

Visualisierung und Analyse multi-dimensionaler Datensätze

Dirk J. Lehmann*

Georgia Albuquerque†

Martin Eisemann‡

Andrada Tatu§

Daniel Keim¶

Heidrun Schumann||

Marcus Magnor**

Holger Theisel††

Zusammenfassung

Für multi-dimensionale Datensätze existieren eine Reihe von automatischen Analysemethoden und Visualisierungstechniken, um ihnen innewohnende Zusammenhänge und Charakteristika aufzudecken. Die zunehmende Größe und Komplexität solcher Daten macht es notwendig, beide Ansätze miteinander zu kombinieren. In diesem Artikel stellen wir Ihnen daher etablierte Methoden zur visuellen und zur automatischen Datenanalyse vor und zeigen neuere Ansätze auf, diese sinnvoll miteinander zu kombinieren. Dabei werden alle Erläuterungen anhand anschaulicher Beispiele verdeutlicht und so für den Leser nachvollziehbar.

Schlüsselwörter: Visualisierung, Datenanalyse, Data Mining, Visual Analytics, Statistik

Abstract

Concerning multi-dimensional data sets there exist a lot of visual-based as well as automatical techniques to detect inherent relations and characteristics. Due to the (increasing) size and complexity of such data, it is necessary to combine both approaches. In this article, we therefore present to you established visual-based and automatical data analysis approaches and we reveal modern methods to combine these approaches, with the goal to enhance the data analysis process. All explanations are supported by examples to ease the reader's understanding.

Keywords: Visualization, Data Analysis, Data Mining, Visual Analytics, Statistic

Einleitung

Im Sommer 1854 brach eine der schlimmsten Cholera-Epidemien in London aus; nach bereits vier Ausbrüchen innerhalb von nur 23 Jahren schritt dieser so schnell fort und war so tödlich wie noch keiner zuvor. Ohne Vorankündigung starben innerhalb weniger Tage allein im Stadtteil Soho weit über hundert Menschen. Ein Gegenmittel gab es nicht. Die ÄrztInnen vermuteten, dass

gefährliche Dünste, Miasmen, für die Ausbreitung der Krankheit verantwortlich seien. Woher aber diese Dünste kommen sollten war ein Rätsel. Der einzig mögliche Schutz bestand in der Flucht aus der Stadt.

John Snow, ein Arzt im Londoner Stadtteil Soho, erkannte, dass die gängige Miasmen-Theorie keine Hilfe bot. Stattdessen hatte er die Vermutung, dass sich die Krankheit von nur wenigen Infektionsherden aus verbreitete. Sein Ziel war es, diese zu finden und zu eliminieren, um die Seuche einzudämmen. Doch wo sollte er mit der Suche nach hypothetischen Krankheitsquellen beginnen, zu einer Zeit, als weder die Existenz von bakteriellen Krankheitserregern noch das Konzept von Infektionswegen bekannt waren? Seine einzige Möglichkeit bestand darin, die Suche auf seine Beobachtungen zu stützen.

Er kam auf die Idee, die Wohnorte der Cholera-Opfer in seinem Bezirk in einer Stadtkarte einzutragen, welche berühmt wurde als "Ghost Map" (siehe Abb. 1). Durch diese Darstellung wird sichtbar, dass die Wohnorte der Cholera-Opfer nicht gleichmäßig über Soho verteilt waren, sondern dass es eine klare Häufung auf der Broad Street gab. Dort befand sich eine öffentliche Wasserpumpe, an der sich die Bewohner mit Trinkwasser versorgten. Ein Zufall? John Snow überzeugte die Stadtverwaltung, den Schwengel der Pumpe abzumontieren, was die Bewohner zwang, ihr Wasser an anderen Pumpen zu holen. Innerhalb weniger Tage ging die Opferzahl in Soho drastisch zurück. John Snow hatte mithilfe einer Visualisierung und ihrer richtigen Analyse zahlreichen Menschen das Leben gerettet.

Das Beispiel macht deutlich, dass Datenvisualisierung und -analyse keine neue Wissenschaft ist. Das Verfahren von John Snow findet auch noch heute Anwendung, z.B. in der Kriminalistik, wenn Tatorte und Beweismittel auf Landkarten miteinander in Beziehung gesetzt werden, um versteckte Muster zu erkennen. So geht es grundsätzlich darum, aus unvollständigen Informationen allein auf Grundlage der vorhandenen (beobachteten) Daten auf nützliche, sinnhafte Zusammenhänge zu schließen. Visualisierung ist damit ein Werkzeug zum phänomenologischen Verständnis von Zusammenhängen: anhand seiner "Ghost Map" konnte John Snow den Zusammenhang zwischen Cholera und Trinkwasser postulieren, ohne sich durch wissenschaftliche Grundlagen wie Bakteriologie oder Epidemiologie leiten lassen zu können.

Kernprinzip visueller Analyse ist es, nicht-zufällig erscheinende Zusammenhänge zwischen scheinbar unabhängigen Größen aufzufinden, geleitet durch Intuition und durch unser visuelles System. Dabei unterstützt der Rechner den menschlichen Suchvorgang, indem er z.B. große Datenmengen voranalysiert, auf verdächtige Nicht-Zufälligkeiten hinweist und Daten so visuell präsentiert,

*dirk@isg.cs.uni-magdeburg.de; Universität Magdeburg

†georgia@cg.cs.tu-bs.de; TU Braunschweig

‡eisemann@cg.cs.tu-bs.de; TU Braunschweig

§tatu@dbvis.inf.uni-konstanz.de; Universität Konstanz

¶keim@dbvis.inf.uni-konstanz.de; Universität Konstanz

||schumann@informatik.uni-rostock.de; Universität Rostock

**magnor@cg.cs.tu-bs.de; TU Braunschweig

††theisel@isg.cs.uni-magdeburg.de; Universität Magdeburg

dass ein Mensch etwaige Muster schnell und sicher erfassen kann.

Solche Techniken lassen sich in vielen Bereichen anwenden. Ihre volle Leistungsfähigkeit entfalten sie jedoch an multi-dimensionalen Datensätzen, in denen sich interessante Zusammenhänge verstecken können, die ohne visuelle Analysemethoden verborgen blieben. Ein Beispiel geben Versicherungsfirmer: So kostete die Kfz-Versicherung für einen roten Wagen in den USA lange Zeit deutlich mehr als für ein baugleiches weisses Auto, weil eine Analyse der Unfallstatistiken ergeben hatte, dass rote Autos häufiger verunglücken als andersfarbige Wagen. Doch haben Automobile natürlich noch andere charakterisierende Dimensionen als nur ihre Lackierung. Eine vollständige visuelle Analyse sämtlicher zugelassener Wagen könnte z.B. ergeben, dass rote Autos häufiger von Männern gefahren werden, oder dass PS-starke Motoren nur sehr selten in weissen Autos verbaut werden, oder, oder ... Um der wahren Ursache eines phänomenologischen Zusammenhangs auf den Grund zu gehen, braucht es daher zweierlei: der Suche nach allen scheinbaren Zusammenhängen in einem Datensatz (exhaustive search) sowie eines Menschen, der die gefundenen Zusammenhänge kausal verknüpfen und Scheinzusammenhänge auf ihre wahren Ursachen zurückführen kann. Das genannte historische Beispiel beschreibt einen

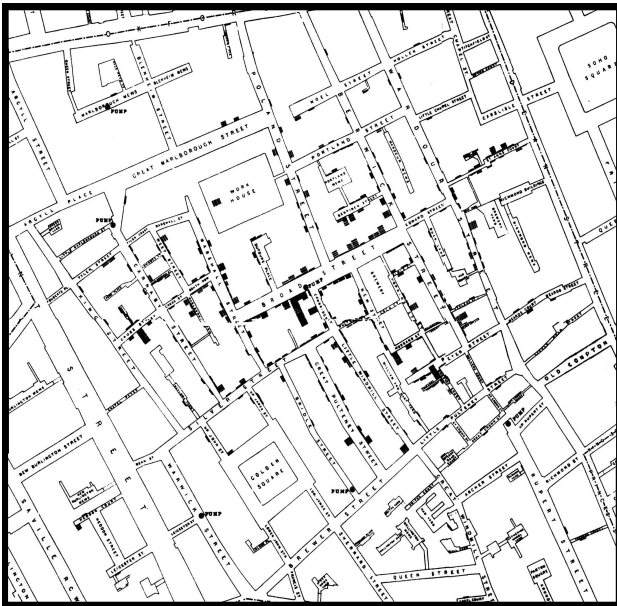


Abbildung 1: In der “Ghost Map” von John Snow sind die Wohnorte der Cholera-Opfer eingetragen; es wird deutlich, dass sich die Todesfälle um eine konkrete Wasserpumpe herum häufen.

einfachen Fall von multi-dimensionalen Daten. Ebenso einfach (und doch wirksam) ist die Wahl der Visualisierungstechnik. Die heutige Situation lässt sich dadurch beschreiben, dass die zu untersuchenden Datensätze immer größer und komplexer werden, und auf der anderen Seite eine stetig wachsende Vielzahl von automatischen und visuellen Analysemethoden zur Verfügung stehen. Eine Kombination von automatischen und visuellen Techniken ist somit notwendig und aktueller Forschungsgegenstand. In den nächsten Kapiteln geben wir eine formale Definition von multi-dimensionalen Daten, beschreiben existierende Standardtechniken zur automatischen und visuellen Datenanalyse und zeigen an Bei-

spielen einige aktuelle Arbeiten zur sinnvollen Kombination solcher Techniken.

Multi-dimensionale Datensätze

Um imstande zu sein ein (visualisiertes) Muster zu interpretieren bzw. um das Gesehene in einen Sinn-Kontext einzuordnen, ist ein allgemeines Verständnis von der Struktur zugrundeliegender multi-dimensionaler Datensätze unerlässlich. Aus diesem Grund werden sie in diesem Abschnitt eingeführt.

Ein intuitives Beispiel eines solchen Datensatzes ist das Resultat einer Messung in einem Zimmer, in welchem eine Anzahl von physikalischen Messgrößen erfasst werden, wie beispielsweise die Temperatur und der Luftdruck. Hierbei spannt das Zimmer einen dreidimensionalen Messraum und die beiden Messgrößen einen zweidimensionalen Messgrößenraum auf: $\mathbb{R}^3 \rightarrow \mathbb{R}^2$. Die Anzahl der gemessenen Paare entspricht der Mächtigkeit der Population.

Ein beliebiger Datensatz ist somit formal charakterisiert durch eine Abbildung f von s -vielen Elementen x_i ; $i = 1, \dots, s$ eines n -dimensionalen Messraumes (spatial domain) auf s -viele Elemente ξ_j ; $j = 1, \dots, s$ eines m -dimensionalen Messgrößenraumes (data domain) und entspricht aus mathematischer Sicht einer diskreten multivariaten vektorwertigen Funktion:

$$x \rightarrow f(x) = \xi : \mathbb{R}^n \rightarrow \mathbb{R}^m.$$

Somit werden den n unabhängigen Dimensionen des Messraumes m abhängige Dimensionen des Messgrößenraumes zugeordnet, wobei die Elementanzahl s die Population der Elemente darstellt.

In der Fachliteratur wird folglich zumeist zwischen den Dimensionen beider Räume unterschieden, wenn auch nicht immer einheitlich [18]. Nicht immer ist das zielführend, weil abhängige Dimensionen auch als unabhängig betrachtet werden können und umgekehrt. Zur Charakterisierung der Dimensionalität eines Datensatzes kann es stattdessen zweckdienlich sein, die Gesamtanzahl der erfassten Dimensionen als Merkmalsraum, unter dem Begriff der *Variablen* $k = n + m$, zu bündeln. Im Weiteren benutzen wir daher den Begriff der Variabel, wenn wir von einer Dimension des Datensatzes sprechen. Es ist unter anderem die Anzahl dieser Variablen, welche die Komplexität des “hervorgerufenen” Visualisierungsproblems bestimmt.

Dateneigenschaften

Ein Datensatz ist jedoch nicht nur durch seine Dimensionalität charakterisiert, sondern auch durch die konkreten Eigenschaften seiner Daten selbst [21]:

- **Ordnung** Ein Datum lässt sich als Tensor i -ter Ordnung beschreiben. Dabei entspricht ein Skalar einem Tensor nullter Ordnung, ein Vektor einem Tensor erster Ordnung, eine Matrix einem Tensor zweiter Ordnung, u.s.w. Weil aber eine Ordnung größer als Null auch durch eine größere Variablenanzahl ausgedrückt werden kann, gehen wir zumeist von skalaren Daten aus.

• Skaleneigenschaft

- **Quantitativität** Die Daten sind Zahlen eines bestimmten Wertebereiches konkreter Zahlenmengen ($\mathbb{Q}$, $\mathbb{Z}$, $\mathbb{R}$, ...).
- **Qualitativität** Unterliegen die Daten einer Ordnungsrelation, wie z.B. größer, kleiner oder gleich, wird von Ordinalität gesprochen; sind sie andererseits textuelle Bezeichner, wie z.B. eine Farbe (rot, grün, blau) oder eine Form (rund, eckig, länglich) handelt es sich um nominelle Daten. Ein nominelles Datum wird zumeist als a priori Klassifikator der Population genutzt, um ihre Elemente eindeutig einer Klasse zuzuordnen. Vorweggreifend sei darauf hingewiesen, dass die Klassenzugehörigkeit innerhalb einer Visualisierung zumeist durch eine klassenkonsistente Farbkodierung kenntlich gemacht wird.

Zusammenfassend läßt sich ein *multi-dimensionaler Datensatz* verstehen, als ein Datensatz mit mindestens zwei (oder mehr) skalaren Variablen. Um jedoch eine Vorstellung von der praktischen Arbeit zu vermitteln sei erwähnt, dass es sich dort zumeist um Datensätze mit 30, 40 oder mehr Variablen handelt, mit einer Population, die durchaus in die Millionen gehen kann.

Letztlich ist ebenfalls auch diese Mächtigkeit der Population charakterisierend für einen Datensatz, weil eine Zunahme gewöhnlich mit einer sich verschlechternden Performance¹ einhergeht. Für Datensätze die mit Standardvisualisierungsmethoden visualisierbar wären bedeutet dies, dass sie ab einer kritischen Mächtigkeit nicht mehr (vollständig) visualisierbar sind, da der zu erwartende Nutzen der Visualisierung den zeitlichen Aufwand nicht mehr rechtfertigt. Diese Problematik ist ein noch immer aktueller Gegenstand der Forschung und in seiner Gesamtheit ungelöst. Teillösungen mit GPU-basierten Ansätzen existieren jedoch schon heute. Sie haben den Vorteil zeitaufwändige Berechnungen parallel (gleichzeitig) auszuführen, anstatt seriell (nacheinander).

Standardmethoden zur Visualisierung multi-dimensioneller Daten

Bei den Methoden zur Datenvisualisierung wird zwischen der Visualisierung von physikalischen Daten (scientific visualization) und abstrakten Daten (information visualization) unterschieden: Dabei sind phys. Daten insbesondere Skalar-, Vektor- und Tensorfelder, resultierend aus Messungen oder Simulationen. Abstrakte Daten können dagegen zumeist als Listen, Bäume und Graphen beschrieben werden, wie beispielsweise die Verlinkungsstruktur zwischen beliebigen Webseiten solche Daten sind. Eine klare Abgrenzung beider Teilgebiete ist nicht immer möglich oder gar notwendig. Dennoch war die historisch bedingte Unterteilung durchaus erfolgreich: Die Fokussierung auf Teilaspekte des Visualisierungsproblems führte in wenigen Jahren (und somit

sehr schnell) zu großen Fortschritten, sowohl in der Theorie als auch in der Anwendung. Gegenwärtig jedoch sind Tendenzen ersichtlich beide Teilgebiete mehr und mehr zu konsolidieren, um synergetisch weitere Fortschritte zu forcieren. Dessen ungeachtet, gibt es beiderseits etablierte Methoden multi-dimensionale Datensätze zu visualisieren. Im Weiteren stellen wir insbesondere drei typische Beispiele zur Visualisierung multi-dimensionaler Daten vor, die in Disziplinen, wie der Systembiologie oder der Meteorologie, Anwendung finden.

Tabellen Eine intuitive und wenig aufwändige Möglichkeit ist die textuelle Darstellung der Daten als Tabellen bzw. Tables (oder auch spreadsheets genannt), wobei Spalten den Variablen und Zeilen den Daten entsprechen. Die Spaltenanzahl korrespondiert mit der Anzahl der Variablen und die Zeilenanzahl mit der Mächtigkeit der Population, wie aus Abb. 2 (a) ersichtlich ist. Obgleich diese Form der Visualisierung einen Datensatz vollständig darstellt, sind weder Zusammenhänge zwischen den Variablen, noch Häufungspunkte der Daten (cluster), ohne größeren kognitiven Aufwand, erkennbar. Zusätzlich überschreitet eine große Population oder auch eine große Variablenanzahl schnell die Darstellungsfähigkeit eines handelsüblichen Monitors.

Graphisch Komprimierte Tabellen Dem letztgenannten Nachteil begegnen graphisch komprimierte Tabellen bzw. Gaphical Compressed Tables erfolgreich, indem sie, anstatt ein Datum textuell darzustellen, eine (nur pixelbreite) qualitative Repräsentation dieser visualisieren. Dadurch kann die benötigte Monitorfläche eines Datums enorm reduziert und insgesamt die Darstellung erheblich komprimiert werden. Abb. 2 (b): Im Vergleich zu den Tables sind sowohl mehr Variablen als auch eine größere Population auf dem Monitor darstellbar. Zusätzlich werden nun auch Zusammenhänge zwischen Variablen zumindest rudimentär erkennbar, gegebenenfalls unterstützt durch spaltenbasierte Sortierungen. Falls notwendig können kontextabhängig textuelle Daten mittels einer nutzerbasierten Selektion rekonstruiert werden (table lens [20]), um derart die Vorteile beider tabellarischen Methoden zu kombinieren.

Streudiagramme Bei den Streudiagrammen bzw. Scatterplots werden die Daten zweier Variablen als Punkte in ein euklidisches Koordinatensystem eingetragen, deren Achsen die beiden Variablen repräsentieren (orthogonale Projektion). Mit ihnen lassen sich bivariate Korrelationen, Cluster, Verteilungseigenschaften sowie Kompaktheit und Streuung der Daten sehr gut analysieren, wie es die Abb. 3 (a) verdeutlicht. Aussagen über multivariate Zusammenhänge (z.B. multivariate Korrelation) sind allerdings kaum möglich, zudem gehen Informationen über die Daten-Anzahl verloren, die auf die gleiche Position im Scatterplot abgebildet werden. Transparenzen zu verwenden kann diesem Effekt bis zu einem gewissen Grad entgegenwirken. Für einen Datensatz mit k Variablen existieren genau l verschiedene Scatterplots: $l = k \frac{(k-1)}{2}$.

Um möglichst übersichtlich verschiedene Visualisierungen, wie z.B. verschiedene Scatterplots eines Datensatzes, darzustellen, bieten sich Visualisierungs-Matrizen bzw. Panel Matrizen an. Dabei handelt es sich um eine

¹Wird von Performance gesprochen, ist je nach Kontext der zeitliche Aufwand und/oder der Speicherverbrauch eines Algorithmus bzw. einer Methode gemeint.

Dim A	Dim B	Dim C	Dim E	Dim F	Dim G	Dim H	Dim I	Dim J	Dim K	Dim L	Dim M	Dim N	Dim O	Dim P
0.0041134	1.5888786	1.3442338	1.5077931	1.5312768	0.8089982	0.8651271	0.8464277	0.2519321	0.3103338	1.3088962	0.7067415	0.2304445	0.2925315	0.7021015
1.0269421	1.1767624	10.2551048	25.503675	6.6001809	5.5673454	5.5313045	1.1274152	0.7601647	10.3574251	2.7854038	1.6012657	2.7624271	2.6454511	1.4444647
2.0497211	1.5301128	15.811862	26.760754	1.7717901	0.9235121	0.924672	1.5426844	6.6353516	2.4848434	1.05879915	4.4464417	25.676328	4.3715446	1.5803002
4.0490391	0.8105668	1.3054403	12.817372	0.24351785	5.058732	0.1322249	1.8964782	11.9138565	18.7445324	6.48281942	4.4756641	6.4769386	1.4769386	1.2666707
5.0420534	0.3064205	6.5180134	17.103706	0.8417415	4.6993106	0.6753457	0.5906657	11.7142357	14.7744214	6.1320211	4.6175811	11.3100001	6.4614643	6.1144101
6.0413502	2.1958797	2.1048588	18.495681	2.0971244	2.0151202	0.1624846	2.4476585	16.043245	24.1183164	5.2113025	2.8148647	18.7807794	5.9720124	4.8801747
7.0319574	0.4676329	7.0347348	1.9178889	1.4484984	1.1715211	0.1785447	4.8139964	20.8514601	7.1912277	2.6405102	2.78444	6.7814115	1.4880075	0.8683775
8.0946644	4.4841727	17.251428	0.9072777	1.4584653	5.3051482	1.3844889	2.7764224	1.1911244	5.3744114	4.1471774	2.5047917	3.5874115	3.2495318	2.7495318
9.1629494	2.2064957	11.4877013	12.509846	0.1923172	4.2943172	0.4786331	1.8241214	10.542384	11.2573159	0.7923247	5.9323247	1.3727593	4.242931	3.2562629
10.0469611	1.2679124	4.32057267	0.6396757	0.9351401	2.2354534	0.3047159	1.0812464	0.0347991	5.2917417	1.2561545	6.8333775	22.846323	6.8847774	1.9471194
11.0152229	0.54414727	2.08064402	21.815756	0.5604232	1.0071987	0.2169919	1.8105149	15.526241	15.4633834	6.59977157	5.3308805	6.0334003	6.6247187	1.9758405
12.01120617	4.1695261	8.11740544	27.181757	1.3024125	4.9724065	0.5651968	4.9577011	16.544931	2.4123474	6.20818635	5.3911382	11.558614	1.9277809	2.0557743
13.0174141	0.7785338	5.10934022	5.0842379	0.2250336	0.6352371	0.5534269	0.5704432	16.9391349	21.9764781	6.4363445	6.4317481	10.521117	1.1987164	4.4055765
14.0055134	4.54637634	7.58875484	10.753888	1.5448995	2.70381521	0.6477099	27.5605454	0.4887931	1.8703768	1.7409975	25.09415	1.1211984	1.4747311	1.4747311
15.0394552	0.82876217	11.5426124	4.1128485	0.1104849	0.9097952	0.0244989	0.4446781	0.4446781	0.4446781	0.4446781	0.4446781	0.4446781	0.4446781	0.4446781
16.0101171	1.11874986	5.48880775	5.9055212	0.2901395	1.0699779	0.4123877	3.2325119	1.6259119	0.8999159	2.1885119	0.205841	17.5467074	1.9591718	1.9591718
17.0148634	1.39117445	2.81681986	5.7801134	0.47915228	1.9484934	0.3061375	0.4425954	2.4421257	22.81656	3.7883338	1.38858	26.861793	2.0888124	1.0035782
18.0300941	0.7883988	8.1930389	2.6313112	0.5311892	2.8991182	0.9097121	4.9550133	16.4077957	25.283774	1.4307767	3.9890329	27.520032	1.8844833	2.6689999
19.0120617	1.7425338	17.5300566	22.94506	2.18795132	0.1188071	0.7137279	0.6179862	20.177002	19.8421861	1.0617942	6.9454512	16.4053975	4.8210794	1.5712628
20.0789589	1.63574128	15.4134732	17.851441	1.0940344	6.3427781	0.4453453	0.5775511	1.8908235	12.5106171	2.0812474	1.2771035	6.1414895	4.8748143	2.688317
21.0191121	1.46858127	7.54389017	13.482485	0.5081175	1.0519979	0.0641541	1.8474236	0.834802	25.588516	6.9071777	4.7679567	6.482115	2.7592984	6.118133
22.01131484	1.55971261	17.8495791	4.1112952	1.941388	5.157895	1.2892328	0.9712625	10.17770	27.0951138	4.0803954	6.6280477	0.6625758	2.9774095	5.9965752
23.0142123	1.4527188	11.4340004	0.4100004	0.000000	0.999999	0.999999	0.999999	0.999999	0.999999	0.999999	0.999999	0.999999	0.999999	0.999999
24.01013594	1.48469384	1.12540222	8.3154496	0.0619671	0.2189791	2.188954	0.0446666	17.8766124	7.7350101	3.8716171	1.0898944	0.18881384	4.4859732	4.4859732
25.0189512	2.10358938	15.3810581	5.3479279	0.5693893	5.8467396	0.946931	0.5988277	0.7885321	12.96988	6.8730988	7.5883445	19.270371	1.6357531	0.3899482

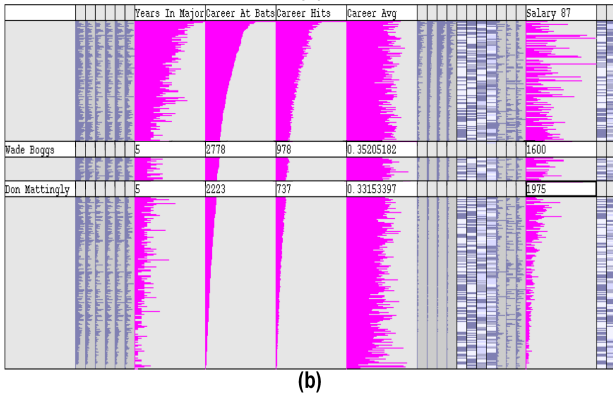


Abbildung 2: Daten Visualisierung mittels Tables und Graphical Compressed Tables: (a) Eine Table visualisiert 15 Variablen mit 25 Daten. (b) Eine Graphical Compressed Table gleicher Auflösung visualisiert mit 25 Variablen und über 300 Daten wesentlich mehr Informationen als die Table auf einmal; weiterhin sind ergänzend textuelle Darstellungen (als table lens) möglich [20].

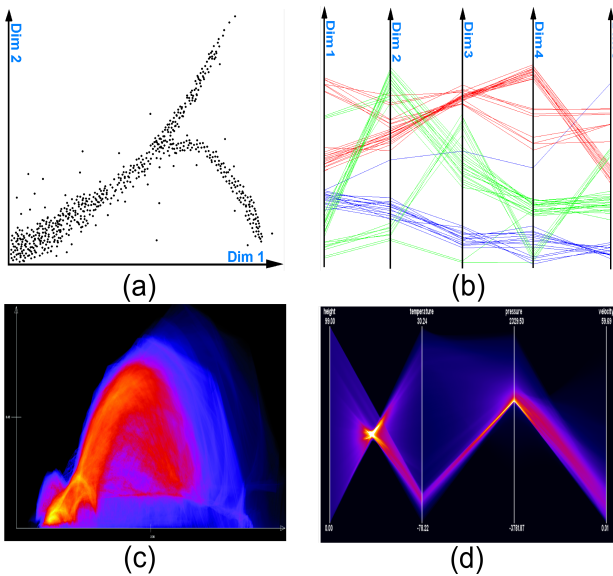


Abbildung 3: Scatterplots und Parallele Koordinaten als diskrete (a-b) und als kontinuierliche Datenvisualisierung [3, 9] (c-d).

Menge von Visualisierungen, die als rechteckiges Schema angeordnet sind. Es ergibt sich somit eine vollständige Visualisierung des Datensatzes. Die Visualisierungs-Matrizen untereinander unterscheiden sich in der Wahl der verwendeten Visualisierungsmethode.

Streudiagramm Matrizen Eine Streudiagramm Matrix bzw. Scatterplot Matrix (SPLOM) [5] eines Datensatzes mit k Variablen ist eine symmetrische $k \times k$ Matrix M , bei der die i -te Spalte und die j -te Zeile

$(0 \leq i, j \leq k-1)$ eindeutig mit Variablen assoziiert sind, und bei der das Matrix-Element der Position $M(i, j)$ ein Scatterplot ist, der die beiden Variablen i und j darstellt. Derart werde alle orthogonalen Projektionen des Datensatzes in der unteren und in der oberen Dreiecksmatrix visualisiert, wie Abb. 4 (a) aufzeigt. Der direkte Vergleich von Scatterplots unterschiedlicher Variablen unterstützt insbesondere die Hypothesenbildung multivariater Zusammenhänge in den Daten.

Parallele Koordinaten Ein Datum wird als Linienzug entlang vertikaler und zueinander paralleler Achsen repräsentiert. Jede Achse korrespondiert mit einer Variable; jeder Achsen-Schnittpunkt des Linienzuges entspricht dem Wert des Datums bezüglich dieser Variable. Somit wird der Datensatz vollständig abgebildet. Aber: für unerfahrene Nutzer sind Parallele Koordinaten [12] nur schwer zu interpretieren. Abb. 3 (b) illustriert dies. Eine Achse steht immer in direkter Verbindung mit zwei Anderen, wodurch es schwierig ist Zusammenhänge zwischen nicht direkt verbundenen Achsen “aufzuspüren”. Folglich ist die Anordnung der Achsen bedeutend, inwieweit und ob Zusammenhänge zwischen den Variablen interpretierbar sind (Anordnungsproblem). Wird zudem berücksichtigt, dass bei k Variablen $k!$ solcher Reihenfolgen existieren, ist die Problematik offensichtlich genau die relevanten Parallelen Koordinaten zu finden.

Parallele Koordinaten Matrizen Um das Anordnungsproblem von Parallelen Koordinaten zumindest teilweise zu lösen, wurde in [1] die Parallelen Koordinaten Matrix (PACOM) eingeführt. Dabei handelt es sich um eine $k \times p$ Matrix, bei der in jeder Zeile alle 3D Achsenkombinationen in Parallelen Koordinaten bezüglich einer (Haupt-)Variablen d dargestellt werden. Dieses wird über die k Spalten aller Variablen $0 \leq d \leq k-1$ fortgesetzt, wie Abb. 4 (b) illustriert. Ein Zeile kann dabei bis zu $p := (k-1)/2$ unterschiedliche Anordnungs-elemente enthalten. Auch eine PACOM ist für den Laien nur schwer interpretierbar.

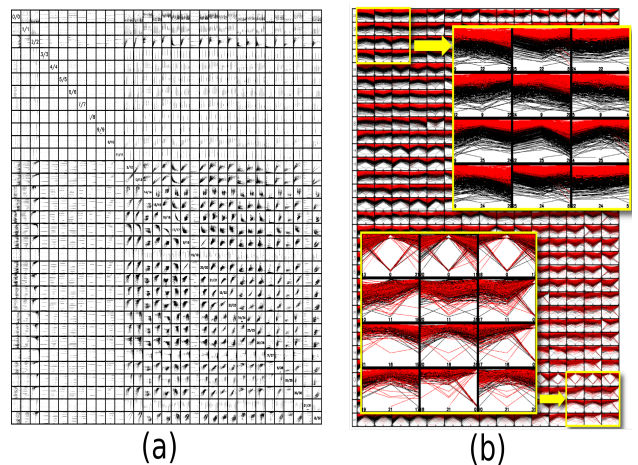


Abbildung 4: SPLOM (a) und PACOM (b) für einen Datensatz mit 32 Variablen in der Gegenüberstellung.

Sowohl für Scatterplots als auch für Parallele Koordinaten existieren zudem kontinuierliche – jedoch weniger performante – Darstellungsmethoden [3, 9] (Abb. 3 (c-d)). Sie erlauben es Lücken in den Daten zu

“überbrücken” und werden von dem Nutzer meist als intuitiver empfunden als diskrete Darstellungen.

Bezüglich der Visualisierungs-Matrizen ist zu beachten, dass sie schlecht mit zunehmender Variablenanzahl skalieren und den Nutzer daher zunehmend überfordern: Es ist kaum mehr möglich zwischen Visualisierungen mit interessanten und uninteressanten Mustern zu unterscheiden oder überhaupt alle sichten zu können.

Abschließend sei betont, dass es viele weitere Visualisierungsmethoden gibt, welche zumeist einen bestimmten Aspekt des Datensatzes besonders gut darstellen. Einige ausgewählte sind in Abb. 5 dargestellt. Dem interessierten Leser sind thematisch weiterführende Werke, wie [27], [5] oder [21] sehr zu empfehlen.

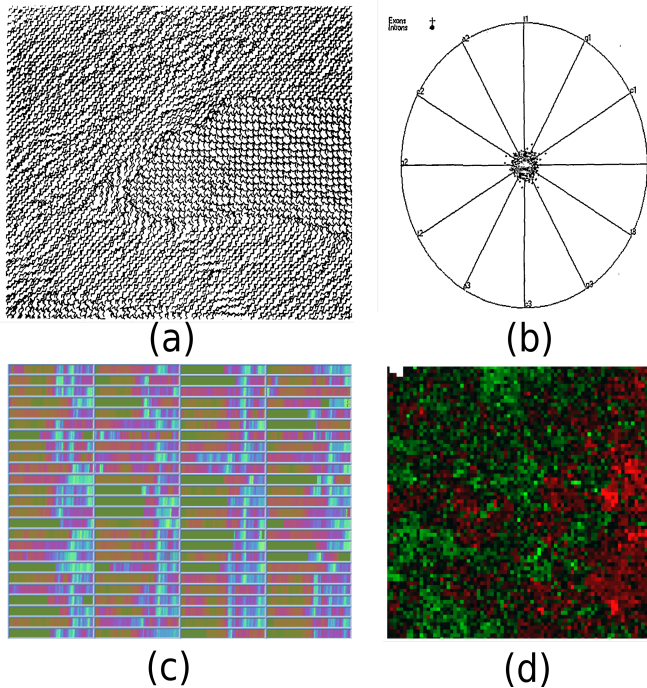


Abbildung 5: Exemplarische Visualisierungen: (a) Iconisierte Darstellung - versteckte Muster treten “zutage” [19], (b) RadViz - große Variablenanzahl im direkten Vergleich [11], (c) Recursive Pattern - sich wiederholende Strukturen werden deutlich [14], (d) Jigsaw Map - zeigt u.a. Cluster einer Variablen [26].

Standardmethoden zur automatischen Datenanalyse

Bei einer Datenvisualisierung besteht leider auch immer das Problem, dass es zu Verwirrungen bei der Interpretation kommen kann (visual clutter): Durch die begrenzte Fläche des Monitors beispielsweise, oder wenn Strukturen, die im Merkmalsraum getrennt sind, sich in der Visualisierung überlappen. Andererseits nehmen aber auch die Anzahl der Dimensionen und das Datenvolumen beständig zu. Es besteht somit ein Bedarf Daten automatisch zu analysieren:

Ziel ist es u.a. Visualisierungs-Methoden zu falsifizieren oder eine multivariate statistische Datenanalyse zu erhalten. Letzteres ist eine erste Annäherung an eine vollständige explorative Analyse. Ziel ist konsequenterweise auch, interessante Teilmengen zu finden, deren Visualisierung sich “lohnt”. Eine vollständige und kontextspezifische Dateninterpretation kann aber auch das

beste automatische Verfahren nicht leisten!

Methoden zur Strukturidentifikation

Im Folgenden werden zwei prominente Datensatz-Strukturen näherer erläutert.

Korrelation Eine Korrelation ist eine (mathematisch-statistische) Beziehung zwischen mehreren Variablen. Die Stärke dieser, wird durch den mittleren quadratischen Fehler (mean square error, kurz MSE) zwischen einer Funktion und den Datenwerten selbst beschrieben. Je kleiner der MSE, umso stärker ist die Korrelation, die durch diese Funktion Ausdruck findet. Da praktisch nicht ersichtlich ist, welche Funktion einen optimalen MSE liefert, wird meist für eine Anzahl (multivariater) Polynome steigenden Grades der minimale MSE berechnet (durch Wahl geeigneter Koeffizienten). Die Funktion mit einem absoluten MSE kleiner einer bestimmten Schwelle wird als Korrelation propagiert, ansonsten gilt, dass keine Korrelation vorliegt. Dieses Vorgehen wird als Regressionsverfahren bezeichnet.

Cluster Beim Auffinden von Clustern (clustering) werden ähnliche Objekte einer gemeinsamen Gruppe (=Cluster) zugeordnet, mit Hilfe von z.T. komplexen Ähnlichkeitsfunktionen. Unterschiedliche Studien haben das Verhalten von Ähnlichkeitsfunktionen im Multidimensionalen analysiert [4, 10]: Sie beschreiben, dass die Distanz des (metrisch) entferntesten Objektes – bezüglich eines Anfrageobjektes – mit steigender Dimensionalität, nicht so schnell zunimmt, wie die Distanz zum (metrisch) nächsten Nachbarobjekt (Fluch der Dimensionalität bzw. curse of dimensionality):

$$\lim_{d \rightarrow \infty} \frac{dist_{max} - dist_{min}}{dist_{min}} = 0.$$

Dies bedeutet, dass die Unterscheidung zwischen dem Nächsten und dem entferntesten Objekt an Bedeutung verliert: Ergebnisse des Clusterings werden somit kontinuierlich schlechter, mit Zunahme der Variablenanzahl. Weiterhin wissen wir, dass Cluster zumeist nur in Unterräumen (subspaces) der Variablen auftreten, was umso wahrscheinlicher ist, je mehr Variablen der Datensatz hat. Daher werden Cluster zumeist nicht mehr global, sondern in lokalen Unterräumen gesucht (*Subspace Clustering*).

Sowohl für Korrelation als auch für Cluster gilt: Sie sind nicht immer eindeutig und die Ergebnisse variieren mit den eingesetzten Verfahren.

Transformation des Merkmalsraumes

In Disziplinen wie der Logistik oder der Bioinformatik umfassen die Daten z.T. hunderte von Variablen. Es ist daher von großem Interesse, Merkmalsräume mit weniger Variablen zu finden, die geeignet sind um möglichst strukturerhaltend eine Transformation der original Daten zu ermöglichen. Üblicherweise wird dabei zwischen dimensionsreduzierenden und dimensionsselektierenden Techniken unterschieden:

Principal Component Analysis Die Principal Component Analysis (PCA) [6] projiziert den Merkmalsraum in Einen, der den größten Teil der Varianz der Daten

enthält. Dabei werden Variablen (Hauptkomponenten bzw. principal components), welche den neuen Merkmalsraum aufspannen, durch die Analyse von Eigenvektoren berechnet.

Multi-dimensional Scaling Unter multi-dimensional Scaling (MDS) [16] wird ein nicht-linearer iterativer Algorithmus verstanden, welcher Daten in ein Verhältnis zu einer Metrik setzt. Derart resultieren Ähnlichkeiten im neuen Merkmalsraum als Cluster, die wiederum mittels Clusteranalyse detektiert werden können. Entgegen der PCA wirkt das MDS bereits als Strukturfilter, gesteuert durch eine entsprechende Wahl der Metrik.

Self-organizing Maps Eine Self-organizing Map (SOM, auch Kohonennetz) [15] ist eine unbeaufsichtigte Lernmethode um den Merkmalsraum auf einen Raum geringerer Dimensionalität zu reduzieren. Sie ist den Methoden der neuronalen Netze zuzuordnen.

Nachteil all dieser Techniken ist, dass die neu generierten Variablen mit ihren Originalen zumeist in nicht linearer Weise assoziiert sind. Somit hat der von ihnen generierte Raum nicht immer eine klar erkennbare Bedeutung für den Nutzer.

Als weiterführende Literatur im thematisch näheren Umfeld seien [7], [17], [8], [22] und [2] genannt. Etwas allgemeiner ist eine Vertiefung in die Disziplinen der multivariaten Statistik, des Data Mining und des Machine Learning sehr zu empfehlen.

Kombination von Methoden der Datenvisualisierung und Datenanalyse: Beispiele und Chancen

Wir haben bisher Methoden aufgezeigt um multi-dimensionale Daten zu visualisieren oder automatisch zu verarbeiten bzw. zu analysieren. Immer einhergehend mit der Problematik einer großen Anzahl von Visualisierungen, die den Nutzer schlicht überfordern; oder automatischen Methoden, die nicht geeignet sind den Kontext mit zu berücksichtigen.

Eine Möglichkeit dieses Dilemma aufzulösen besteht in der zielgerichteten Kombination beider Methoden. Wie ist das möglich? Zum Einen können automatische Methoden helfen geeignete Visualisierungsmethoden auszuwählen; zum Anderen können sie genutzt werden um geeignete Visualisierungen als Ausgangspunkt für eine Mustersuche zu finden. Die letztere Möglichkeit stellen wir exemplarisch in diesem Abschnitt vor.

Es ist dabei das Ziel, automatisch, bestimmte Visualisierungen aus der Gesamtheit aller zu ermitteln, wie z.B. in [13]: Insbesondere solche, welche ein Visualisierungsziel (Korrelation, Cluster, Assoziation, etc.) vermeintlich am besten darstellen. Die Methoden, die im Bildraum der Visualisierungen selbst operieren, werden als *Quality Measures* bezeichnet

Fünf Quality Measures, am Beispiel der Scatterplots und der Parallelen Koordinaten multi-dimensionaler Daten, führen wir nun auf, die bezüglich der Korrelation, Klassensepariertheit und der Cluster hin bewerten.

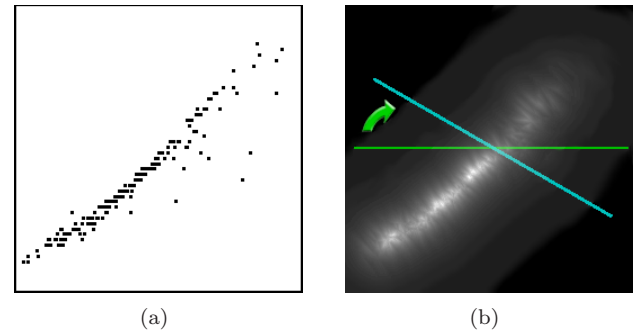


Abbildung 6: Scatterplot Beispiel mit Dichtefeld: Für jedes Pixel wird die Masseverteilung entlang verschiedener Richtungen – hier als blaue Linie dargestellt – berechnet und jeweils der minimalste Wert wird gespeichert.

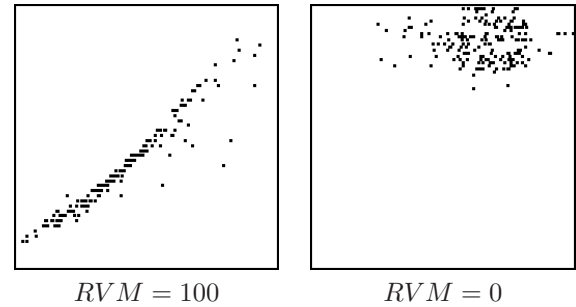


Abbildung 7: Bewertung von Scatterplots bezüglich der Korrelation seiner Variablen: Ein hoher RVM entspricht einem Scatterplot mit stark korrelierenden Variablen, ein niedriger RVM-Wert hingegen deuten auf schwach korrelierende Variablen hin.

Rotating Variance Measure (RVM) [24] ist ein Maß, um lineare und nicht-lineare Korrelationen in Scatterplots zu bewerten. Um das RVM zu berechnen, wird zunächst ein kontinuierliches Dichtefeld aus dem Scatterplot ermittelt. Für ein Pixel $\mathbf{p}$ der Position $\mathbf{x} = (x, y)$ wird die maximale Distanz r zum nächsten Punkt im Scatterplot berechnet, zudem die lokale Dichte $\rho = 1/r$. Dieser Schritt ist für weitere Berechnungen essentiell und schließt Ausreißer aus der Bewertung aus. Stark korrelierende Dichtefelder zeigen in der Regel eine auffällige schmale, längliche Struktur mit hohen Dichtewerten, während sonst viele verteilte lokale Maxima im Dichtefeld zu erkennen sind. Um diese Verteilung zu messen, wird die Massenverteilung entlang verschiedener Messrichtungen um das Pixel $\mathbf{p}$ berechnet (Abb. 6). Der beste Wert jeder Bildspalte und Richtung wird als Referenz für das RVM verwendet (Gleichung 1); je größer, desto besser ist die Korrelation, wie aus Abb. 7 ersichtlich ist:

$$\text{RVM} = \frac{1}{\sum_x \min_y \nu(x, y)}, \quad (1)$$

mit der Massenverteilung $\nu(x, y)$.

Hough Space Measure (HSM) [24] ist ein Maß um Parallele Koordinaten auf Cluster hin zu bewerten. Ein Cluster im Raum der Parallelen Koordinaten kann als eine Häufung von Geraden mit ähnlicher Lage definiert werden. Unter Verwendung dieser Transformation erhalten wir für jeden nicht-Hintergrund Pixel eine sinusförmige Kurve in einer 2D Ebene, dem sogenannten Hough- oder Akkumulatorraum. Ein Schnitt dieser Kurven deutet darauf hin, dass die zugehörigen Pixel auf einer Geraden im Bildraum liegen. Abb. 8 zeigt zwei Bei-

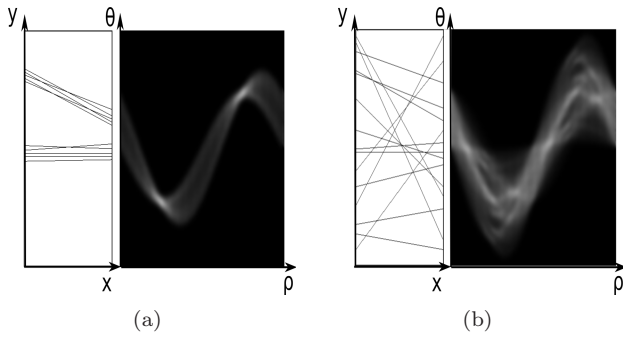


Abbildung 8: Beispiele von Parallelen Koordinaten und ihren korrespondierenden Houghräumen: (a) enthält zwei wohldefinierte Cluster von Geraden und ist für die Cluster-Erkennung besser geeignet als (b), die keine Cluster enthält.

spiele vor und nach einer Hough-Transformation. Abb. 8(a) enthält zwei wohldefinierte Geraden-Cluster und ist für die Cluster-Erkennung besser geeignet als Abb. 8(b), die keine Cluster enthält. Die hellen Bereiche der Ebene stellen hier Cluster von Geraden mit ähnlichen Parametern dar. Der Akkumulatorraum ist aufgeteilt in $w \times h$ Zellen. Eine "gute" Visualisierung enthält wohldefinierte Cluster, wenn es Zellen mit hohen Werten im Houghraum gibt. Um solche Zellen zu erkennen, berechnen wir den Median m als Schwellwert, der die Akkumulatorfunktion $h(x)$ in zwei identische Teile teilt:

$$\frac{\sum h(x)}{2} = \sum g(x) \quad , \quad \text{mit} \quad g(x) = \begin{cases} x & \text{wenn } x \leq m; \\ m & \text{sonst.} \end{cases}$$

Das endgültige Maß wird über die Menge der Akkumulatorzellen, die einen höheren Wert als m haben, berechnet.

$$\text{HSM}_{i,j} = 1 - \frac{n_{\text{cells}}}{wh},$$

wobei i, j den Indizes der jeweiligen Variablen entsprechen. Der errechnete HSM-Wert ist hoch, für Bilder die wohldefinierte Geraden-Cluster enthalten und niedrig für Bilder, die keine Cluster enthalten.

Class Density Measure (CDM) [24] ist ein Maß um die Klassenseparierung von Scatterplots zu messen. Klassen sind durch eine konsistente Farbkodierung kenntlich gemacht. Bei einer gegebenen Menge an Scatterplots eines Datensatzes gilt es die Plots zu selektieren, welche die Klasse am besten separieren. Durch die Farbkodierung können die Klassen sehr leicht in individuelle Bilder aufgetrennt werden. Es werden nun, zum RVM analoge, Dichtefelder benutzt um die gegenseitige Überlappung zwischen den Klassen zu berechnen. Die Überlappung ist die Summe der absoluten Differenz der Dichtefelder aller paarweisen Kombinationen der Klassen:

$$\text{CDM} = \sum_{k=1}^{M-1} \sum_{l=k+1}^M \sum_{i=1}^P \|\mathbf{p}_k^i - \mathbf{p}_l^i\| \quad ,$$

wobei M der Menge der Dichtefelder, $\mathbf{p}_k^i$ dem i -tem Pixel im k -tem Dichtefeld und P der Menge an Pixeln entsprechen. Abb. 9 zeigt ein Beispiel mit den am besten und den am schlechtesten bewerteten Scatterplots eines Datensatzes.

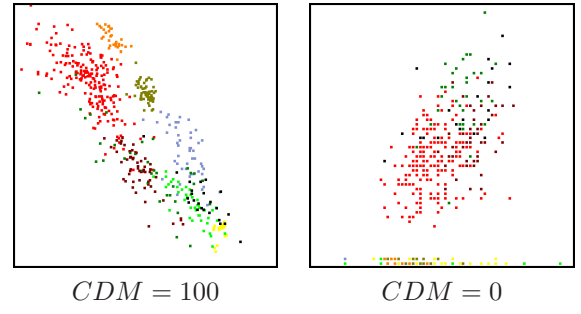


Abbildung 9: Bewertung von Scatterplots mit (farbkodierten) Klassen. Ein hoher CDM entspricht einem Scatterplot mit gut separierter Klassendarstellung, ein niedriger CDM hingegen deutet auf eine starke Überlappung zwischen den Klassen hin.

Distance Consistency Measure (DSC) [23] Jeder Datenwert $x_i; i = \{1, \dots, s\}$ eines Scatterplots erhält eine Marke, die "true" ist, wenn der Abstand zwischen x_i und seinem Klassenzentroiden $c_o(x_i)$ kleiner ist als der Abstand zu allen anderen Klassen-Zentroiden. Ansonsten ist die Marke "false". Ein Klassenzentroid meint den Schwerpunkt aller Werte die zu einer Klasse c gehören. Das DSC ist nun der Anteil an Marken mit der Belegung "true", bezüglich aller s Datenwerte:

$$\text{DSC} = \frac{|x : \text{Marke}(x, c_o(x)) = \text{true}|}{s}.$$

Je größer das DSC, desto besser sind Klassen voneinander separiert, wie Abb. 10 (a-b) aufzeigt. Es eignet sich insbesondere für kompakte Klassen.

Distribution Consistency Measure (DC) [23] In einer ε -Umgebung jedes Datenwertes x_i werden die Anzahl $p_c(x_i)$ der Werte gleicher Klasse c gezählt. Die Entropie

$$H(x_i)_c = - \sum \frac{p_c}{\sum p_c} \log_2 \frac{p_c}{\sum p_c}$$

beschreibt nun die "Dichte" der Klasse c innerhalb dieser Umgebung. Nach dem Aufsummieren dieses Maßes nach Gleichung 2 kann eine globale Aussage über das Verteilungs- und Separierungsverhalten der Klassen getroffen werden:

$$\text{DC} = 100 - \frac{1}{Z} \sum_{i=1}^s \sum_c p_c \underbrace{\left(- \sum \frac{p_c}{\sum p_c} \log_2 \frac{p_c}{\sum p_c} \right)}_{H(x_i)_c} \quad (2)$$

mit der Normierung $\frac{1}{Z} = \frac{100}{\log_2(k) \sum x_i \sum_c p_c}$. Je größer das $\text{DC} \in \{0, \dots, 100\}$, desto besser sind die Klassen separiert; wobei das Maß in diesem Fall sehr gut für nicht konvexe Klassenverteilungen geeignet ist, wie aus Abb. 10 (c-d) ersichtlich.

Quality Measures zeigen erstmals das Potential auf, welches die Kombination von Teildisziplinen für die Visualisierung und Datenanalyse bietet und geben damit die Richtung zukünftiger Forschungen vor.

Ausblick

Die hier vorgestellten Ansätze beschreiben selbstverständlich nur einen kleinen Ausschnitt der Methoden zur visuellen Analyse multi-dimensionaler Datensätze.

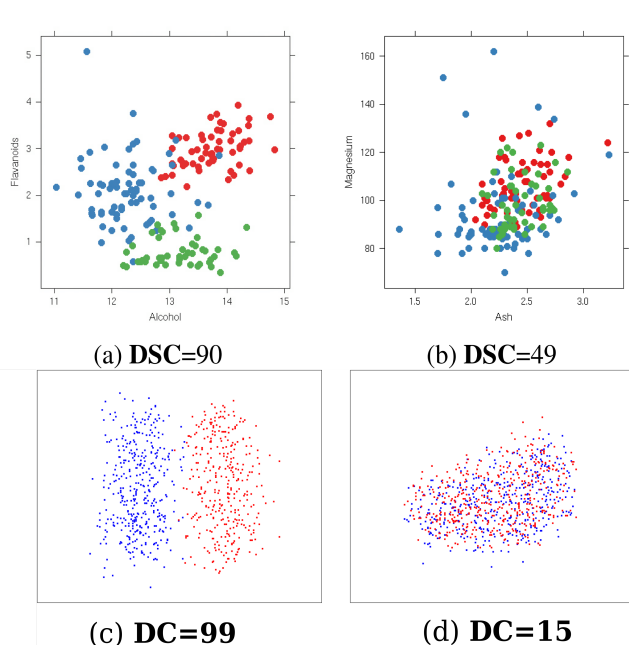


Abbildung 10: Bewertung von Separation farbkodierter Klassen (rot, blau, grün) in Scatterplots nach [23]: Ein hoher DSC/DC entspricht einem Scatterplot mit gut separierter Klassendarstellung (a und c); niedrigere DSC/DC Werte hingegen deuten auf eine schlechte Separierung hin (b und d).

Nicht erwähnt wurde die Einbeziehung applikationsspezifischen Wissens, welche zu spezialisierten Ansätzen z.B. für medizinische oder biologische Daten führen (siehe weitere Artikel in diesem Heft). Auch nicht diskutiert werden konnten Fragen der Performance, speziell die Frage, welche Möglichkeiten die rasante Entwicklung der Graphik-Hardware bietet. Ebenso ergeben sich spezielle Fragestellungen, wenn die Zeitabhängigkeit der Daten explizit untersucht wird. Die eigentliche Stärke visueller Datenanalysemethoden zeigt sich allerdings erst an interaktiven Softwaresystemen, bei denen unterschiedlichste Visualisierungen, Analysemethoden, Selektions- und Interaktionstechniken durch den Nutzer beliebig kombiniert eingesetzt werden können, um einen Datensatz (idealerweise) in Echtzeit zu explorieren und zu analysieren. Geprägt wurde hierfür u.a. in [25] der Begriff Visual Analytics, welches z.Zt. im Umfeld der Visualisierung und des Mensch-Computer-Interfaces eines der größten und mit am stärksten wachsenden Forschungsfelder darstellt: An deren Ende steht eine ferne Vision von einem System, das in der Lage ist alle interessanten Visualisierungen für jedes beliebige Visualisierungsziel eines beliebigen Datensatzes on demand liefern zu können.

Literatur

- [1] G. Albuquerque, M. Eisemann, D. J. Lehmann, H. Theisel, and M. Magnor. Quality-based visualization matrices. *Proceedings of Vision, Modeling, and Visualization (VMV 2009)*, 2009.
- [2] D. Asimov. The grand tour: a tool for viewing multidimensional data. *Journal on Scientific and Statistical Computing*, 6(1):128–143, 1985.
- [3] S. Bachthaler and D. Weiskopf. Continuous Scatterplots. *IEEE Transactions on Visualization and Computer Graphics*, 14(28):1428–1435, 2008.
- [4] K. S. Beyer, J. Goldstein, R. Ramakrishnan, and U. Shaft. When is "nearest neighbor" meaningful? In *ICDT '99: Proceedings of the 7th International Conference on Database Theory*, pages 217–235, London, UK, 1999. Springer-Verlag ISBN 3-540-65452-6.
- [5] W. S. Cleveland. *Visualizing Data*. Hobart Press, Summit, New Jersey ISBN 0-9634884-0-6, 1993.
- [6] B. S. Everitt and G. Dunn. *Applied Multivariate Data Analysis*. Arnold, 1991.
- [7] M. A. Fisher, J. H. Friedman, and J. W. Tukey. *Prim-9: An interactive multi-dimensional data display and analysis system*, volume In W. S. Cleveland, editor. Chapman and Hall, New York, 1987.
- [8] J. H. Friedman. Exploratory projection pursuit. *Journal of the American Statistical Association*, (82):249 – 266, Mar. 1987.
- [9] J. Heinrich and D. Weiskopf. Continuous Parallel Coordinates. *IEEE Transactions on Visualization and Computer Graphics (Proceedings Visualization / Information Visualization 2009)*, 15(6), 2009.
- [10] A. Hinneburg, C. C. Aggarwal, and D. A. Keim. What is the nearest neighbor in high dimensional spaces? In *VLDB '00: Proceedings of the 26th International Conference on Very Large Data Bases*, pages 506–515, San Francisco, CA, USA, 2000. Morgan Kaufmann Publishers Inc.
- [11] P. Hoffman, G. Grinstein, K. Marx, I. Grosse, and E. Stanley. Dna visual and analytic data mining. *Proceedings of the 8th conference on Visualization*, page 437 ff, 1997.
- [12] A. Inselberg. *Parallel Coordinates*. Springer Verlag Berlin, ISBN 0387215077, 2009.
- [13] S. Johansson and J. Johansson. Interactive dimensionality reduction through user-defined combinations of quality metrics. *IEEE Transactions on Visualization and Computer Graphics*, 15(6):993–1000, 2009.
- [14] D. Keim, M. Ankerst, and H. Kriegel. Recursive pattern: A technique for visualizing very large amounts of data. *Proc. Visualization 1995 IEEE Computer Society Press*, pages 279–287, 1995.
- [15] T. Kohonen. *Self Organizing Maps*. Springer Verlag, 1995.
- [16] A. Mead. *Review of the development of multidimensional scaling methods*, volume 33. The Statistician, 1992.
- [17] D. S. Moore and G. P. McCabe. *Introduction to the Practice of Statistics*. W. H. Freeman & Co., New York, NY, USA, 1999.
- [18] T. Nocke. *Visuelles Data Mining und Visualisierungsdesign für die Klimaforschung*. Dissertation, Universität Rostock, Fakultät für Informatik und Elektrotechnik, 2007.
- [19] R. M. Picket and G. Grinstein. Iconographics displays for visualizing multidimensional data. *Proc. IEEE Conference on Systems, Man and Cybernetics*, pages 514–519, 1988.
- [20] R. Rao and S. K. Card. The table lens: Merging graphical and symbolic representations in an interactive focus+context visualization for tabular information. In *Proceedings of the ACM SIGCHI Conference on Human Factors in Computing Systems*, 1994.
- [21] H. Schumann and W. Müller. *Visualisierung: Grundlagen und allgemeine Methoden*. Springer Verlag, ISBN 3-540-64944-1, 2000.
- [22] J. Seo and B. Shneiderman. A rank-by-feature framework for interactive exploration of multi-dimensional data. *Information visualization*, 4(2):96–113, 2005.
- [23] M. Sips, B. Neubert, J. P. Lewis, and P. Hanrahan. Selecting good views of high-dimensional data using class consistency. *Computer Graphics Forum (Proc. EuroVis 2009)*, 28(3):831–838, 2009.
- [24] A. Tatu, G. Albuquerque, M. Eisemann, J. Schneidewind, H. Theisel, M. Magnor, and D. Keim. Combining automated analysis and visualization techniques for effective exploration of high dimensional data. *IEEE Symposium on Visual Analytics Science and Technology*, pages 59–66, 2009.
- [25] J.J. Thomas and K.A. Cook. A Visual Analytics Agenda. *IEEE Computer Graphics and Applications*, 10-13, 2006.
- [26] M. Wattenberg. A note on space-filling visualizations and space-filling curves. *Proc. of the 2005 IEEE Symposium on Information Visualization*, 2005.
- [27] P. C. Wong and R. D. Bergeron. 30 years of multidimensional multivariate visualization. *IEEE Computer Society Press*, pages 37–44, 1992.