, Katja Buhler, Jai Y. Yu, und Barry J. Dickson. 2009. BrainGazer – Visual Queries for Neurobiology Research. *IEEE Transactions on Visualization and Computer Graphics* 15 (6): 1497-504. doi: 10.1109/TVCG.2009.121.
- Card, Stuart K., Jock D. Mackinlay, und Ben Shneiderman. 1999a. Information Visualization. In *Readings in Information Visualization: Using Vision to Think*, 1-34. Morgan Kaufmann Series in Interactive Technologies. San Francisco, CA, USA: Morgan Kaufmann Publishers Inc. <<http://books.google.de/books?id=wdh2gqWfQmgC>>.
- Card, Stuart K., Jock D. Mackinlay, und Ben Shneiderman, Hg. 1999b. *Readings in information visualization – Using vision to think*. Morgan Kaufmann Series in Interactive Technologies. San Francisco, CA, USA: Morgan Kaufmann Publishers Inc. <<http://dl.acm.org/citation.cfm?id=300679.300826>>.
- Carpendale, Sheelagh. 2008. Evaluating Information Visualizations. In *Information Visualization*, hg. v. Andreas Kerren, John T. Stasko, Jean-Daniel Fekete und Chris North, 19-45. Berlin, Heidelberg: Springer. doi: 10.1007/978-3-540-70956-5_2.

- Carpenter, Patricia A., und Priti Shah. 1998. A model of the perceptual and conceptual processes in graph comprehension. *Journal of Experimental Psychology Applied* 4 (2): 75-100.
- Chandler, Paul, und John Sweller. 1991. Cognitive Load Theory and the Format of Instruction. *Cognition and Instruction* 8 (4): 293-332.
- Chawda, B., B. Craft, P. Cairns, D. Heesch, und S. Rüger. 2005. Do attractive things work better? An exploration of search tool visualisations. In *Proceedings of '05 British HCI group annual conference*, hg. v. L. McKinnon, O. Bertelsen und N. Bryan-Kinns, 46-9. HCI2005. Swindon, UK: British HCI Group.
- Chi, Ed H., Peter Pirolli, Kim Chen, und James Pitkow. 2001. Using information scent to model user information needs and actions and the Web. In *Proceedings of the SIGCHI conference on Human factors in computing systems*, 490-97. CHI '01. New York, NY, USA: ACM. doi: 10.1145/365024.365325.
- Cleveland, William S., und Robert McGill. 1984. Graphical Perception: Theory, Experimentation, and Application to the Development of Graphical Methods. *Journal of the American Statistical Association* 79 (387): 531-54.
- Cleveland, William S., und Robert McGill. 1986. An Experiment in Graphical Perception. *International Journal of Man-Machine Studies* 25 (5): 491-501.
- Dion, Karen, Ellen Berscheid, und Elaine Walster. 1972. What is beautiful is good. *Journal of Personality and Social Psychology* 24 (3): 285-90.
- Dörk, Marian, Sheelagh Carpendale, Christopher Collins, und Carey Williamson. 2008. VisGets: Coordinated Visualizations for Web-based Information Exploration and Discovery. *IEEE Transactions on Visualization and Computer Graphics* 14 (6): 1205-12. doi: 10.1109/TVCG.2008.175.
- Dörk, Marian, Daniel M. Gruen, Carey Williamson, und M. Sheelagh T. Carpendale. 2010. A Visual Backchannel for Large-Scale Events. *IEEE Transactions on Visualization and Computer Graphics* 16 (6): 1129-38.
- Eagly, Alice H., Richard D. Ashmore, Mona G. Makhijani, und Laura C. Longo. 1991. What is beautiful is good, but...: A meta-analytic review of research on the physical attractiveness stereotype. *Psychological Bulletin* 110 (1): 109-28.
- Eibl, Maximilian. 2002. DEViD: a media design and software ergonomics integrating visualization for document retrieval. *Information Visualization* 1 (2): 139-57. doi: 10.1057/palgrave.ivs.9500017.
- Eibl, Maximilian. 2003. *Visualisierung im Document Retrieval - Theoretische und praktische Zusammenführung von Softwareergonomie und Graphik Design*. IZ-Forschungsbericht. Informationszentrum Sozialwissenschaften – IZ Bonn.
- English, Jennifer, Marti Hearst, Rashmi Sinha, Kirsten Swearingen, und Ka-Ping Yee. 2002. Flexible Search and Navigation Using Faceted Metadata. In *Proceedings of the ACM SIGIR Conference on Research and Development of Information Retrieval*, 11-5. SIGIR '02.
- Fekete, Jean-Daniel. 2004. The InfoVis Toolkit. In *Proceedings of the IEEE Symposium on Information Visualization*, 167-74. INFOVIS '04. Washington, DC, USA: IEEE Computer Society. doi: 10.1109/INFOVIS.2004.64.
- Fuhr, Norbert. 1996. Ziele und Aufgaben der Fachgruppe Information Retrieval der Gesellschaft für Informatik <http://www.uni-hildesheim.de/fgir/index.php?option=com_content&task=view&id=14&Itemid=41> (Zugegriffen am 12. September 2012).

- Gershon, Nahum, und Ward Page. 2001. What storytelling can do for information visualization. *Communications of the ACM* 44 (8): 31-7. doi: 10.1145/381641.381653.
- Hanani, Uri, Bracha Shapira, und Peretz Shoval. 2001. Information Filtering: Overview of Issues, Research and Systems. *User Modeling and User-Adapted Interaction* 11 (3): 203-59. doi: 10.1023/A:1011196000674.
- Hassenzahl, Marc. 2004. The interplay of beauty, goodness, and usability in interactive products. *Human-Computer Interaction* 19 (4): 319-49. doi: 10.1207/s15327051hci1904_2.
- Hassenzahl, Marc. 2008. Aesthetics in interactive products: Correlates and consequences of beauty. In *Product experience*, hg. v. Hendrik N. J. Schifferstein und Paul Hekkert, 287-99. Elsevier.
- Hassenzahl, Marc, und Noam Tractinsky. 2006. User experience – a research agenda. *Behaviour & Information Technology* 25 (2): 91-7. doi: 10.1080/01449290500330331.
- Havre, Susan, Beth Hetzler, und Lucy Nowell. 2000. ThemeRiver: Visualizing Theme Changes over Time. In *Proceedings of the IEEE Symposium on Information Visualization 2000*, 115-23. INFOVIS '00. Washington, DC, USA: IEEE Computer Society <<http://dl.acm.org/citation.cfm?id=857190.857680>>.
- Heer, Jeffrey, Stuart K. Card, und James A. Landay. 2005. Prefuse: a toolkit for interactive information visualization. In *Proceedings of the SIGCHI conference on Human factors in computing systems*, 421-30. CHI '05. New York, NY, USA: ACM. doi: 10.1145/1054972.1055031.
- Heer, Jeffrey, Fernanda B. Viégas, Martin Wattenberg, und Maneesh Agrawala. 2009. Point, Talk, and Publish: Visualisation and the Web. In *Trends in Interactive Visualization*, hg. v. Lakhmi Jain, Xindong Wu, Robert Liere, Tony Adriaansen und Elena Zudilova-Seinstra, 1-15. Advanced Information and Knowledge Processing. London: Springer. doi: 10.1007/978-1-84800-269-2_12.
- Hegner, Marcus. 2003. *Methoden zur Evaluation von Software*. Arbeitsbericht. Bonn: Informationszentrum Sozialwissenschaften <http://www.gesis.org/fileadmin/upload/forschung/publikationen/gesis_reihen/iz_arbeitsberichte/ab_29.pdf>.
- Hemmje, Matthias. 1995. LyberWorld: a 3D graphical user interface for fulltext retrieval. In *Proceedings ACM SIGCHI '95*, 417-8. SIGCHI '95. ACM.
- Hemmje, Matthias, Clemens Kunkel, und Alexander Willett. 1994. LyberWorld – a visualization user interface supporting fulltext retrieval. In *Proceedings of the 17th annual international ACM SIGIR conference on research and development in information retrieval*, hg. v. W. Bruce Croft und C. J. van Rijsbergen, 249-59. SIGIR '94. New York, NY, USA: Springer <<http://dl.acm.org/citation.cfm?id=188490.188563>>.
- Hienert, Daniel. 2014. A Model for the Integration of Interactive Visualizations in the Process of Information Searching and Linking on the Web. *Historical Social Research* 39 (3): 179-192. doi: 10.12759/hsr.39.2014.3.179-192.
- Hilbert, Martin, und Priscila López. 2011. The World's Technological Capacity to Store, Communicate, and Compute Information. *Science* 332 (6025): 60-5. doi: 10.1126/science.1200970.
- Holmqvist, Lars Erik. 1997. Focus+context visualization with flip zooming and the zoom browser. In *CHI '97 extended abstracts on Human factors in computing systems*.

- tems: looking to the future, 263-64. CHI EA '97. New York, NY, USA: ACM. doi: 10.1145/1120212.1120383.
- Honkela, Timo, Samuel Kaski, Krista Lagus, und Teuvo Kohonen. 1998. WEB-SOM – Self-Organizing Maps of Document Collections. *Neurocomputing* 1-3: 101-17. doi: 10.1016/S0925-2312(98)00039-3.
- Huynh, David F., David R. Karger, und Robert C. Miller. 2007. Exhibit: light-weight structured data publishing. In *Proceedings of the 16th international conference on World Wide Web*, 737-46. WWW '07. Banff, Alberta, Canada: ACM. doi: 10.1145/1242572.1242672.
- Ingwersen, P., und K. Järvelin. 2007. On the holistic cognitive theory for information retrieval: drifting outside the border of the laboratory framework. In *Studies in the Theory of Information Retrieval (ICTIR 2007), Foundation for Information Society*, hg. v. S. Sominić und F. Kiss, 135-47. Budapest, Hungary.
- Inselberg, A. 1997. Multidimensional detective. In *Proceedings of the 1997 IEEE Symposium on Information Visualization*, 100-7. INFOVIS '97. Washington, DC, USA: IEEE Computer Society <<http://dl.acm.org/citation.cfm?id=857188.857631>>.
- Jog, N. K., und B. Shneiderman. 1995. Starfield visualization with interactive smooth zooming. *Proceedings of the third IFIP WG2.6 working conference on Visual database systems 3 (VDB-3)*, 3-14. London, UK: Chapman & Hall, Ltd. <<http://dl.acm.org/citation.cfm?id=265394.265401>>.
- Kappe, Frank, Georg Droschl, Wolfgang Kienreich, Vedran Sabol, Jutta Becker, Keith Andrews, Michael Granitzer, Klaus Tochtermann, und Peter Auer. 2002. InfoSky: Eine neue Technologie zur Erforschung großer, hierarchischer Wissensräume. *KnowTech 2002, 4. Konferenz zum Einsatz von Knowledge Management in Wirtschaft und Verwaltung (www.knowtech2002.de)*. Munich, Germany.
- Keim, Daniel A. 2002. Information Visualization and Visual Data Mining. *IEEE Transactions on Visualization and Computer Graphics* 8 (1): 1-8 doi: 10.1109/2945.981847.
- Keskin, Can, und Volker Vogelmann. 1997. Effective visualization of hierarchical graphs with the cityscape metaphor. *Proceedings of the 1997 workshop on New paradigms in information visualization and manipulation*, 52-7. NPIV '97. New York, NY, USA: ACM. doi: 10.1145/275519.275531.
- Kim, Jinwoo, Joeun Lee, und Dongseong Choi. 2003. Designing emotionally evocative homepages: an empirical study of the quantitative relations between design factors and emotional dimensions. *International Journal of Human-Computer Studies* 59 (6): 899-940. doi: 10.1016/j.ijhcs.2003.06.002.
- Kosslyn, S. M. 1989. Understanding charts and graphs. *Applied Cognitive Psychology* 3: 185-226.
- Krause, Jürgen. 1996. *Visualisierung und graphische Benutzungsoberflächen*. IZ-Arbeitsbericht. Arbeitsbericht Informationszentrum Sozialwissenschaften – IZ Bonn und Universität Koblenz-Landau/Institut für Informatik Koblenz <<http://scholar.google.com/scholar?hl=en&btnG=Search&q=intitle:Visualisierung+und+graphische+Benutzungsoberflächen#0>>.
- Kuhlen, Rainer. 1995. *Informationsmarkt: Chancen und Risiken der Kommerzialisierung von Wissen*. Konstanz: UVK, Universitätsverlag <http://gso.gbv.de/DB=2.1/CMD?ACT=SRCHA&SRT=YOP&IKT=1016&TRM=ppn+183276760&sourceid=fbw_bibsonomy>.

- Kurosu, Masaaki, und Kaori Kashimura. 1995. Apparent usability vs. inherent usability: experimental analysis on the determinants of the apparent usability. *Conference companion on Human factors in computing systems*, 292-3. CHI '95. New York, NY, USA: ACM. doi: 10.1145/223355.223680 <<http://doi.acm.org/10.1145/223355.223680>>.
- Lam, H., E. Bertini, P. Isenberg, Catherine Plaisant, und S. Carpendale. 2011. *Seven guiding scenarios for information visualization evaluation*. Technical Report. Department of Computer Science. University of Calgary.
- Lamping, John, Ramana Rao, und Peter Pirolli. 1995. A focus+context technique based on hyperbolic geometry for visualizing large hierarchies. *Proceedings of the SIGCHI conference on Human factors in computing systems*, 401-8. CHI '95. New York, NY, USA: ACM Press/Addison-Wesley Publishing Co. doi: 10.1145/223904.223956.
- Lassila, Ora, und Ralph R. Swick. 1999. *Resource Description Framework (RDF) Model and Syntax Specification*. W3C Recommendation. W3C <<http://www.w3.org/TR/1999/REC-rdf-syntax-19990222/>>.
- Lavie, Talia, und Noam Tractinsky. 2004. Assessing dimensions of perceived visual aesthetics of web sites. *International Journal of Human-Computer Studies* 60 (3): 269-98. doi: 10.1016/j.ijhcs.2003.09.002.
- Lewandowski, Dirk. 2005. *Web information retrieval*. Frankfurt am Main: Deutsche Gesellschaft für Informationswissenschaft und Informationspraxis <<http://www.durchdenken.de/lewandowski/web-ir>>.
- Liu, Zhicheng, Nancy Nersessian, und John Stasko. 2008. Distributed Cognition as a Theoretical Framework for Information Visualization. *IEEE Transactions on Visualization and Computer Graphics* 14 (6): 1173-80. doi: 10.1109/TVCG.2008.121.
- Lohse, Gerald Lee. 1993. A cognitive model for understanding graphical perception. *Human-Computer Interaction* 8 (4): 353-88. doi: 10.1207/s15327051hci0804_3.
- Lynch, Kevin. 1960. *The Image of the City*. The MIT Press.
- Mackinlay, Jock D., Pat Hanrahan, und Chris Stolte. 2007. Show Me: Automatic Presentation for Visual Analysis. *IEEE Transactions on Visualization and Computer Graphics* 13 (6): 1137-44. doi: <http://dx.doi.org/10.1109/TVCG.2007.705> 94.
- Mackinlay, Jock D., George G. Robertson, und Stuart K. Card. 1991. The perspective wall: detail and context smoothly integrated. *Proceedings of the SIGCHI conference on Human factors in computing systems: Reaching through technology*, 173-6. CHI '91. New York, NY, USA: ACM. doi: 10.1145/108844.108870.
- Mandl, Thomas. 2006a. Implementation and evaluation of a quality-based search engine. *Proceedings of the seventeenth conference on Hypertext and hypermedia*, 73-84. HYPERTEXT '06. New York, NY, USA: ACM. doi: 10.1145/1149941.1149957.
- Mandl, Thomas. 2006b. *Die automatische Bewertung der Qualität von Internet-Seiten im Information Retrieval*. Habilitation, Universität Hildesheim, Fachbereich III – Institut für Angewandte Sprachwissenschaft.
- Marchionini, Gary. 2006. Exploratory search: from finding to understanding. *Communications of the ACM* 49 (4): 41-6. doi:10.1145/1121949.1121979.
- Mayer, Richard E. 2001. *Multimedia Learning*. Cambridge University Press.
- Mayer, Richard E. 2005. Cognitive Theory of Multimedia Learning. *The Cambridge handbook of multimedia learning*, hg. v. Richard E. Mayer, 31-48. Cambridge University Press <<http://books.google.com/books?hl=en&lr=&id=duWx8>>.

- fxkkk0C&oi=fnd&pg=PA31&dq=Cognitive+Theory+of+Multimedia+Learning
&ots=x62cq2Ag4C&sig=km1JhOxQR2OGZunazZVGY6qDU3Q>.
- McKeon, Matt. 2009. Harnessing the Information Ecosystem with Wiki-based Visualization Dashboards. *IEEE Transactions on Visualization and Computer Graphics* 15 (6): 1081-8. doi:10.1109/TVCG.2009.148.
- Meier, Beat. 1999. *Differentielle Gedächtniseffekte: Implizite und explizite Erfahrungsnachwirkungen aus experimenteller und psychometrischer Perspektive*. Internationale Hochschulschriften 308. Münster, New York, München, Berlin: Waxmann.
- Miller, George A. 1956. The Magical Number Seven, Plus or Minus Two: Some Limits on Our Capacity for Processing Information. *The Psychological Review* 63 (2): 81-97.
- Nardi, Bonnie A., und Craig L. Zarter. 1993. Beyond Models and Metaphors: Visual Formalisms in User Interface Design. *Journal of Visual Languages and Computing* 4 (1): 5-33.
- Norman, Donald A. 1999. *The invisible computer: why good products can fail, the personal computer is so complex, and information appliances are the solution*. MIT Press <http://gso.gbv.de/DB=2.1/CMD?ACT=SRCHA&SRT=YOP&IKT=1016&TRM=ppn+303799986&sourceid=fbw_bibsonomy>.
- North, Chris. 2005. Information Visualization. *Handbook of Human Factors and Ergonomics*, 3. Aufl., hg. v. G. Salvendy, 1222-46. New York: John Wiley & Sons.
- North Chris. 2006. Toward Measuring Visualization Insight. *IEEE Computer Graphics and Applications* 26 (3): 6-9. doi: 10.1109/MCG.2006.70.
- North, Chris, Nathan Conklin, Kiran Indukuri, und Varun Saini. 2002. Visualization schemas and a web-based architecture for custom multiple-view visualization of multiple-table databases. *Information Visualization* 1 (3-4): 211-28. doi: 10.1057/palgrave.ivs.9500020.
- North, Chris, und Ben Shneiderman. 2000. Snap-together visualization: a user interface for coordinating visualizations via relational schemata. *Proceedings of the working conference on Advanced visual interfaces*, 128-35. AVI '00. New York, NY, USA: ACM <<http://doi.acm.org/10.1145/345513.345282>>.
- Olston, Christopher, und Ed H. Chi. 2003. ScentTrails: Integrating browsing and searching on the Web. *ACM Transactions on Computer-Human Interaction* 10 (3): 177-97. doi: 10.1145/937549.937550.
- Paivio, Allan. 1986. *Mental Representations: A Dual Coding Approach*. Oxford Psychology Series 9. Oxford University Press.
- Peterson, Lloyd, und Margaret Jean Peterson. 1959. Short-term retention of individual verbal items. *Journal of Experimental Psychology* 58 (3): 193-8. doi: 10.1037/h0049234.
- Pinker, Steven. 1990. A theory of graph comprehension. In *Artificial Intelligence and the future of testing*, hg. v. R. Feedle, 73-126. Lawrence Erlbaum Associates Inc.
- Pirolli, Peter, und Stuart K. Card. 1999. Information Foraging. *Psychological Review* 106: 643-75.
- Pirolli, Peter, und Ramana Rao. 1996. Table lens as a tool for making sense of data. *Proceedings of the workshop on Advanced visual interfaces*, 67-80. AVI '96. New York, NY, USA: ACM <<http://doi.acm.org/10.1145/948449.948460>>.

- Plaisant, Catherine. 2004. The challenge of information visualization evaluation. *Proceedings of the working conference on Advanced visual interfaces*, 109-16. AVI '04. New York, NY, USA: ACM. doi: 10.1145/989863.989880.
- Ratwani, Raj M., Gregory J. Trafton, und Deborah A. Boehm-Davis. 2008. Thinking graphically: Connecting vision and cognition during graph comprehension. *Journal of Experimental Psychology: Applied* 14 (1): 36-49.
- Reas, Casey, und Benjamin Fry. 2003. Processing: a learning environment for creating interactive Web graphics. *ACM SIGGRAPH 2003 Sketches & Applications*, 1. SIGGRAPH '03. New York, NY, USA: ACM. doi: 10.1145/965400.965440.
- Rekimoto, Jun, und Mark Green. 1993. The Information Cube: Using Transparency in 3D Information Visualization. *Proceedings of the Third Annual Workshop on Information Technologies & Systems*, 125-32. WITS'93.
- Rijsbergen, C. J. van. 1979. *Information retrieval* 2. Aufl. Butterworths, London, Boston.
- Robertson, George G., Jock D. Mackinlay, und Stuart K. Card. 1991. Cone Trees: animated 3D visualizations of hierarchical information. *CHI '91: Proceedings of the SIGCHI conference on Human factors in computing systems*, 189-94. New York, NY, USA: ACM.
- Russell, Daniel M., Mark J. Stefik, Peter Pirolli, und Stuart K. Card. 1991. Cibe trees: animated 3D visualizations of hierarchical information. *CHI '91: Proceedings of the SIGCHI conference on Human factors in computing systems*, 189-94. New York, NY, USA: ACM. doi: 10.1145/169059.169209.
- Schaer, Philipp. 2013a. Applied Informetrics for Digital Libraries: An Overview of Foundations, Problems and Current Approaches. *Historical Social Research* 38 (3): 267-81.
- Schaer, Philipp. 2013b. Information Retrieval und Informetrie: Zur Anwendung informetrischer Methoden in digitalen Bibliotheken. *Historical Social Research* 38 (3): 282-354.
- Schierl, Thomas. 2001. *Text und Bild in der Werbung: Bedingungen, Wirkungen und Anwendungen bei Anzeigen und Plakaten*. Halem.
- Schnotz, Wolfgang. 2005. An Integrated Model of Text and Picture Comprehension. In *The Cambridge Handbook of Multimedia Learning*, hg. v. Richard E. Mayer, 49-67. Cambridge University Press.
- Schnotz, Wolfgang, und Maria Bannert. 2003. Construction and interference in learning from multiple representation. *Learning and Instruction* 13 (2): 141-56.
- Shneiderman, Ben. 1996. The Eyes Have It: A Task by Data Type Taxonomy for Information Visualizations. *Proceedings of the 1996 IEEE Symposium on Visual Languages*, 336-43. VL '96. Washington, DC, USA: IEEE Computer Society <<http://dl.acm.org/citation.cfm?id=832277.834354>>.
- Shneiderman, Ben. 2009. *Treemaps for space-constrained visualization of hierarchies* <<http://www.cs.umd.edu/hcil/treemap-history/index.shtml>> (Zugegriffen am 10. September 2012).
- Smith, Greg, Mary Czerwinski, Brian Meyers, Daniel Robbins, George Robertson, und Desney S. Tan. 2006. FacetMap: A Scalable Search and Browse Visualization. *IEEE Transactions on Visualization and Computer Graphics* 12 (5): 797-804. doi: 10.1109/TVCG.2006.142.
- Snowdon, Dave, und Kai-Mikael Jää-Aro. 1997. A Subjective Virtual Environment for Collaborative Information Visualization. *Virtual Reality Universe '97*, 2-5.

- Sparacino, Flavia, Glorianna Davenport, and Alex Pentland. 1996. City of News: cataloguing the World Wide Web through Virtual Architecture. *SIGGRAPH 99, Visual Proceedings, Emerging Technologies*. SIGGRAPH '99. Los Angeles, CA, USA.
- Squire, Larry R. 1992. Memory and the Hippocampus: A Synthesis from Findings with Rats, Monkeys, and Humans. *Psychological Review* 99 (2): 195-231.
- Squire, Larry R., Barbara Knowlton, and Gail Musen. 1993. The Structure and Organization of Memory. *Annual Review of Psychology* 44 (1): 453-95.
- Sweller, John. 2003. Evolution of human cognitive architecture. *Psychology of Learning and Motivation* 43: 215-66. Elsevier <<http://usaoll.org/iddtheorywb/pdf/Sweller-2003.pdf>>.
- Sweller, John. 2005. Implications of Cognitive Load Theory for Multimedia Learning. In *The Cambridge Handbook of Multimedia Learning*, hg. v. Richard E. Mayer, 19-30. Cambridge University Press.
- Tractinsky, Noam. 1997. Aesthetics and apparent usability: empirically assessing cultural and methodological issues. *Proceedings of the SIGCHI conference on Human factors in computing systems*, 115-22. CHI '97. New York, NY, USA: ACM. doi: 10.1145/258549.258626.
- Tractinsky, Noam. 2005. Does Aesthetics Matter in Human-Computer Interaction? In *Mensch und Computer 2005: Kunst und Wissenschaft – Grenzüberschreitung der interaktiven Art*, hg. v. Christian Stary, 29-42. München: Oldenbourg Verlag.
- Tractinsky, Noam, Adi S. Katz, and Dror Ikar. 2000. What is beautiful is usable. *Interacting with Computers* 13 (2): 127-45. doi: 10.1016/S0953-5438(00)00031-X.
- Tractinsky, Noam, and Oded Lowengard. 2007. Web-Store Aesthetics in E-Retailing: A Conceptual Framework and Some Theoretical Implications. *Academy of Marketing Science Review [Online-Journal]* 11 (1).
- Trafton, J. Gregory, Susan S. Kirschenbaum, Ted L. Tsui, Robert T. Miyamoto, James A. Ballas, and Paula D. Raymond. 2000. Turning pictures into numbers: extracting and generating information from complex visualizations. *International Journal of Human-Computer Studies* 53 (5): 827-50.
- Trafton, J. Gregory, Sandra P. Marshall, Farilee Mintz, and Susan B. Trickett. 2002. Extracting explicit and implicit information from complex visualizations. In *Proceedings of the Second International Conference on Diagrammatic Representation and Inference*, hg. v. M. Hegarty, B. Meyer und H. Narayanan, 206-20. DIAGRAMS '02. London, UK: Springer.
- Trafton, J. Gregory, and Susan B. Trickett. 2001. A New Model of Graph and Visualization Usage. In *Proceedings of the twenty-third annual conference of the cognitive science society*, hg. v. J. D. Moore und K. Stenning, 1048-53. Mahwah, NJ: Erlbaum.
- Tufte, Edward Rolf. 1999. *Beautiful evidence*. Cheshire, Connecticut: Graphics Press <http://www.worldcat.org/search?qt=worldcat_org_all&q=0961392177>.
- Tulving, Endel. 1972. Episodic and semantic memory. In *Organization of memory*, hg. v. Endel Tulving und W Donaldson, 381-403. New York: Academic Press.
- Tulving, Endel. 1985. Memory and consciousness. *Canadian Psychology/Psychologie canadienne* 26 (1). doi: 10.1037/h0080017.
- Tulving, Endel, und Daniel L. Schacter. 1990. Priming and Human Memory Systems. *Science* 247: 301-6.
- Viegas, Fernanda B., Martin Wattenberg, Frank van Ham, Jesse Kriss, und Matt McKeon. 2007. ManyEyes: a Site for Visualization at Internet Scale. *IEEE*

- Transactions on Visualization and Computer Graphics* 13 (6): 1121-8. doi: 10.1109/TVCG.2007.70577.
- Voorhees, Ellen M. 2002. The Philosophy of Information Retrieval Evaluation. *Revised Papers from the Second Workshop of the Cross-Language Evaluation Forum on Evaluation of Cross-Language Information Retrieval Systems*, 355-70. CLEF '01. London, UK: Springer <<http://dl.acm.org/citation.cfm?id=648264.753539>>.
- Voorhees, Ellen M., und Donna K. Harman. 2005. *TREC: Experiment and evaluation in information retrieval*, Bd. 63. Cambridge, MA: MIT Press <http://scholar.google.de/scholar.bib?q=info:KfjP-HMJmtkJ:scholar.google.com/&output=citation&hl=de&as_sdt=0,5&ct=citation&cd=0>.
- Wang Baldonado, Michelle Q., Allison Woodruff, und Allan Kuchinsky. 2000. Guidelines for using multiple views in information visualization. *Proceedings of the working conference on Advanced visual interfaces*, 110-9. AVI '00. Palermo, Italy. doi: 10.1145/345513.345271 <<http://portal.acm.org/citation.cfm?doid=345513.345271>>.
- Ware, Colin. 2005. Visual queries: the foundation of visual thinking. In *Knowledge and Information Visualization*, hg. v. Sigmar-Olaf Tergan und Tanja Keller, 27-35. Berlin, Heidelberg: Springer <<http://dl.acm.org/citation.cfm?id=2206936.2206940>>.
- White, Ryan W., Bill Kules, Steven M. Drucker, und m.c. schraefel. 2006. Supporting Exploratory Search. *Communications of the ACM* 49 (4): 36-9. doi: 10.1145/1121949.1121978.
- Willett, Wesley, Jeffrey Heer, und Maneesh Agrawala. 2007. Scented Widgets: Improving Navigation Cues with Embedded Visualizations. *IEEE Transactions on Visualization and Computer Graphics* 13 (6): 1129-36. doi: 10.1109/TVCG.2007.70589.
- Wittrock, Merlin C. 1989. Generative processes of comprehension. *Educational Psychologist* 24 (4): 345-76.
- Wolff, Christian. 1996. *Graphisches Faktenretrieval mit Liniendiagrammen*. Schriften zur Informationswissenschaft 24. Konstanz: Universitätsverlag Konstanz.
- Wolff, Christian, und Christa Womser-Hacker. 1997. Graphisches Faktenretrieval mit vager Anfrageinterpretation. *Theorien, Modelle und Implementierungen integrierter elektronischer Informationssysteme*, 251-64. Konstanz.
- Wong, Pak Chung, Beth Hetzler, Christian Posse, Mark Whiting, Susan Havre, Nick Cramer, Anuj Shah, Mudita Singhal, Alan Turner, und Jim Thomas. 2004. IN-SPIRE InfoVis 2004 Contest Entry. *Proceedings of the IEEE Symposium on Information Visualization*, 216. INFOVIS '04. Washington, DC, USA: IEEE Computer Society. doi: 10.1109/INFOVIS.2004.37.
- Wood, Jo, Jason Dykes, Aidan Slingsby, und Keith Clarke. 2007. Interactive Visual Exploration of a Large Spatio-temporal Dataset: Reflections on a Geovisualization Mashup. *IEEE Transactions on Visualization and Computer Graphics* 13 (6): 1176-83. doi: 10.1109/TVCG.2007.70570.