

Fraktale – Chaos – Künstliche Intelligenz

Christian Posthoff, Bernd Steinbach

12. August 2024

Inhaltsverzeichnis

Stichwortverzeichnis	iv
Vorwort	1
1 Einleitung	3
1.1 Felix Hausdorff	3
1.2 Benoît Mandelbrot	8
1.3 Selbstähnlichkeit	11
1.4 Fraktale auf dem Schachbrett	14
1.5 Affine Abbildungen	17
2 Grundlagen	21
2.1 Vorläufer	21
2.2 Die Cantor-Menge	26
2.3 Die Hausdorff - Dimension	28
2.4 Der Feigenbaum-Attraktor	29
2.5 Die Länge der Koch - Kurve	30
2.6 Sierpinski - Strukturen	32
2.7 Von Flächen und Körpern abgeleitete Fraktale	34
2.8 Perkolationen	36
2.9 Das Wachstum von Fraktalen	42
2.10 Brownsche Bewegungen	47
2.11 Risse und Bruchstellen	50
2.12 Fraktale Dimensionen	53
3 Fraktale in der Medizin	56
3.1 Enzephalographie	56
3.2 Das Herz	60
3.3 Das Auge	62
3.4 Einige Schlussfolgerungen	63
3.5 Die Lunge	65
3.6 Krebsdiagnose	67
3.7 Die Haut	69
3.8 Fraktale in der Anatomie	72

3.9	Epidemien	73
3.10	Digitalisierung und Künstliche Intelligenz	91
4	Fraktale Strukturen in der menschlichen Gesellschaft	96
4.1	Menschengruppen und ihr Verhalten	96
4.2	Das Kommunikationsverhalten von Gruppen	99
5	Fraktale in der Natur	104
5.1	Fraktale in der unbelebten Natur	104
5.2	Fraktale in der Pflanzenwelt	108
5.3	Fraktale in der Tierwelt	109
6	Fraktale Kunst	115
6.1	Fraktale Musik	115
6.2	Fraktale Bilder	122
6.3	Fraktale in der Literatur	122
7	Fraktale in Wirtschaft und Technik	125
7.1	Die Fraktale Fabrik	125
7.2	Fraktale in der Finanzwelt	128
7.3	Die Programmierung von Fraktalen - ihr Nutzen für die Ausbildung	139
8	Chaostheorie	141
8.1	Nichtlineare Systeme	143
8.2	Die Erzeugung von Chaos	151
8.3	Der Schmetterlingseffekt	153
8.4	Chaos in der Natur	169
8.5	Attraktion	180
8.6	Bifurkation	187
9	Zellulare Automaten	193
9.1	Conway's Game of Life	197
9.2	Dreidimensionale Fraktale	208

Abbildungsverzeichnis

1.1	Felix Hausdorff (1868 - 1942)	4
1.2	Das Grab von Felix Hausdorff auf dem Poppelsdorfer Friedhof in Bonn	7
1.3	Benoît Mandelbrot (1868 - 1942)	8
1.4	Die Mandelbrot - Menge	9
1.5	Zur Winterolympiade 2022 erhielt jedes Land seine eigene Schneeflocke	9
1.6	Eine Variante einer Kochschen Schneeflocke	10
1.7	Ein fraktales Kunstwerk	11
1.8	Spirale	13
1.9	Alle Lösungen für vier Damen auf einem Schachbrett 4×4	15
1.10	Eine Lösung für fünf Damen auf einem Brett 5×5	15
1.11	Eine Lösung des Dameproblems auf einem Brett 25×25	16
1.12	Verschiebung und Drehung	19
1.13	Konstruktion von Hugues Vermeiren	20
2.1	Eine zweidimensionale Hilbert - Kurve	21
2.2	Die optimale Gestaltung eines Warteraumes	22
2.3	Eine Peano-Kurve	22
2.4	Ein Menger - Schwamm	23
2.5	Die ersten Schritte der Konstruktion des Pythagoras-Baums	23
2.6	Ein Pythagoras - Baum	23
2.7	Newton - Fraktal zur Funktion $z^3 - 2z + 2$	25
2.8	Eine dreidimensionale Darstellung von Covid-19 - Viren	25
2.9	Cantorscher Staub	26
2.10	Ein Sandsturm - Cantorstaub in der Realität	27
2.11	Der Saturn und seine Ringe	28
2.12	Ein Feigenbaum - Diagramm	30
2.13	Die Konstruktion der Kochkurve	31
2.14	Ein Parkett aus Kochschen Flocken	32
2.15	Ein Sierpinski - Dreieck	33
2.16	Ein Sierpinski-Teppich	33
2.17	Dreidimensionale Flocken	34
2.18	Eine dreidimensionale Verallgemeinerung des Sierpinski-Dreiecks	34
2.19	Fraktale Kreise und Kugeln	35
2.20	Ineinander geschachtelte gekräuselte Kugeloberflächen	36

2.21	Die fraktale Struktur des Gehirns	36
2.22	Der Rössler - Attraktor	37
2.23	Ein Perkolations-Fraktal	37
2.24	Waldbrände in Sibirien	39
2.25	Eine interstellare Gaswolke	40
2.26	Verschmutzung einer Umgebung durch ausgelaufenes Öl	41
2.27	Der Jerusalem-Würfel	41
2.28	Drachenkurve	42
2.29	Ein Eden-Fraktal	43
2.30	Das Wachstum von Krebszellen	45
2.31	Der größte Stalaktit der Welt	45
2.32	Die Feengrotten bei Saalfeld	46
2.33	Milch auf kaltem Wasser	47
2.34	Verklumpung im Inneren eines Körpers	49
2.35	Der Fettberg in einem Londoner Abwasserkanal	49
2.36	Eine Lichtenberg - Figur	50
2.37	Ladungen zwischen zwei Leitern	50
2.38	Eine zerbrochene Eisdecke	51
2.39	Ein Riss in einer alten Mauer	53
2.40	Ein Riss in einem Autoreifen	53
2.41	Risse in einer Straße	54
2.42	Der Zusammenstoß von Kontinentalplatten	54
2.43	Die Überdeckung eines Punktes	54
3.1	Aufzeichnungen eines EEG-Gerätes	57
3.2	Die fraktale Struktur des menschlichen Herzens	61
3.3	Netzhautgefäße - fraktal strukturiert	62
3.4	Ein gesunder Lungenflügel und sein Aussehen nach zehnjährigem Rauchen	65
3.5	Die fraktale Struktur der Blutgefäße in der Lunge	67
3.6	Ein weit fortgeschrittener Tumor	69
3.7	Auch die Haut besitzt eine Fraktale Struktur	70
3.8	Die fraktale Struktur der DNA	71
3.9	Modell eines Fußes (für die Anatomie-Ausbildung)	73
3.10	Eine an der Pest erkrankte Frau	75
3.11	Eine mittelalterliche Gesichtsmaske zum Schutz gegen Infektionen .	80
3.12	Eine Ebola-Behandlungsstation in Westafrika	81
3.13	Eine an Pocken erkrankte indianische Frau	84
3.14	Opfer der Cocoliztli-Epidemie in Mexico	86
3.15	Typhus-Erreger	86
3.16	Der Erreger der Schlafkrankheit	89

4.1	Panik am Ausgang eines Fußball-Stadions	97
5.1	Fraktale Berge	105
5.2	Das Delta der Lena	106
5.3	Die Küste von England	107
5.4	Robert Brown (1773 - 1858)	108
5.5	Pilze an einem Baumstumpf	108
5.6	Ein gut entwickelter Kaktus	109
5.7	Korallen	113
5.8	Die Tarnung wird oft in die fraktale Struktur einbezogen	114
6.1	Die Begleitung eines Volksliedes durch eine Trommel	118
6.2	Das Lorenz - Fraktal	119
6.3	Schwingungen von Musikinstrumenten	120
6.4	Fraktales Kunstwerk von Sebastian Baumer	122
6.5	Ein fraktaler Teppich	123
6.6	Fraktale Keramik	123
7.1	Struktur einer fraktalen Fabrik	128
7.2	Der DAX - Kurs im Verlaufe eines Jahres	128
7.3	Der DAX - Kurs im Verlaufe von 10 Jahren	129
7.4	Der DAX - Kurs seit 1988	129
7.5	Das Kerzendiagramm für einen Tag	130
7.6	Die Darstellung einer Aufwärtsbewegung	131
7.7	Die Darstellung einer Abwärtsbewegung	132
7.8	Eine Stauchungszone	133
7.9	Widerstandszonen	134
7.10	Eine Morning-Star-Formation	135
7.11	Eine Evening-Star- Formation	136
7.12	Der Wechselkurs für Euro - Dollar	137
8.1	Ein Vierfarbengitter	151
8.2	Kombinationen von blauen Punkten	152
8.3	Kombination von gblen Punkten	152
8.4	Der Schmetterlingseffekt	155
8.5	Das Dreikörperproblem	156
8.6	Eine Gnu-Herde in Tansania	157
8.7	Ameisen auf einer Melone	168
8.8	Ein Heuschreckenschwarm	169
8.9	Ein Autostau	169
8.10	Die Folgen eines schweren Erdbebens in Haiti	170
8.11	Das Auge eines Orkans	171
8.12	Ein Tornado in Oklahoma	175

8.13	Eine Überschwemmung in Japan	178
8.14	Kleine Strudel sind hübsch und harmlos	178
8.15	Turbulenzen sind selten gefährlich	179
8.16	Darstellung der Pendelbewegung	181
8.17	Ein quadratischer Hénon-Attraktor	183
8.18	Ein De Jong - Attraktor	183
8.19	Ein Schwarzes Loch	185
8.20	Darstellung eines Torus-Attraktors	186
8.21	Die Stimmgabel-Bifurkation	190
8.22	Die Sattel-Bifurkation	190
8.23	Die Hopf-Bifurkation	191
9.1	Die Moore-Nachbarschaft	194
9.2	Die Von Neumann-Nachbarschaft	195
9.3	Eine Einband-Turing-Maschine	196
9.4	Regeln von Conway	199
9.5	Stabile Zellsysteme	201
9.6	Oszillierende Objekte	201
9.7	Das Bild einer Ameisenstraße	203
9.8	Langtons Ameisenstraße	204
9.9	Eine einzelne stabile Welle (Soliton)	205
9.10	Ein Tsunami trifft auf die Küste	205
9.11	Wellen, die sich überlagern, erzeugen ein absolutes Chaos	206
9.12	Eine zweidimensionale apollonische Figur	209
9.13	Ein Schmuckstück, das das apollonische Prinzip dreidimensional ver- wendet	209
9.14	Eine dreidimensionale Mandelbrotstruktur	210
9.15	Das Tetrabrot	211
9.16	Eine gefüllte Julia-Menge	211
9.17	Ein Mandala	212
9.18	Ein weiteres Mandala	212

Vorwort

Lernverfahren der Künstlichen Intelligenz, die oft unter dem Begriff *Deep Learning* zusammengefasst werden, sind seit etwa zehn bis fünfzehn Jahren außerordentlich erfolgreich. Die untersuchten Objekte (das Nervensystem des Menschen, sein Blutkreislauf, Erkrankungen wie Epilepsie, Parkinson, Demenz, Tumore u.a.) besaßen alle eine *fraktale Struktur* [43], [44]. Diesem Gebiet wollen wir uns in diesem Buch zuwenden. Man kann die These aufstellen, dass Fraktale die Verbindung zwischen Naturwissenschaften und Künstlicher Intelligenz herstellen. Unabhängig von Fragen der Künstlichen Intelligenz erklären sie aber auch viele Erscheinungen, die von großem Interesse für die menschliche Gesellschaft sind. Bei der Steuerung komplexer Systeme ist die Zusammenarbeit mit der Kybernetik sehr fruchtbar, ebenso spielt die Informationstheorie eine große Rolle. Die großen Störungen in den Informationssystemen im Juli 2024 sind ein hervorragendes Beispiel für die Entstehung chaotischer Entwicklungen. Das Auftreten von panikartigem Verhalten gehört hierher, ebenso wie die Entstehung von Epidemien aller Art.

Die klassische Mathematik – etwa verkörpert durch Analytische Geometrie und Analysis – ist nicht in der Lage, viele in der Natur auftretende Formen zu beschreiben: Sandstürme, Staubwolken, Erdbeben, Küstenlinien, die Oberfläche eines Blumenkohls, ein Elektroenzephalogramm (EEG), das Wachstums eines Tumors, die Oberfläche und die zeitliche Entwicklung einer Düne und ihre Wanderung, das sind alles Erscheinungen, die man gut kennt, die aber mit der Schulmathematik nicht erfasst werden können.

Sucht man nach Eigenschaften, die bei allen diesen Erscheinungen auftreten, so findet man nach einigem Beobachten der Natur und nach dem Lesen entsprechender Literatur, dass das Prinzip der Selbstähnlichkeit eine zentrale Rolle spielt. Eine Düne, ein Haufen von Bausand, ein Sandhäufchen zum Spielen im Sandkasten – sie haben alle das gleiche Aussehen. Wenn man sie fotografiert, kann man sie durch Vergrößern oder Verkleinern ineinander überführen. Betrachtet man die Adern eines Menschen, so sieht man zunächst die großen, starken Adern, von denen kleinere Adern abgehen. Von jeder kleineren Ader gehen noch kleinere Adern ab usw. Auch hier kann man durch Vergrößern oder Verkleinern die Strukturen ineinander überführen.

Betrachtet man DAX - Kurven auf einem Bildschirm, dann sehen die Kurven für eine Woche, einen Monat oder ein Jahr immer gleich aus. Unterscheidbar werden sie durch Zuordnung einer fraktalen Dimension.

Die Möglichkeit, viele Eigenschaften durch eine fraktale Dimension zu charakterisieren, ist ein großer Schritt in Richtung der mathematischen Beschreibung fraktaler Strukturen. Und nicht zuletzt spielt die ständige Weiterentwicklung der Computer eine große Rolle. Der aktuelle Stand erlaubt es, viele fraktale Strukturen zu simulieren und zu untersuchen, wie die verwendeten Modelle mit den Erscheinungen der realen Welt übereinstimmen.

Ein ganz wesentlicher Bestandteil der Chaostheorie ist der *Schmetterlingseffekt*. Er ist auf den ersten Blick nicht so bekannt, bei vielen Gelegenheiten erinnert man sich an ihn. Eine winzige Veränderung an einer Stelle kann große, eventuell katastrophale Veränderungen an anderen Stellen bewirken. Wir werden oft darauf zurückkommen.

Es ist sehr interessant zu sehen, dass viele bedeutende Mathematiker des 17., 18., 19. und 20. Jahrhunderts zur Theorie der Fraktale beigetragen haben, ohne diese Struktur explizit zu erwähnen: Gottfried Wilhelm Leibniz (1646 – 1716), Georg Cantor (1845 - 1918), Karl Weierstraß (1815 - 1897), S. N. Bose (1894 - 1974), Albert Einstein (1879 - 1955), Emmy Nöther (1882 - 1935), David Hilbert (1862 - 1943) u.v.a.m. Explizit ist die Entwicklung der Fraktale aber mit den Namen von Felix Hausdorff (1868 - 1942) und Benoît Mandelbrot verbunden.

Die Bücher [1], [24] und [25] bieten eine ausgezeichnete Grundlage, sind aber kaum noch verfügbar. Und der Einfluss der Computer ist natürlich außerordentlich bedeutsam für die weitere Entwicklung. Sie bedarf auch der interdisziplinären Zusammenarbeit unterschiedlicher Gebiete.

Chemnitz, Dezember 2024

Christian Posthoff

Bernd Steinbach

1 Einleitung

In den letzten 150 Jahren haben sich alle Gebiete, die das heutige Niveau von Wissenschaft und Technik bestimmen, stark entwickelt; die Informatik, die Elektronik, die Informationstechnik u.a. sind neu entstanden, ebenso die Raumfahrt sowie die Anwendungen der Kernspaltung und der Kernfusion. Diese Entwicklungen wurden begleitet von einer ständigen Erweiterung der Anforderungen an die Mathematik, die sich ebenfalls ständig neuen Problemen stellen musste.

Eines dieser neuen mathematischen Gebiete, das sich auf Grund der neuen Probleme entwickelte, wird unter dem Begriff „**Fraktale**“ zusammengefasst. Seine mathematischen Grundlagen gehen auf den deutschen Mathematiker **Felix Hausdorff** (1868 - 1942) und den französischen Mathematiker **Benoît Mandelbrot** (1924 - 2010) zurück. Sie beruhen unter anderem darauf, dass viele Erscheinungen in der Welt nicht geradlinig und glatt verlaufen, wie es vor allem die **Euklidische Geometrie** darstellt, sondern irgendwie zerzaust und krumm aussehen. Die zu ihrer Untersuchung notwendigen mathematischen Grundlagen und Begriffe sollen hier dargestellt werden, mit der Absicht, sie parallel zur Entwicklung der Künstlichen Intelligenz zu begreifen, da diese beiden Gebiete an vielen Stellen zusammenarbeiten müssen.

1.1 Felix Hausdorff

Felix Hausdorff (geboren am 8. November 1868 in Breslau; gestorben am 26. Januar 1942 in Bonn) war ein deutscher Mathematiker. Er gilt als Mitbegründer der **Allgemeinen Topologie** und lieferte wesentliche Beiträge zur **Maßtheorie**, zur **Funktionalanalysis** und zur **Algebra**. Neben seinem Beruf wirkte er unter dem Pseudonym **Paul Mongré** auch als philosophischer Schriftsteller und Literat. Er wurde von den Nationalsozialisten verfolgt und nahm sich 1942 das Leben, um einer Einlieferung in ein Konzentrationslager zu entgehen.

Hausdorff schrieb nach der Habilitation noch je eine Arbeit über Optik, über Nicht-euklidische Geometrie und über Hyperkomplexe Zahlensysteme sowie zwei Arbei-

ten über Wahrscheinlichkeitstheorie. Sein Hauptarbeitsgebiet wurde jedoch bald die Mengenlehre, vor allem die Theorie der geordneten Mengen. Es war anfangs ein philosophisches Interesse, welches ihn um 1897 dazu führte, Georg Cantors Arbeiten zu studieren. Bereits im Sommersemester 1901 hielt Hausdorff eine Vorlesung über Mengenlehre. Dies war eine der ersten Vorlesungen über Mengenlehre überhaupt, nur Ernst Zermelos Kolleg in Göttingen im Wintersemester 1900/1901 entstand ein wenig früher.



Abbildung 1.1: Felix Hausdorff (1868 - 1942)

Cantor selbst hat nie über Mengenlehre gelesen. Zum Sommersemester 1910 wurde Hausdorff als planmäßiger Extraordinarius an die Universität Bonn berufen. In Bonn begann er mit einer Vorlesung über Mengenlehre, die er im Sommersemester 1912 wesentlich überarbeitet und erweitert, wiederholte. Im Sommer 1912 begann auch die Arbeit an seinem opus magnum, dem Buch „Grundzüge der Mengenlehre“, das im April 1914 erschien.

Hausdorff wurde zum Sommersemester 1913 als Ordinarius an die Universität Greifswald berufen. Diese Universität war die kleinste unter den preußischen Uni-

versitäten. Auch das mathematische Institut war klein; im Sommersemester 1916 und im Wintersemester 1916/17 war Hausdorff der einzige Mathematiker in Greifswald. Dies brachte es mit sich, dass er in der Lehre durch die Grundvorlesungen fast vollständig ausgelastet war. Es bedeutete eine wesentliche Verbesserung seiner wissenschaftlichen Situation, dass Hausdorff 1921 wieder nach Bonn berufen wurde. Hier konnte er eine thematisch weitgespannte Lehrtätigkeit entfalten und immer wieder über neueste Forschungen vortragen. Besonders bemerkenswert ist beispielsweise eine Vorlesung über Wahrscheinlichkeitstheorie vom Sommersemester 1923, in der er diese Theorie axiomatisch - maßtheoretisch begründete, und dies zehn Jahre vor A. N. Kolmogoroffs „Grundbegriffe der Wahrscheinlichkeitsrechnung“. In Bonn hatte Hausdorff mit Eduard Study und später mit Otto Toeplitz herausragende Mathematiker als Kollegen und Freunde.

Der Antisemitismus wurde mit der Machtübernahme der Nationalsozialisten Staatsdoktrin. Von dem 1933 erlassenen „Gesetz zur Wiederherstellung des Berufsbeamtentums“ war Hausdorff zunächst nicht unmittelbar betroffen, da er schon vor 1914 deutscher Beamter war. Es blieb jedoch auch ihm vermutlich nicht erspart, dass eine seiner Vorlesungen von nationalsozialistischen Studentenfunktionären gestört wurde. So brach er seine Vorlesung „Infinitesimalrechnung III“ vom Wintersemester 1934/35 am 20. November ab. Da an der Bonner Universität in diesen Tagen eine Arbeitstagung des Nationalsozialistischen Deutschen Studentenbundes stattfand, die festlegte, dass der Schwerpunkt der Arbeit im laufenden Semester das Thema „Rasse und Volkstum“ sei, liegt die Vermutung sehr nahe, dass Hausdorffs Abbruch der Vorlesung mit diesem Ereignis zusammenhängt, denn er hat nie sonst in seiner langen Laufbahn als Hochschullehrer eine Vorlesung abgebrochen.

Zum 31. März 1935 wurde Hausdorff nach einigem Hin und Her schließlich doch noch regulär emeritiert. Ein Wort des Dankes für 40 Jahre erfolgreiche Arbeit im deutschen Hochschulwesen fanden die damals Verantwortlichen nicht. Er arbeitete unermüdlich weiter und publizierte neben der erweiterten Neuauflage seiner Mengenlehre noch sieben Arbeiten zur Topologie und zur deskriptiven Mengenlehre, die alle in polnischen Zeitschriften erschienen.

Auch der Nachlass Hausdorffs zeigt, dass er in den immer schwieriger werdenden Zeiten ständig mathematisch arbeitete und die aktuelle Entwicklung auf den ihn interessierenden Gebieten zu verfolgen suchte. Dabei unterstützte ihn Erich Bessel-Hagen (1848 - 1946) selbstlos, indem er nicht nur der Familie Hausdorff in Freundschaft die Treue hielt, sondern auch Bücher und Zeitschriften aus der Institutsbibliothek besorgte, die Hausdorff als Jude nicht mehr betreten durfte.

Vergeblich versuchte Hausdorff 1939 über den Mathematiker Richard Courant ein Forschungsstipendium in den USA zu erhalten, um doch noch emigrieren zu können.

Mitte 1941 schließlich wurde damit begonnen, die Bonner Juden in das Kloster „Zur ewigen Anbetung“ in Bonn-Endenich zu deportieren. Von dort erfolgten später die Transporte in die Vernichtungslager im Osten. Nachdem Felix Hausdorff, seine Frau und die bei ihnen lebende Schwester seiner Frau, Edith Pappenheim, im Januar 1942 den Befehl zur Übersiedlung in das Endenicher Lager erhalten hatten, schieden sie gemeinsam am 26. Januar 1942 durch Einnahme einer Überdosis Veronal aus dem Leben. Ihre letzte Ruhestätte befindet sich auf dem Poppelsdorfer Friedhof in Bonn. Seinen handschriftlichen Nachlass übergab er zwischen der Bestellung ins Zwischenlager und der Selbsttötung dem Ägyptologen und Presbyter Hans Bonnet (1887 - 1972), der diesen trotz Zerstörung seines Hauses durch einen Bombentreffer weitestgehend retten konnte.

Manche Bonner Juden machten sich möglicherweise über das Sammellager Endenich noch Illusionen – Hausdorff selbst nicht. E. Neuenschwander entdeckte im Nachlass Bessel-Hagen auch den Abschiedsbrief, den Hausdorff an den jüdischen Rechtsanwalt Hans Wollstein geschrieben hatte; hier Anfang und Ende des Briefes:

„Lieber Freund Wollstein! Wenn Sie diese Zeilen erhalten, haben wir Drei das Problem auf andere Weise gelöst – auf die Weise, von der Sie uns beständig abzubringen versucht haben. Das Gefühl der Geborgenheit, das Sie uns vorausgesagt haben, wenn wir erst einmal die Schwierigkeiten des Umzugs überwunden hätten, will sich durchaus nicht einstellen, im Gegenteil! Was in den letzten Monaten gegen die Juden geschehen ist, erweckt begründete Angst, dass man uns einen für uns erträglichen Zustand nicht mehr erleben lassen wird.“

Nach dem Dank an Freunde und nachdem er in großer Gefasstheit letzte Wünsche bezüglich Bestattung und Testament geäußert hat, schreibt Hausdorff weiter:

„Verzeihen Sie, dass wir Ihnen über den Tod hinaus noch Mühe verursachen; ich bin überzeugt, dass Sie tun, was Sie tun können (und was vielleicht nicht sehr viel ist). Verzeihen Sie uns auch unsere Desertion! Wir wünschen Ihnen und allen unseren Freunden, noch bessere Zeiten zu erleben. Ihr treu ergebener Felix Hausdorff“

Dieser letzte schriftliche Wunsch Hausdorffs erfüllte sich nicht: Rechtsanwalt Wollstein wurde in Auschwitz ermordet.

Bereits 1918 definierte er die nach ihm benannte **Hausdorff-Dimension**, ein Maß für die Rauheit oder, genauer gesagt, für die **fraktale Dimension** eines Objektes. Beispielsweise ist die Hausdorff-Dimension eines einzelnen Punktes gleich Null, eines Linienabschnitts gleich 1, eines Quadrats gleich 2 und eines Würfels gleich 3. Das heißt, für Punktmengen, die eine glatte Form oder eine Form mit einer gerin-



Abbildung 1.2: Das Grab von Felix Hausdorff auf dem Poppelsdorfer Friedhof in Bonn

gen Anzahl von Ecken definieren - die Formen der traditionellen Geometrie und Wissenschaft - ist die Hausdorff-Dimension eine ganze Zahl, die mit der üblichen Bedeutung der Dimension übereinstimmt, auch bekannt als topologische Dimension. Es wurden jedoch auch Formeln entwickelt, die die Berechnung der Dimension anderer, weniger einfacher Objekte ermöglichen, bei denen man allein auf Grund ihrer Eigenschaften der Skalierung und Selbstähnlichkeit zu dem Schluss kommt, dass bestimmte Objekte keine ganzzahlige Hausdorff-Dimensionen haben. Auf Grund der bedeutenden technischen Fortschritte von Abram Samoilovitch Besicovitch (1891 - 1970), die die Berechnung von Dimensionen für hochgradig unregelmäßige oder „grobe“ Mengen ermöglichten, wird diese Dimension gemeinhin auch als Hausdorff - Besicovitch-Dimension bezeichnet. Eine umfangreiche und liebevoll geschriebene Darstellung des gesamten Schaffens von Felix Hausdorff ist in [46] enthalten.

1.2 Benoît Mandelbrot

Der französisch-US-amerikanische Mathematiker und Sachbuchautor ist der Vater der fraktalen Geometrie. Als IBM-Fellow am Thomas J. Watson Research Center, als Sterling Professor für Mathematik an der Yale University und Mitarbeiter am Pacific Northwest National Laboratory, wurde er Pionier der Theoretischen Physik sowie der Finanzmathematik. Der Sohn eines litauisch – jüdischen Kleiderhändlers und einer Ärztin aus Polen lebte als Schüler während der deutschen Besatzung in Frankreich und erwarb seine Kenntnisse in Mathematik und Geometrie hauptsächlich im Selbststudium. Nach Studien und Forschungen in Paris entwickelte Benoît Mandelbrot als Mitarbeiter von IBM in den USA während der 1970er und 1980er Jahre seine Theorie der *Fraktalen Geometrie* und prägte den Begriff „fraktal“.



Abbildung 1.3: Benoît Mandelbrot (1868 - 1942)

Durch die Entdeckung eines formenreichen geometrischen Gebildes, der so genannten **Mandelbrot-Menge**, und den Nachweis von Fraktalen in der Natur gelang dem Mathematiker die Mitbegründung einer völlig neuen Wissenschaftsdisziplin; der

Chaostheorie. Diese Gebilde oder Muster besitzen im Allgemeinen keine ganzzahlige Hausdorff - Dimension, sondern einen gebrochenen Wert und weisen zudem einen hohen Grad von Skaleninvarianz bzw. Selbstähnlichkeit auf. Das ist beispielsweise der Fall, wenn ein Objekt aus mehreren verkleinerten Kopien seiner selbst besteht. Geometrische Objekte dieser Art unterscheiden sich in wesentlichen Aspekten von gewöhnlichen glatten Figuren.

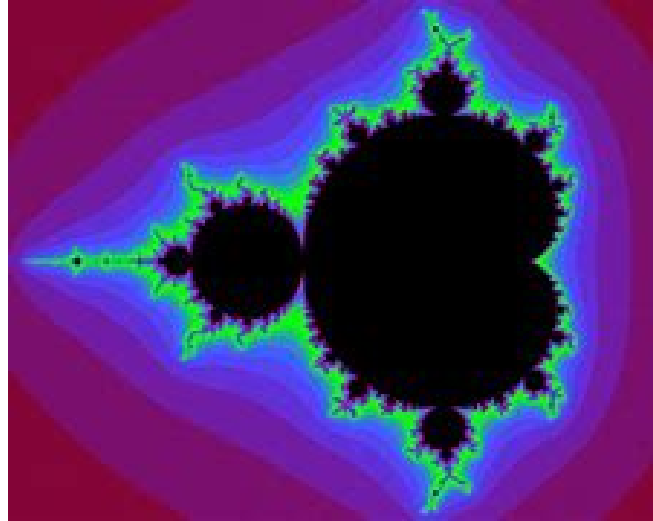


Abbildung 1.4: Die Mandelbrot - Menge



Abbildung 1.5: Zur Winterolympiade 2022 erhielt jedes Land seine eigene Schneeflocke

Fraktale sind geometrische Gestalten oder Formen, die in natürlichen Objekten dargestellt werden, von einem Farnblatt über ein Spinnennetz oder eine Schneeflocke bis hin zu größeren Phänomenen wie Wolken, Wirbelstürmen oder sogar Galaxien

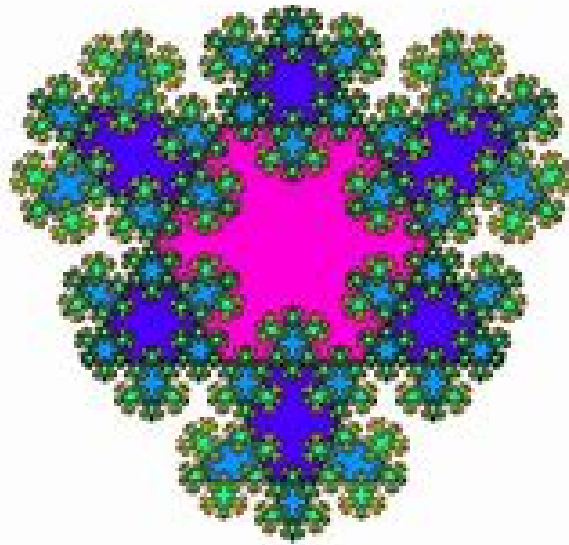


Abbildung 1.6: Eine Variante einer Kochschen Schneeflocke

im All. Die Entdeckung von Fraktalen ermöglichte es, unregelmäßige Formen zu entwickeln, die bis dahin nicht durch die Mathematik definiert oder dargestellt werden konnten. Die klassische Geometrie behandelte Quadrate, Rechtecke, Würfel, Kugeln u.v.a.m. schon seit Jahrtausenden, für Wolken, Baumkronen, Küstenlinien, Flussläufe und deren Delta gab es keine Möglichkeit, sie mathematisch zu beschreiben. Man sieht bereits in Abb. 1.6 sehr schön, dass die Ränder der Schneeflocken-Teile aus geraden Linien bestehen, die Abbildung insgesamt aber anfängt, die Ebene auszufüllen.

Mandelbrot selbst startete mit dem Problem, die Länge der Küstenlinie von Großbritannien zu bestimmen: *How long is the Coast of Britain?* Man geht in folgender Weise vor: es wird ein Punkt an der Küste, etwa Brighton an der Südküste, fixiert. Von dort aus verwendet man den geradlinigen Abstand von 10 km nach links und nach rechts. Diese Prozedur setzt man um die ganze Insel herum fort, bis man sich irgendwo in Schottland trifft. Die Endpunkte werden nicht genau zusammentreffen, aber die Summe all dieser Linien von 10 km ist eine Näherung für die Länge der Küste. Buchten oder Kaps werden ignoriert, man geht davon aus, dass sie sich im wesentlichen ausgleichen. Wenn man genügend Zeit hat, kann man den Maßstab verkleinern, etwa auf 1 km. Die Länge wird jetzt genauer bestimmt, ihr Wert wird größer sein. Die wirkliche Länge hängt also vom Maßstab ab. Einen relativ exakten Wert würde man erhalten, wenn man die gesamte Küste zu Fuß abschreitet, was natürlich an einigen Stellen auf Schwierigkeiten stoßen würde. Heute verwendet man Satellitenbilder und die Bildverarbeitung mit Computerprogrammen. Eine sehr verbreitete fraktale Form, die sowohl in Fraktalen als auch in der gesamten

Natur zu finden ist, ist die Spirale. Dies ist eine der häufigsten Formen, die als Fraktal dargestellt werden; man sieht sie häufig im Internet. Diese mathematische Form ist auch in der gesamten Natur verbreitet.



Abbildung 1.7: Ein fraktales Kunstwerk

Sie wird beschrieben durch die Gleichung

$$T_w(x) = \sum_{n=0}^{\infty} w^n \cdot s(2^n x)$$

mit einem Parameter w und $s(x) = \min_{n \in \mathbf{Z}} |x - n|$ (die Entfernung von x zur nächsten ganzen Zahl) Die Funktionen $T_w(x)$ sind im Intervall $[0, 1]$ stetig, aber nicht differenzierbar.

1.3 Selbstähnlichkeit

Eine strenge Selbstähnlichkeit liegt dann vor, wenn eine Veränderung des verwendeten Maßstabes zu einem gleichartigen Objekt führt. Liegt zum Beispiel ein Kreis mit dem Radius r vor, dann gelten für den Umfang U und den Flächeninhalt F die Gleichungen

$$U = 2 \cdot r \cdot \pi \tag{1.1}$$

$$F = \pi \cdot r^2. \tag{1.2}$$

Verdoppelt man den Radius und verwendet $r' = 2 \cdot r$, dann gelten die gleichen Formeln für U' und F' , man muss nur r durch r' ersetzen. Man kann sogar einen Punkt mit einbeziehen, indem man $r' = 0$ verwendet. Benutzt man für alle Kreise den gleichen Mittelpunkt, dann handelt es sich um konzentrische Kreise.

In einem kartesischen Koordinatensystem verwendet man für eine Gerade die Gleichung

$$y = mx + n. \tag{1.3}$$

Für jeden Punkt (x_0, y_0) ist $m = \frac{y_0}{x_0}$ und beschreibt den Anstieg der Geraden. Eine Änderung von m ändert den Anstieg, eine Änderung von n verschiebt die Gerade nach oben oder unten.

Kugeln werden durch die Gleichung

$$x^2 + y^2 + z^2 = r^2 \quad (1.4)$$

definiert. Hier kann man wieder durch Änderung des Radius selbstähnliche Figuren erzeugen. Für $r = 0$ kann man sogar den Nullpunkt als Kugel ansehen.

Insgesamt kann man sagen, dass alle ein-, zwei- oder dreidimensionalen Figuren, die durch geeignete Gleichungen der Analytischen Geometrie beschrieben werden, Ausgangspunkt streng selbstähnlicher Figuren sein können.

Abbildung 1.8 zeigt eine Spirale, die von der Gleichung

$$r = 2 - 2 \sin \theta$$

erzeugt wird. Für $\theta = \pi/2$ ist $\sin \theta = 1$, also $r = 0$. Für $\theta = 3\pi/2$ ist $\sin \theta = -1$, also $r = 4$. Auf Grund der Periodizität der Sinusfunktion kehrt sie dann wieder zum Nullpunkt und durchläuft die Kurve immer wieder. Man kann jetzt die ersten zwei Werte durch einen anderen Wert ersetzen, ebenso den Faktor 2, und für θ sind ebenfalls Varianten $\theta + \alpha$ möglich, wobei α eine beliebige reelle Zahl sein kann. Damit kann man wieder eine große Menge selbstähnlicher Kurven erzeugen.

Spiralen findet man in vielen Bereichen von Natur und Technik vor. Die schönsten Spiralförmigkeiten zeigt uns der Himmel. Es sind Gebilde aus Milliarden Sonnen, in denen sich kosmische Materie als Spiralen ordnet. Erst kürzlich (im Jahr 2000) hat man festgestellt, dass sich im Zentrum jeder Galaxie ein Schwarzes Loch befindet, das einen spiralförmigen Strudel erzeugt! Nach neuesten Erkenntnissen ist das schwarze Loch Mittelpunkt jeder Galaxie und erzeugt die Spiralförmigkeit. Es ist sowohl ein verschlingender Schlund (Chaos) als auch ein Sternenbildner (Ordnung).[29]

Auch kleinste Einheiten der Materie, subatomare Partikel mit elektrischer Ladung, bewegen sich auf spiralförmigen Bahnen, wenn ein Magnetfeld sie ablenkt.

Die Spiralförmigkeit ist eine Folge von Anpassungen an Lebensbedingungen: Schmetterlinge zum Beispiel besitzen lange Rüssel, um den Nektar aus den Blütenkelchen saugen zu können. Bei einigen tropischen Arten können diese Saugrüssel in Anpassung an langröhrige Blüten 25 Zentimeter lang sein. Während des Fluges wäre ein solch langes Organ natürlich hinderlich. Es zu einer Spirale aufzurollen ist die einfachste Art, mit dem Problem fertig zu werden. In vielen Fällen führt auch das

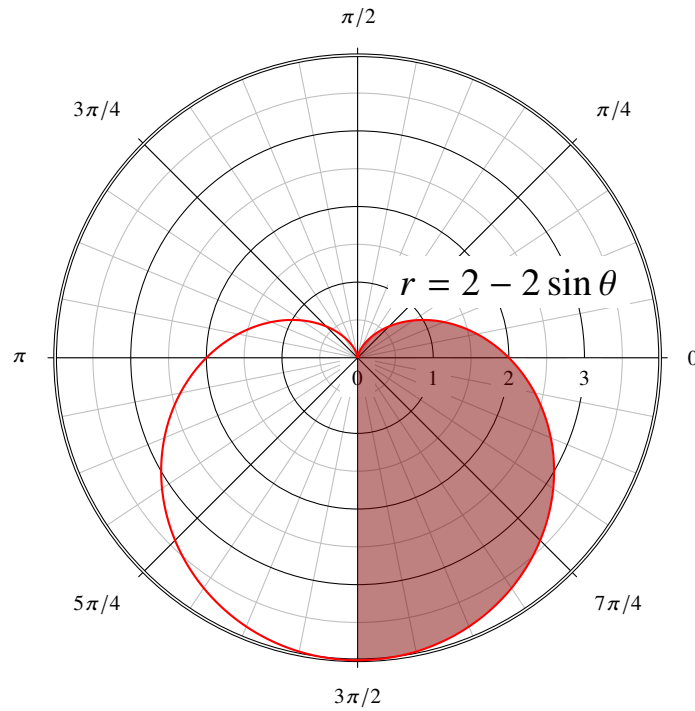


Abbildung 1.8: Spirale

Verhalten eines Lebewesens zu Spiralen: Für die Köcherfliege, die ihre winzigen kugelförmigen Eier auf der Unterseite von Schwimmblättern ablegt, war die Spirale dafür der günstigste Weg.

Im Reich der Pflanzen hat der spiralförmige Wuchs große funktionelle Vorteile: In Knospen kann auf engstem Raum, mehrfach schützend umhüllt, eine große Zahl von Blüten und Blättern bereitgehalten und mit einer winzigen Drehung entfaltet werden. Viele Pflanzen wachsen spiralförmig nach oben, und viele Blätter ordnen sich in Spiralen, um das Licht optimal zu nutzen. Bäume wie Fichten offenbaren bei genauem Studium einen spiralförmigen Wuchs – von den Wurzeln bis hinauf zum Wipfel

Die versteinerten spiralförmigen Schalen der seit langem ausgestorbenen Ammoniten und die ebenfalls versteinerten kleinste Schalen tragende Tiere, die winzigen einzelligen Foraminiferen, offenbaren noch nach Jahrmillionen Wachstumsprozesse und deren funktionellen Sinn.

Schneckenhäuser sind, trotz ihres Variantenreichtums, nach strengen geometrischen Gesetzen aufgebaut. Wohl an keinem anderen Lebewesen wird sinnfälliger, dass Gestalt in eine genetisch vorbestimmte Form hineinwächst. Jede Zelle weiß zu jeder Zeit um das Ganze, um die Idee, die es wachsen lässt. So entstehen nach artspezifischen Bauplänen verschiedenartige Gestalten, alle von hoher Präzision und Schön-

heit. Die meisten Schneckenhäuser sind rechtsdrehend, nur sehr wenige Arten sind linksdrehend. Sehr selten treten innerhalb einer Art falsch-gewundene, spiegelbildlich symmetrische Exemplare auf, sogenannte Inverse – bei der rechtsdrehenden Weinbergschnecke eine unter 5000. Dass spiegelbildliche Gegenstücke eines Schneckenhauses höchst selten existieren, ist Ausdruck eines grundlegenden Unterschiedes zwischen belebter und unbelebter Natur. Zunächst ließe sich vermuten, dass Rechts- und Linksformen statistisch gleichberechtigt auftreten.

Die belebte Natur bevorzugt indes eine der beiden Möglichkeiten: Schon Zuckermoleküle treten stets in der Rechtsform auf, die von Aminosäuren dagegen in der Linksform – sofern sie an Lebensvorgängen beteiligt sind. Die spiegelbildliche Gleichberechtigung wird allerdings auch in der unbelebten Natur verletzt: das Zerfallsschema von Teilchen, die beim radioaktiven Betazerfall von Atomkernen entstehen, erweist sich ebenfalls als asymmetrisch.

Die Gehörschnecke darf als Muschelhorn mit umgekehrter Funktion gelten. Sie wird von der Basilarmembran längs geteilt, auf der das eigentliche Sinnesorgan, das Gortische Organ (*Organon spirale*), liegt. Dort wirken Haarzellen als Rezeptoren und setzen die Schwingungen der Schallwellen in elektrische Impulse um, die dann zum Gehirn geleitet werden.

Wirbelstürme und Wasserstrudel laufen in der nördlichen Erdhemisphäre entgegen dem Uhrzeigersinn, in der südlichen dagegen mit ihm. Damit kann man den Äquator exakt lokalisieren. Dort entstehen überhaupt keine Bewegungen des Wassers.

1.4 Fraktale auf dem Schachbrett

Man findet fraktale Strukturen auch an Stellen vor, wo es auf den ersten Blick nicht zu erwarten ist. Das Damenproblem ist eine kombinatorische Aufgabe: es sollen jeweils acht Damen auf einem Schachbrett so aufgestellt werden, dass sich keine zwei Damen schlagen können, das heißt, es dürfen keine zwei Damen auf derselben Reihe, Linie oder Diagonalen stehen.

Im Mittelpunkt steht dabei die Frage nach der Anzahl der möglichen Lösungen. Im Falle des klassischen 8×8 -Schachbrettes gibt es 92 verschiedene Möglichkeiten, die Damen entsprechend aufzustellen. Betrachtet man Lösungen als gleich, die sich durch Spiegelung oder Drehung des Brettes auseinander ergeben, verbleiben noch 12 Basis-Lösungen.

Das Problem kann auf quadratische Schachbretter beliebiger Größe verallgemeinert werden: Dann gilt es, n Damen auf einem Brett von $n \times n$ Feldern zu positionieren (mit n als Parameter aus den natürlichen Zahlen), unter den gleichen Bedingungen.

Der kleinste Wert mit einer Lösung des Problems ist $n = 4$. Dort existiert nur eine wesentliche Lösung, die zweite Lösung kann man aus dieser ersten Lösung durch Spiegelung an der mittleren Horizontalen oder Vertikalen erzeugen. Dies zeigt das folgende Diagramm:

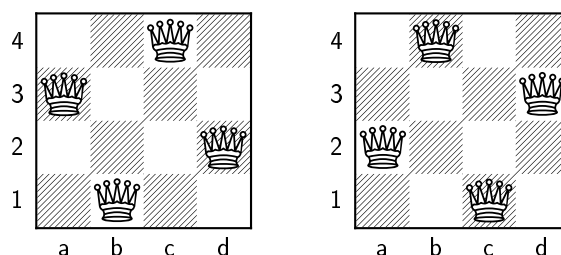


Abbildung 1.9: Alle Lösungen für vier Damen auf einem Schachbrett 4×4

Um eine fraktale Struktur zu konstruieren, kann man $n = 5$ verwenden. Abb. 1.10

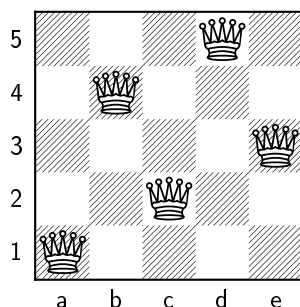


Abbildung 1.10: Eine Lösung für fünf Damen auf einem Brett 5×5

zeigt eine mögliche Lösung. Im Sinne der Selbstähnlichkeit kommt nur eine Vergrößerung dieser Lösung vor. Zunächst ersetzt man jedes einzelne Feld des Brettes 5×5 durch ein ganzes Brett 5×5 und erhält damit ein Brett 25×25 . Auf dieses neue größere Brett überträgt man die Lösung 5×5 in folgender Weise [1]:

- in der kleinen Lösung steht auf a1 eine Dame: man überträgt die Lösung in die entsprechenden Felder.
- Für die Felder b4, c2, d5 und e3 verfährt man in gleicher Weise.
- Da alle weiteren Felder des kleinen Quadrates leer sind, bleiben alle anderen Felder des großen Quadrates auch leer.

Eine Lösung für das große Problem durch Probieren zu finden, dürfte die menschlichen Möglichkeiten schon überschreiten.

Das Ergebnis folgt der Regel, dass in jeder Reihe und in jeder Spalte genau eine Dame steht, das ganze Brett überdeckt ist, aber dass zwischen den Akteuren ein Sicherheitsabstand vorhanden ist. Jeweils fünf Damen leben quasi auf einer Insel, in sicherem Abstand zur nächsten Insel. Eine solche Strategie kann beispielsweise zur Bekämpfung einer Epidemie angewandt werden.

Die Eigenschaften der horizontalen und vertikalen Spiegelung sowie der Spiegelung an einem zentralen Punkt werden wir bei Fraktalen immer antreffen. Bei vielen Interpretationen und Anwendungen ist nicht nur die Tatsache eines Sicherheitsabstandes zwischen den belegten Positionen wichtig, sondern auch der Fakt, dass jede Zeile und jede Spalte genau ein Element enthält. Kurioserweise hat die Lösung für 5×5 noch die Eigenschaft, dass die fünf markierten Felder einem Springer eine Rundreise erlaubt, bei der er nur markierte Felder benutzt.

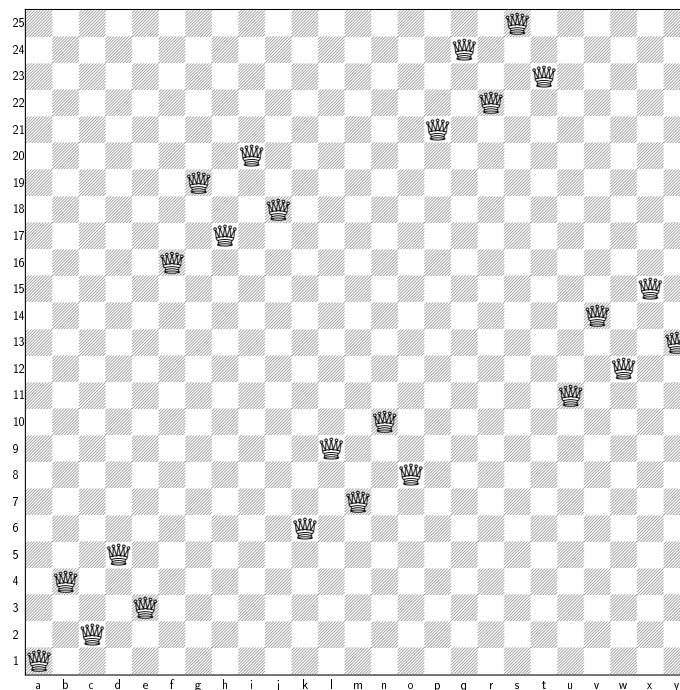


Abbildung 1.11: Eine Lösung des Dameproblems auf einem Brett 25 x 25

Aus dieser Lösung lassen sich nun viele weitere Lösungen konstruieren: man ersetzt beispielsweise jedes Feld durch diese Lösung 25 x 25 und erhält damit eine Lösung für ein Brett der Größe 625 x 625. Lösungen für kleinere Bretter erhält man, wenn man ein Element auswählt und die zugehörige Spalte und Zeile gleichzeitig streicht. Das Finden einiger Lösungen für Bretter beliebiger Größe wird damit möglich. Die Anzahl der möglichen Lösungen erhält man auf diese Weise natürlich hat. Das ist aber nicht so schlimm. Bei der Lösung eines schwierigen praktischen Problems

reicht ja die Realisierung einer Lösung aus.

Die meisten strategischen Spiele (Skat, Schach, Go, chinesisches Schach Xiangqi, japanisches Schach Shogi) weisen Selbstähnlichkeiten auf. Besonders bei Go tritt sie in Erscheinung, nach dem Setzen eines Steines sieht die Position genauso aus wie zuvor; man erkennt nicht so ohne weiteres, dass sich wichtige Eigenschaften der Stellung geändert haben.

1.5 Affine Abbildungen

Die im vorhergehenden Abschnitt aufgetretenen Begriffe der Spiegelung, der Verschiebung u. a. sind Spezialfälle von affinen Abbildungen. Die sollen jetzt allgemein dargestellt werden.

Eine Abbildung ist allgemein eine Funktion f , die einen Raum $\mathbf{X}$ auf einen Raum $\mathbf{Y}$ abbildet, formal $f : \mathbf{X} \rightarrow \mathbf{Y}$, dabei wird jeder Punkt des Raumes $\mathbf{X}$ auf einen Punkt des Raumes $\mathbf{Y}$ abgebildet. Eine affine Abbildung $\mathbf{w}$ in der Ebene hat demnach die allgemeine Form

$$\mathbf{w} : \mathbf{R}^2 \rightarrow \mathbf{R}^2. \quad (1.5)$$

Zu den affinen Abbildungen zählen insgesamt die Translation, die Skalierung, die Rotation, die Spiegelung und die Transvektion (Scherung). Alle diese Abbildungen und deren Kombinationen lassen sich durch eine Matrix A und einen Verschiebungsvektor $\vec{v}$ darstellen:

$$A = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \quad (1.6)$$

$$\vec{v} = \begin{pmatrix} e \\ f \end{pmatrix} \quad (1.7)$$

, zusammengefasst

$$w : \vec{p}' = A\vec{p} + \vec{v} \quad (1.8)$$

$$w : \begin{pmatrix} p'_x \\ p'_y \end{pmatrix} = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \cdot \begin{pmatrix} p_x \\ p_y \end{pmatrix} + \begin{pmatrix} e \\ f \end{pmatrix} \quad (1.9)$$

beschreiben.

Um die Koordinaten eines Bildpunktes (x, y) zu berechnen, muss man zunächst die Matrixmultiplikation

$$A \cdot \vec{x} = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \cdot \begin{pmatrix} p_x \\ p_y \end{pmatrix} = \begin{pmatrix} a \cdot p_x + b \cdot p_y \\ c \cdot p_x + d \cdot p_y \end{pmatrix} \quad (1.10)$$

ausführen. Anschließend addiert man $\vec{v}$ und erhält

$$A \cdot \vec{x} = \begin{pmatrix} a & b \\ c & d \end{pmatrix} \cdot \begin{pmatrix} p_x \\ p_y \end{pmatrix} = \begin{pmatrix} a \cdot p_x + b \cdot p_y + e \\ c \cdot p_x + d \cdot p_y + f \end{pmatrix} \quad (1.11)$$

für die beiden Koordinaten des Bildpunktes.

Es existieren nun verschiedene affine Abbildungen, die alle selbstähnliche Strukturen erzeugen.

1. Translation um einen Vektor $\vec{v} = \begin{pmatrix} v_x \\ v_y \end{pmatrix}$. Wenn man für A die Einheitsmatrix verwendet, dann werden die Ursprungskoordinaten verschoben.

$$w_t : \begin{pmatrix} p'_x \\ p'_y \end{pmatrix} = \begin{pmatrix} 1 & 0 \\ 0 & 1 \end{pmatrix} \cdot \begin{pmatrix} p_x \\ p_y \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \end{pmatrix} = \begin{pmatrix} p_x + v_x \\ p_y + v_y \end{pmatrix} \quad (1.12)$$

2. Veränderung der Größe (Skalierung) Hier verändert man die Größe eines Objektes in der x - und der y -Richtung. Sind die Änderungsfaktoren gleich, dann bleibt die ursprüngliche Form erhalten, sind sie unterschiedlich, dann kann beispielsweise aus einem Quadrat ein Rechteck werden oder umgekehrt.

$$w_t : \begin{pmatrix} p'_x \\ p'_y \end{pmatrix} = \begin{pmatrix} a & 0 \\ 0 & b \end{pmatrix} \cdot \begin{pmatrix} p_x \\ p_y \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \end{pmatrix} \quad (1.13)$$

Dabei sollen a und b positive Zahlen sein.

3. Rotation um den Nullpunkt entgegen dem Uhrzeigersinn mit einer anschließenden Verschiebung

$$w_t : \begin{pmatrix} p'_x \\ p'_y \end{pmatrix} = \begin{pmatrix} \cos \alpha & -\sin \alpha \\ \sin \alpha & \cos \alpha \end{pmatrix} \cdot \begin{pmatrix} p_x \\ p_y \end{pmatrix} + \begin{pmatrix} v_x \\ v_y \end{pmatrix} \quad (1.14)$$

4. Kontraktion: Das ist ebenfalls eine sehr interessante Abbildungsform. Sie kann angewendet werden, wenn man einen Abstand $d(x, y)$ zwischen zwei

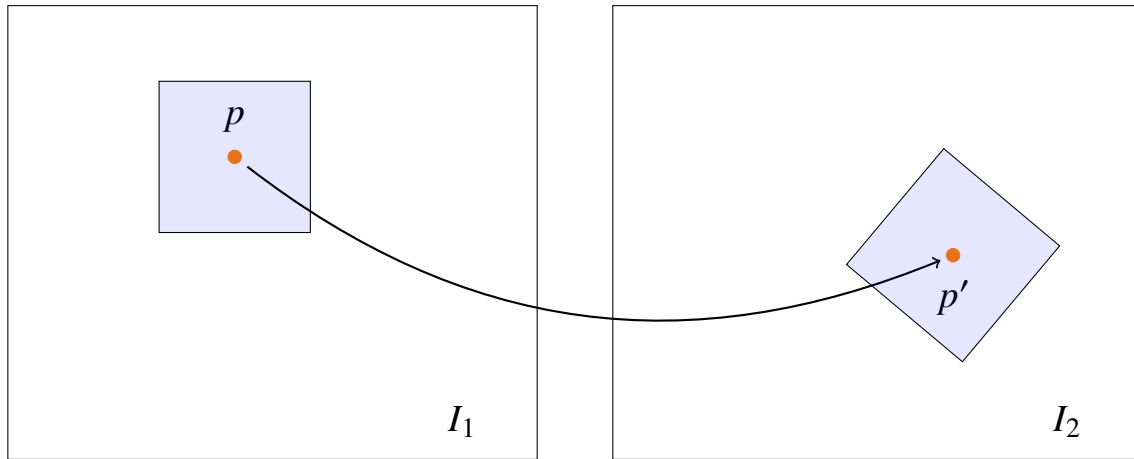


Abbildung 1.12: Verschiebung und Drehung

Elementen definieren kann. Sind in der Ebene zwei Punkte $P_1 = (x_1, y_1)$ und $P_2 = (x_2, y_2)$ gegeben, so beträgt ihr Abstand

$$d(x, y) = +\sqrt{(x_2 - x_1)^2 + (y_2 - y_1)^2} \quad (1.15)$$

Hat man nun eine Abbildung λ , bei der sich stets der Abstand um einen Faktor $\lambda < 1$ verringert, dann wird die Figur immer kleiner und zieht sich schließlich auf den Nullpunkt zusammen. Alle dazwischen liegenden Figuren sind selbstähnlich.

Die verschiedenen Arten von affinen Abbildungen lassen sich auch mühelos miteinander koppeln. Die nachstehende Zeichnung zeigt eine Verschiebung, gefolgt von einer Drehung.

Die nächste Zeichnung stammt von Hugues Vermeiren und zeigt ebenfalls die Kombination verschiedener Abbildungen.

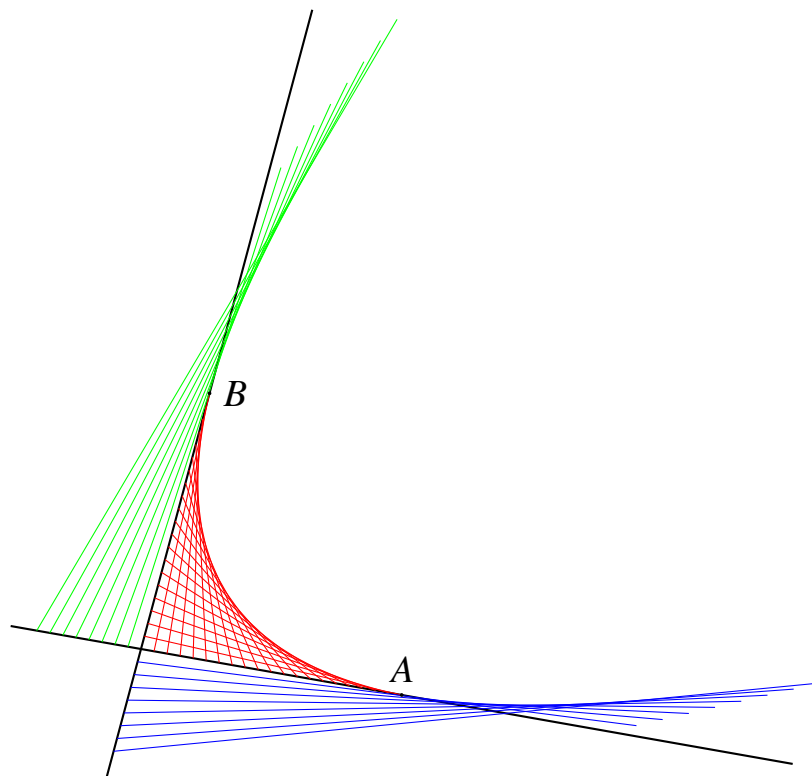


Abbildung 1.13: Konstruktion von Hugues Vermeiren

2 Grundlagen

2.1 Vorläufer

Die frühesten Fraktale stammen von D. Hilbert und G. Peano, ohne dass sie bei ihrer ersten Beschreibung bereits als Fraktale bezeichnet wurden. Die **Hilbert-Kurve** ist eine **spiegelsymmetrische** und **raumfüllende** Kurve in der zweidimensionalen Ebene. Sie lässt sich problemlos auf höhere Dimensionen verallgemeinern.



Abbildung 2.1: Eine zweidimensionale Hilbert - Kurve

Die **Peano-Kurve** ist eine **punktsymmetrische** und **raumfüllende** Kurve in der zweidimensionalen Ebene. Sie lässt sich problemlos auf höhere Dimensionen verallgemeinern. Man beginnt mit der Unterteilung eines Quadrats in neun gleich große Quadrate, die in einer S-Kurve durchlaufen werden. Im nächsten Schritt wird jedes dieser Quadrate wieder unterteilt und die entstehenden Quadrate erneut in S-Kurven durchlaufen, die als neue Kurve zusammengehängt werden usw.

Nach dem Schema in Abb. 2.2 werden oft Warteschlangen konfiguriert, die den vorhandenen Platz optimal ausnutzen, aber eine bestimmte Ordnung erzwingen.

Das **Sierpinski-Dreieck** wird mit Hilfe von selbstähnlichen Dreiecken in der zweidimensionalen Ebene definiert, die gleichseitige Dreiecke oder auch allgemeine Dreiecke sein können. Die dreidimensionale Variante ist das **Sierpinski-Tetraeder**. Auch das Sierpinski-Dreieck kann auf drei Dimensionen erweitert werden, es entstehen **Sierpinski-Simplexe**. Wegen ihrer angeblich höchst seltsamen Eigenschaf-

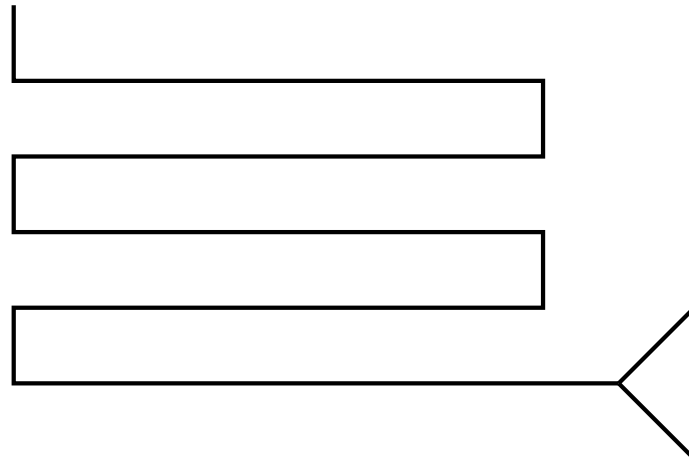


Abbildung 2.2: Die optimale Gestaltung eines Warteraumes

ten wurden fraktale Kurven früher auch **Monsterkurven** genannt. Fraktale Muster

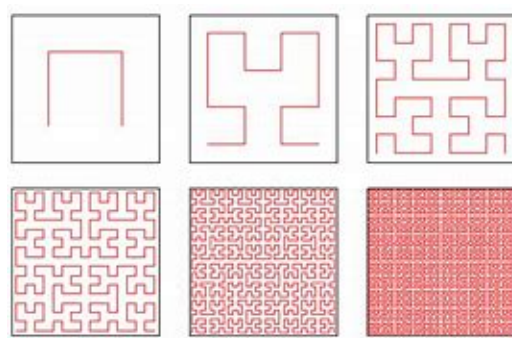


Abbildung 2.3: Eine Peano-Kurve

werden oft durch rekursive Operationen erzeugt. Auch einfache Erzeugungsregeln ergeben nach wenigen **Rekursionsschritten** schon komplexe Muster.

In Abb. 2.4 ist ein **Pythagoras-Baum** zu sehen. Auf einer Grundlinie wird ein Quadrat konstruiert. Auf diesem Grundelement (Stamm) wird auf der Oberseite ein Thaleskreis gezeichnet und dieser beliebig geteilt. Der entstehende Punkt wird mit dem Grundelement verbunden, so dass ein rechtwinkliges gleichschenkliges Dreieck entsteht. Aus den beiden entstandenen Schenkeln des Dreiecks wird wieder jeweils ein Quadrat konstruiert, ein Thaleskreis aufgezeichnet, dieser geteilt, ein rechtwinkliges Dreieck konstruiert und so wieder zu einem Quadrat erweitert. Dieser Vorgang wird beliebig oft wiederholt.

Im Folgenden sei die Seitenlänge des ersten Quadrats (der „Stamm“) gleich 1. Wenn die Innenwinkel des ersten rechtwinkligen Dreiecks gleich 45° , 45° und 90° , die Seitenlängen also gleich $\sqrt{2}/2$, $\sqrt{2}/2$ und 1 sind, ist der Pythagoras-Baum symmetrisch. Die Symmetrieachse ist die Mittelsenkrechte der Hypotenuse des ersten rechtwinkligen und gleichschenkligen Dreiecks. Das Ausgangsquadrat hat die Höhe 1. Das

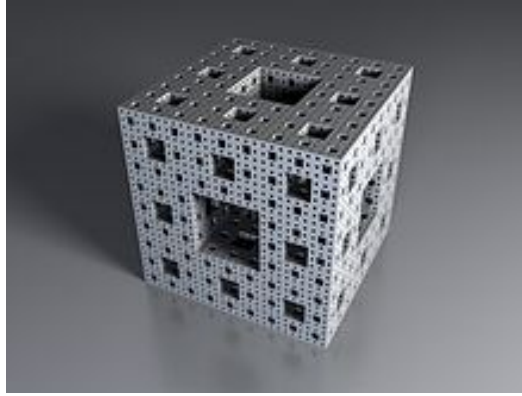


Abbildung 2.4: Ein Menger - Schwamm

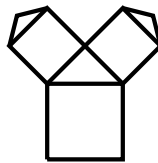


Abbildung 2.5: Die ersten Schritte der Konstruktion des Pythagoras-Baums

Dreieck über diesem Quadrat besitzt die Höhe $1/2$. Damit ist die Höhe bis zur mittleren Ecke der kleineren Quadrate gleich 2. Die Gesamthöhe h , die man bei Fortführung dieses Verfahrens erreichen kann, ergibt sich als Summe einer unendlichen Reihe

$$h = 2(1 + \frac{1}{2} + \frac{1}{4} + \frac{1}{8} + \dots) = 4. \quad (2.1)$$

Man kann die Höhe des Baumes beliebig vergrößern, wenn man nicht $a = 1$ wählt, sondern einen größeren Wert. Dann ist die Höhe $h = 4 \cdot a$. Für die Breite ergibt sich $b = 6 \cdot a$. Eine vollständige Umrundung des Baumes ergibt eine Länge von $\approx 25,313 \cdot a$.



Abbildung 2.6: Ein Pythagoras - Baum

Das Aussehen der Pythagoras-Bäume kann in vielfacher Weise modifiziert werden. Durch Anwendung verschiedener Farben für Quadrate und Dreiecke kann man sehr

schöne Bilder erzeugen. Von der Webseite <https://www.redusoft.de/info/download-demo.html> kann man eine sehr schöne Demo-Version des Programmes **MathProf 5.0** herunterladen. Dort findet man viele weitere interessante Programme für Mathematik und Physik.

Man sieht hier wieder die Effektivität der biologischen Konstruktion: auf einer Fläche von $24 \cdot a^2$ erreicht man das 25-fache einer biologisch aktiven Oberfläche.

Ein weiteres Fraktal ist das **Newton-Fraktal**. Es bildet die Menge der komplexen Zahlen in sich ab, d.h. die Variablen und die Funktionswerte sind im allgemeinen komplexe Zahlen. Es wird durch die Funktion

$$f_z \rightarrow f(z) = z - \frac{p(z)}{p'(z)} \quad (2.2)$$

definiert und wird zum Auffinden der Nullstellen der Funktion p verwendet:

$$f(z) = z^3 - 2z + 2, \quad (2.3)$$

die das Newtonverfahren zum Auffinden von Nullstellen der Funktion $z^3 - 2z + 2$ beschreibt. In Abhängigkeit vom Startpunkt z_0 kann die Folge

$$z, f(z), f^2(z), f^3(z), \dots$$

ein ganz unterschiedliches Verhalten zeigen. Es muss beachtet werden, dass $f^n(z)$ die n -fache Anwendung der Funktion $f(z)$ bedeutet.

Das Verfahren beginnt mit einem Wert z_0 und berechnet den Wert $z_{k+1} = f(z_k)$. Die Farben drücken die Konvergenz-Geschwindigkeit aus sowie die drei Nullstellen. Startwerte, die in den beige gezeichneten Gebieten liegen, konvergieren gegen die gleiche Nullstelle (in Abb. 2.7 in hellbeige), analog für das grüne und das blaue Gebiet. Die Nullstellen zum grünen bzw. blauen Gebiet liegen symmetrisch zur waagerechten Symmetrieachse rechts (Abb. 2.7). Je schneller ein Startwert zu seiner Nullstelle konvergiert, desto heller ist er eingefärbt. Die Werte in den roten Bereichen konvergieren nicht gegen eine Nullstelle, sondern werden vom anziehenden Zyklus $\{0,1\}$ eingefangen. Das Newtonfraktal – im Bild als helle Struktur erkennbar – ist nicht beschränkt. In den drei zu erkennenden Richtungen reicht es bis nach ∞ .

Da man das Newton - Verfahren auf beliebige komplexe Funktionen anwenden kann, kann man auch jede Menge wunderschöner Fraktale erzeugen. Die Färbung kommt dadurch zustande, dass man nach einer bestimmten Anzahl von Schritten die Farbe

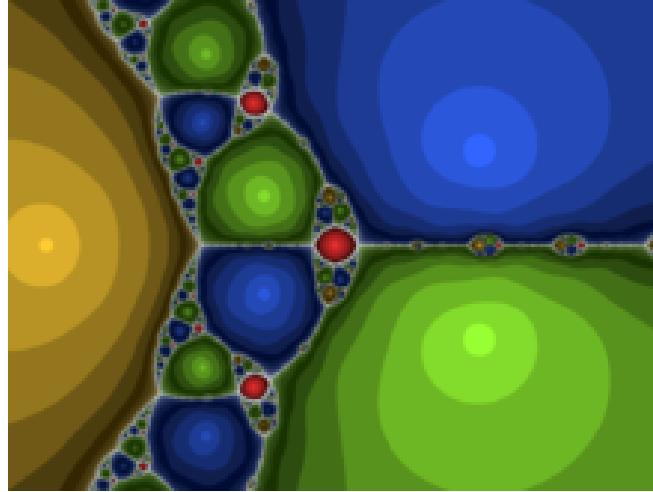


Abbildung 2.7: Newton - Fraktal zur Funktion $z^3 - 2z + 2$

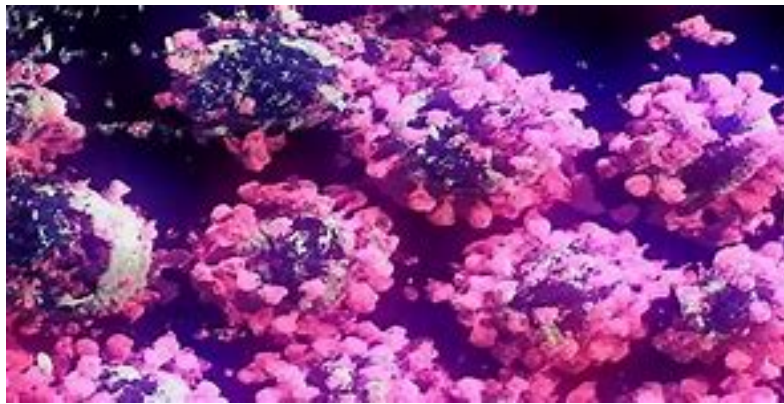


Abbildung 2.8: Eine dreidimensionale Darstellung von Covid-19 - Viren

wechselt. Wählt man die Farbfolge von hell nach dunkel, kann man direkt sehen, auf welche Weise man sich einer Nullstelle nähert.

Viele Probleme in Wissenschaft und Technik konnten in der Vergangenheit erfolgreich bewältigt werden. Die klassische Mathematik geht davon aus, dass alle Objekte exakt (scharf) definiert und Aussagen stets so präzise formuliert sind, dass sie einen Sachverhalt genau widerspiegeln oder ihn verfehlen, d. h., dass jede Aussage entweder wahr oder falsch ist. Parabeln beschreiben beispielsweise die Bahnkurve eines Körpers, der schräg nach oben geworfen wird. Die Bahnen von Kometen können ebenfalls näherungsweise als Parabel dargestellt werden, wenn sie an der Erde vorbeifliegen. Ähnliche Anwendungen existieren für Ellipsen und Hyperbeln.

Die meisten Kurven und Flächen unserer Umwelt können jedoch nicht exakt durch Gleichungen der Analytischen Geometrie beschrieben werden. Die Oberfläche eines Blumenkohls, die Ränder von zerbrochenem Porzellan, Spinnweben u.v.a.m. sind einige Beispiele dafür.

2.2 Die Cantor-Menge

Unter der **Cantor-Menge**, auch **Cantorsches Diskontinuum**, **Cantor-Staub** oder **Wischmenge** genannt, versteht man eine bestimmte Teilmenge der Menge der reellen Zahlen mit besonderen topologischen, maß- und mengentheoretischen sowie geometrischen Eigenschaften. Sie ist nach dem Mathematiker Georg Cantor (1845 - 1918) benannt.

Die Konstruktion der Cantor-Menge ist eine einfache Möglichkeit, fraktale Strukturen zu konstruieren:

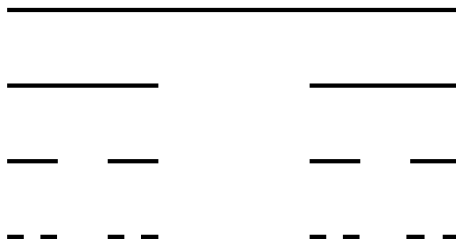


Abbildung 2.9: Cantorscher Staub

Man beginnt mit einer Strecke der Länge 1 und entfernt das mittlere Drittel; man erhält die abgeschlossenen Intervalle $[0, 1/3]$ und $[2/3, 1]$, die Punkte des offenen Intervalls $]1/3, 2/3[$ wurden entfernt. Mit den verbliebenen Strecken verfährt man in gleicher Weise, es wird immer das mittlere Drittel entfernt. Als Streckenlängen erhält man der Reihe nach

$$\frac{1}{3}, \quad \frac{1}{9}, \quad \frac{1}{27}, \quad \frac{1}{81} \quad \dots$$

Die Anzahl der Elemente auf jedem Niveau wächst exponentiell, die Strecken werden immer kürzer, und sie rücken immer mehr zusammen. Der Ausdruck „Staub“ hat durchaus seine Berechtigung. Es ist gleichgültig, auf welchem Niveau man sich befindet, jeder Punkt ist nach links und rechts von seinen Nachbarn getrennt. Trotzdem enthält der Staub genau so viele Elemente wie das ursprüngliche Intervall $[0, 1]$. Die Abschnitte, die man im Konstruktionsprozess eliminiert hat, haben nicht dazu geführt, dass es weniger Elemente geworden sind. Das ist eine wesentliche Eigenschaft von Mengen, die die Mächtigkeit des **Kontinuums** (der Menge der reellen Zahlen) besitzen. Überabzählbare Teilmengen der Menge der reellen Zahlen sind gleich mächtig.

Man betrachtet die beiden abgeschlossenen Intervalle $[0, 1]$ und $[0, 10]$. Die Abbildung

$$y = 10 \cdot x \tag{2.4}$$

$$x = y/10 \tag{2.5}$$



Abbildung 2.10: Ein Sandsturm - Cantorstaub in der Realität

ordnet jedem Element des einen Intervalls genau ein Element des anderen Intervalls zu. Also müssen diese beiden Intervalle gleich mächtig sein. Die Mächtigkeit der Menge der natürlichen Zahlen hat die gleiche Mächtigkeit wie die der rationalen Zahlen und wird mit $\aleph_0$ bezeichnet, die der reellen Zahlen mit $\aleph_1$. $\aleph$ ist der erste Buchstabe des hebräischen Alphabets.

Die Hausdorff-Dimension der Cantormenge ist gleich $\ln 2 / \ln 3 \approx 0,6309$; sie liegt also unter 1, aber sie ist größer als 0. Auf jeder Ebene hat man 3^n Abschnitte der Länge n , 2^n Abschnitte reichen zum Überdecken. Der Staub ist also ein Mittelding zwischen Punkt und eindimensionaler Kurve (Abb. 2.8).

Die Cantor-Menge ist kompakt, abgeschlossen, besitzt keine isolierten Punkte und total unzusammenhängend (ein Diskontinuum) und nirgends dicht. Ihr Flächeninhalt ist gleich null. Sie ist selbstähnlich und besitzt eine nicht ganzzahlige Hausdorff-Dimension; Ihre Mächtigkeit ist gleich $\aleph_1$.

Die Vorgehensweise, die zur Cantor-Menge führte, beruhte darauf, dass man den mittleren Teil einfach eliminierte.

Die Vorgehensweise zur Konstruktion des Cantor-Staubes kann noch weiter modifiziert werden. Bei jeder Iteration wird die Länge des linken Intervalls durch eine Zufallsvariable C_1 definiert, die die Länge eines Teilintervalls des ursprünglichen Segments definiert. Das Gleiche gilt für das rechte Intervall mit einer Zufallsvariablen C_2 . Die Hausdorff-Dimension s erfüllt dann die Gleichung $E(C_1^s + C_2^s) = 1$. $E(X)$ ist der Erwartungswert von X . In der Literatur wird für $E(C_1) = 0,5$ und $E(C_2) = 0,3$ eine Hausdorff-Dimension von 0,7499 angegeben.

Für die Konstruktionen von Mengen, die dem Cantor-Staub ähnlich sind, kann man



Abbildung 2.11: Der Saturn und seine Ringe

auch Potenzen rationaler Zahlen verwenden, zum Beispiel die Folge

$$(1/2)^n = (1, 1/2, 1/4, 1/8, \dots) .$$

Man geht von $x = 1$ nach links und lässt jedes zweite Intervall weg, als erstes das zwischen $1/2$ und $1/4$, danach das zwischen $1/8$ und $1/16$ usw. Die weggelassenen Intervalle werden immer kleiner. Interessant ist auch, wenn man Potenzen von $1/3$ verwendet. Die Elemente der damit konstruierten Menge drängen ebenfalls immer mehr gegen den Nullpunkt, aber viel schneller.

Auch im Weltraum kommen ähnliche Strukturen vor.

In Abb. 2.9 ist der Saturn-Nebel dargestellt. Man sieht an seinem Äquator konzentrische Kreise, die aus einer Art Staub bestehen. Sie liegen dicht beieinander, aber die Teilchen scheinen nicht von einem Kreis auf einen anderen zu wechseln.

2.3 Die Hausdorff - Dimension

Am Ende des vorher gehenden Abschnittes wurde für den Cantorstaub eine Hausdorff - Dimension von $\approx 0,6309$ angegeben. Jetzt soll beschrieben werden, wie diese gebrochenen Dimensionen allgemein zustande kommen.

Für eine Punktmenge endlicher Ausdehnung in einem dreidimensionalen Raum betrachtet man die Anzahl n der Kugeln mit dem Radius r , die mindestens erforderlich sind, um die Punktmenge zu überdecken. Diese Mindestanzahl ist eine

Funktion $n(r)$ des Radius r . Je kleiner der Radius ist, umso größer ist n . Aus der Potenz von r , mit der $n(r)$ für den Limes r gegen Null anwächst, berechnet sich die Hausdorff-Dimension D :

$$n(r) \sim \frac{1}{r^D}. \quad (2.6)$$

Die Lösung dieser Gleichung ergibt dann

$$D = -\lim_{r \rightarrow 0} \frac{\log n}{\log r}. \quad (2.7)$$

Anstelle von Kugeln können ebenso gut Würfel oder vergleichbare Objekte verwendet werden. Bei Punktemengen in der Ebene können auch Kreise zur Überdeckung verwendet werden. Bei Punktmengen in mehr als drei Dimensionen müssen entsprechend höherdimensionale Kugeln verwendet werden.

2.4 Der Feigenbaum-Attraktor

Die verwendete Gleichung wurde ursprünglich 1837 von Pierre François Verhulst als demographisches mathematisches Modell eingeführt. Die Gleichung ist ein Beispiel dafür, wie komplexes, chaotisches Verhalten aus einfachen nichtlinearen Gleichungen entstehen kann. Die zugehörige Dynamik kann anhand eines sogenannten Feigenbaum-Diagramms veranschaulicht werden (Abb. 2.10) [4]. Sie beruht auf der folgenden Gleichung:

$$x_{n+1} = r \cdot x_n \cdot (1 - x_n) = -r \cdot x_n^2 + r \cdot x_n. \quad (2.8)$$

x_0 ist ein Wert zwischen 0 und 1, r liegt zwischen 0 und 4. Entsprechend dem für r gewählten Wert verhalten sich die Wertefolgen von x verschieden. Dabei hängt dieses Verhalten nicht vom Startwert x_0 ab, sondern nur von r :

- Mit r von 0 bis 1 erhält man immer 0 nach einigen Iterationen.
- Mit r zwischen 1 und 3 stellt sich der Grenzwert $1 - \frac{1}{r}$ ein.
- Mit r zwischen 3 und $1 + \sqrt{6} \approx 3,45$ wechselt die Folge bei fast allen Startwerten, ausgenommen 0, 1 und $1 - \frac{1}{r}$, zwischen zwei Werten, **Attraktoren** genannt.

- Mit r zwischen $1 + \sqrt{6} \approx 3,45$ und $\approx 3,54$ wechselt die Folge bei fast allen Startwerten zwischen vier Attraktoren.
- Wird r größer als 3,54, stellen sich erst 8, dann 16, 32 usw. Attraktoren ein. Bei $r \approx 3,57$ beginnt das Chaos.

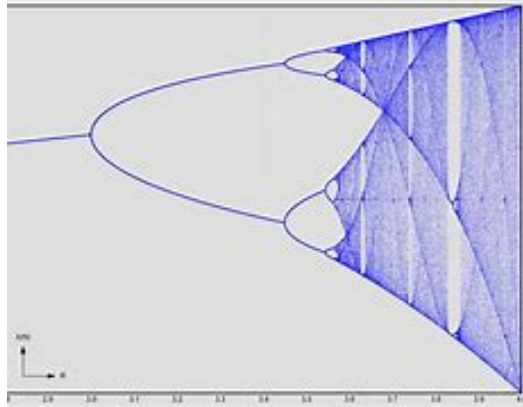


Abbildung 2.12: Ein Feigenbaum - Diagramm

Ein Feigenbaum-Diagramm (Abb. 2.12) entsteht durch Aufzeichnen der Wertepaare (r, x_n) . Man wählt für das Diagramm eine bestimmte Anzahl von Iterationen n , die dann für alle berechneten Punkte des Diagramms beibehalten wird. Für jeden Wert r wird dann eine Reihe von x_0 -Startwerten zwischen 0 und 1 gewählt. Für jeden Startwert wird die Gleichung zusammen mit dem gewählten r x_n -mal iteriert. Bei den so berechneten Werten x_n wird dann bei r im Diagramm jeweils ein Punkt eingezeichnet. Dies wiederholt man für den ganzen Bereich von r .

Jeder Punkt im Diagramm ist das Resultat von einer festgelegten Zahl von n Iterationen für einen bestimmten Startwert x_0 und einem bestimmten Parameter r . In den Bereichen, wo einzelne Linien sichtbar sind, konvergieren fast alle Startwerte bei einem bestimmten r -Wert auf die gleichen Endwerte x_n zu. Diese Werte bzw. Linien bezeichnet man deshalb als **Attraktoren**. In den anderen Bereichen erhält man für jeden Startwert einen anderen Endwert. Diese Endwerte liegen jedoch nicht auf einer Reihe, sondern verteilen sich (pseudo-) zufällig.

2.5 Die Länge der Koch - Kurve

Bei der Konstruktion einer Koch - Kurve (Abb. 2.13) beginnt man mit einer Strecke der Länge 1. Diese wird in drei Teile geteilt. Der mittlere Teil wird weggelassen, stattdessen werden zwei Strecken der Länge $1/3$ angetragen. Die Gesamtlänge der Kurve beträgt jetzt $4/3$.

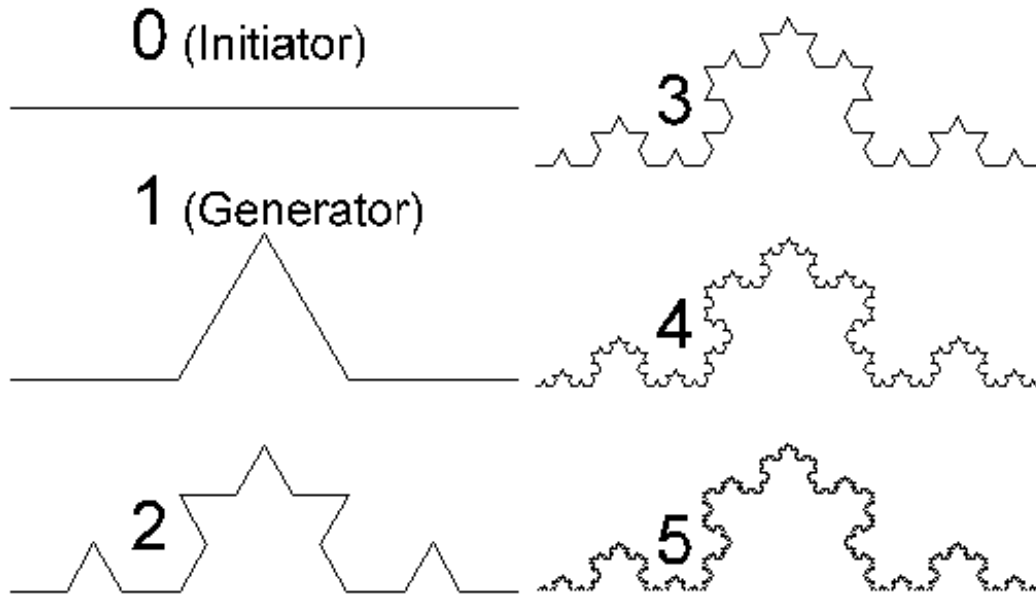


Abbildung 2.13: Die Konstruktion der Kochkurve

Bei jedem Iterationsschritt wird jede Strecke des Streckenzugs durch vier Strecken mit $1/3$ der Streckenlänge ersetzt (Abb. 2.13). Die Kurve wird also mit jedem Iterationsschritt um den Faktor $4/3$ länger. Nach dem Iterationsschritt n ist die Kurvenlänge also auf $(4/3)^n$ angewachsen, wenn die Anfangsstrecke die Länge 1 besaß. Demzufolge nimmt die Länge immer weiter zu, wenn man die Prozedur fortführt. Die Fläche unterhalb der Kurve ist hingegen begrenzt. Sie nähert sich mit wachsendem n dem Wert $\sqrt{3}/20 \approx 0,0866$ an. Die Koch-Kurve ist nach ihrer Konstruktionsvorschrift streng selbstähnlich, das heißt, es erscheinen bei beliebiger Vergrößerung immer wieder die gleichen Strukturen. Sie hat eine Hausdorff-Dimension von

$$\frac{\log 4}{\log 3} \approx 1,262.$$

Die Kurve wird also immer länger, aber die verwendete Fläche bleibt gleich. Viele biologische Systeme verwenden diese Struktur, sie ermöglicht, die Anzahl bestimmter Elemente zu erhöhen, ohne dass das benutzte Volumen oder die benutzte Fläche größer wird.

Die euklidische Ebene kann mit Kochschen Schneeflocken zweier verschiedener Größen parkettiert werden. Diese Parkettierung ist periodisch, spiegelsymmetrisch, punktsymmetrisch, drehsymmetrisch und translationssymmetrisch (Abb. 2.13).

Dabei sind die Abmessungen der gelben Schneeflocken um den Faktor $\sqrt{3} \approx 1,73$ größer als die der grünen Schneeflocken. Die gesamte Fläche, die durch gelbe Flecken überdeckt wird, ist dreimal so groß wie die Fläche, die durch grüne Flecken

überdeckt wird.

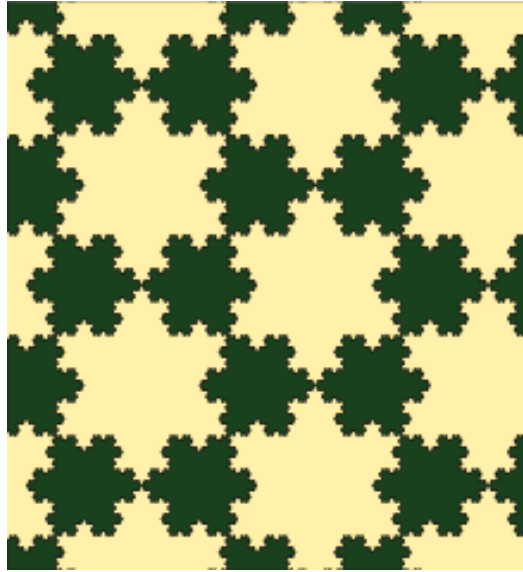


Abbildung 2.14: Ein Parkett aus Kochschen Flocken

Die Konstruktionen der Koch-Kurve stammt aus dem Jahr 1904. Helge von Koch war ein schwedischer Mathematiker (1870 - 1924). Er konstruierte die nach ihm benannte Koch-Kurve, eines der ersten Fraktale, als Beispiel für eine unendlich lange, an keiner Stelle differenzierbare Kurve.

2.6 Sierpinski - Strukturen

Das Sierpinski-Dreieck ist ein 1915 von dem polnischen Mathematiker Waclaw Sierpiński (1882 - 1969) beschriebenes Fraktal. Es handelt sich um eine selbstähnliche Teilmenge eines meist gleichseitigen Dreiecks.

In dem großen gleichseitigen Dreieck verbindet man die Mittelpunkte der drei Seiten miteinander und schneidet das mittlere dunkle Dreieck aus. Es bleiben an den Ecken drei kleinere, ebenfalls gleichseitige, Dreiecke übrig, mit denen man in gleicher Weise verfährt. Es entstehen immer mehr dunkle Dreiecke, die alle innerhalb des Ausgangsdreiecks liegen (Abb. 2.15). Die Fläche des Ausgangsdreiecks beträgt

$$F = \frac{\sqrt{3}}{4} \cdot a^2 \approx 0,433 \cdot a^2 . \quad (2.9)$$

Die vier entstehenden Dreiecke - das dunkle in der Mitte und die drei schattierten

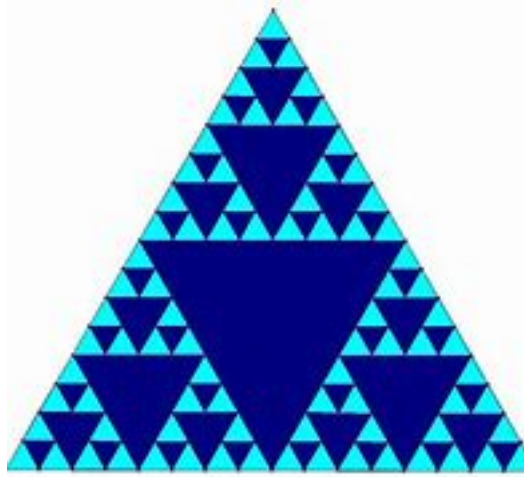


Abbildung 2.15: Ein Sierpinski - Dreieck

an den Eckpunkten besitzen jeweils eine Fläche von

$$F' = \frac{\sqrt{3}}{16} \cdot a^2 \approx 0,108 \cdot a^2 . \quad (2.10)$$

Man zieht jetzt die Fläche des inneren Dreiecks von der Gesamtfläche ab und erhält

$$F^1 = \frac{\sqrt{3}}{4} \cdot a^2 - \frac{\sqrt{3}}{16} \cdot a^2 \approx 0,325 \cdot a^2 . \quad (2.11)$$

Will man jetzt die Kanten zu einem bestimmten Zweck verwenden, dann hat man jetzt die Außenkanten mit der Länge $3a$ zur Verfügung und zusätzlich drei innere Kanten der Länge $a/2$, die das dunkle Dreieck begrenzen, also insgesamt eine Länge von $4,5 \cdot a$. Führt man das Verfahren fort, so wird die Gesamtlänge der zur Verfügung stehenden Kanten immer größer, ohne dass man eine größere Fläche verwenden muss.

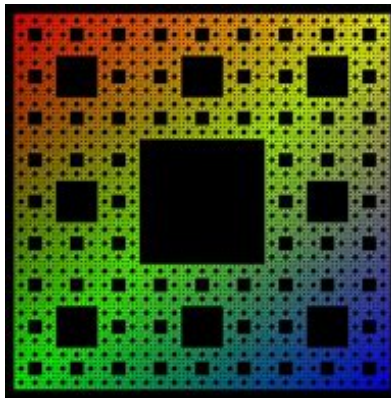


Abbildung 2.16: Ein Sierpinski-Teppich



Abbildung 2.17: Dreidimensionale Flocken

Der Sierpinski-Teppich entsteht, wenn man anstelle der Dreiecke Quadrate verwendet. Die fraktale Dimension beträgt $\ln 8 / \ln 3$.

Man kann auch eine Pyramide als Ausgangspunkt werden. Aus einer gegebenen großen Pyramide werden immer kleiner werdende Pyramiden entfernt, ähnlich wie es beim Sierpinski-Dreieck geschieht.



Abbildung 2.18: Eine dreidimensionale Verallgemeinerung des Sierpinski-Dreiecks

2.7 Von Flächen und Körpern abgeleitete Fraktale

Viele weitere geometrische Figuren können als Ausgangspunkt für die Definition von Fraktalen dienen. Sehr einfach ist das für Kreise und Kugeln (Abb. 2.19). Man sieht erneut das Prinzip, innerhalb eines relativ geringen Volumens eine große Oberfläche zu bilden.

Besonders biologische Strukturen folgen diesem Prinzip. Abb. 2.20 zeigt einen Krautkopf - möglichst viele Blätter innerhalb eines relativ kleinen Volumens.



Abbildung 2.19: Fraktale Kreise und Kugeln

Auch beim Kopf und beim Gehirn kann man noch die Kugeln im Hintergrund ahnen, sie sind aber noch weiter verfeinert (Abb. 2.21).

Der **Rössler-Attraktor** (nach Otto E. Rössler, * 1940) ist ein Attraktor, der durch das folgende Differentialgleichungssystem definiert wird:

$$\frac{dx}{dt} = -(y + z) \quad (2.12)$$

$$\frac{dy}{dt} = x + ay \quad (2.13)$$

$$\frac{dz}{dt} = b + (x - c) \cdot z. \quad (2.14)$$

Laut Otto E. Rössler wurde dieses Modell durch die Betrachtung einer Bonbon-Knetmaschine (taffy puller) auf Coney Island inspiriert, die ihre Toffeemasse wiederholt dehnt und faltet. Der **Rössler-Attraktor** beschreibt kein real existierendes physikalisches System. Es handelt sich um ein System von Differentialgleichungen, das bestimmte chaotische Effekte veranschaulichen soll.

Abb. 2.22 zeigt sehr schön, dass nur die Gleichung 2.14 durch die Multiplikation $x \cdot z$ nichtlinear ist. In der $x - y$ -Ebene geht es ganz ruhig zu. Die fraktale Dimension des Rössler-Attraktors ist etwas größer als 2. Für $a = 0,1$, $b = 0,1$ und $c = 14$ wird sie zwischen 2,01 und 2,02 geschätzt (siehe Abb. 2.21). Der Wert der fraktalen Dimension zeigt an, dass es sich hier im wesentlichen um eine Fläche handelt mit einem kleinen Abstecher in die dritte Dimension.



Abbildung 2.20: Ineinander geschachtelte gekräuselte Kugeloberflächen



Abbildung 2.21: Die fraktale Struktur des Gehirns

2.8 Perkolationen

Die Perkolationstheorie stellt eines der einfachsten Modelle für ungeordnete Systeme dar. Sie wurde entwickelt, um ungeordnete Medien mathematisch zu behandeln, in denen die Unordnung durch eine zufällige Variation des Konnektivitätsgrades definiert ist. Das Hauptkonzept der Perkolationstheorie ist das Vorhandensein einer Perkolationsschwelle, oberhalb derer sich die physikalischen Eigenschaften des gesamten Systems drastisch ändern.

Man kann den Punkten oder den Flächen ganz unterschiedliche physikalische Eigenschaften zuschreiben und damit bestimmte physikalische Eigenschaften modellieren. Beide Arten zeigen, dass die Verbindung zwischen den Elementen eine Rolle spielen kann. Man kann beispielsweise Kanten zwischen den Knoten oder verbundene Flächen als elektrisch leitend ansehen. Solange keine Verbindung zwischen gegenüber liegenden Außenseiten existiert, wirkt das System als Isolator. Wenn

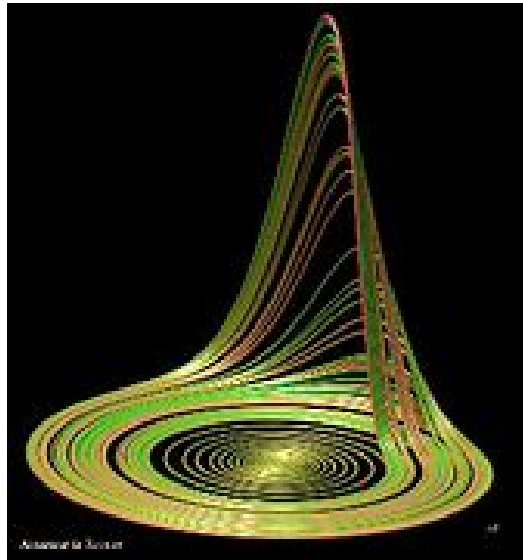


Abbildung 2.22: Der Rössler - Attraktor

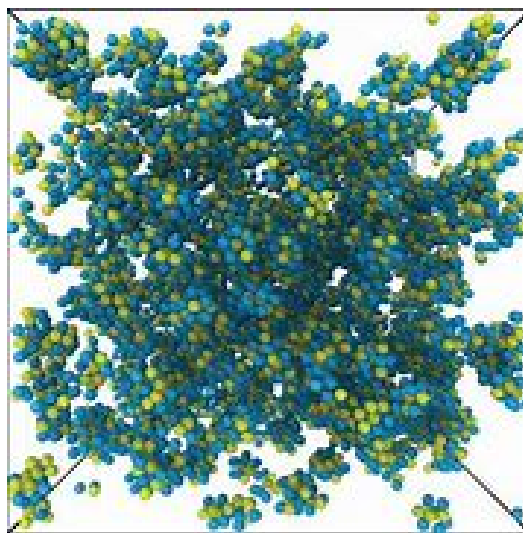


Abbildung 2.23: Ein Perkulations-Fraktal

aber die leitenden Inseln zwei gegenüberliegende Außenkanten verbinden, handelt es sich um einen Leiter. Dazu müssen wenigstens fünf Kanten von links nach rechts oder fünf schwarze Quadrate von links nach rechts vorhanden sein. Das ist aber das Minimum. Es ist leicht zu sehen, dass es viele Strukturen gibt, wo fünf Elemente nicht ausreichen. Es gibt aber auch einen Schwellwert, der sichert, dass es sich in jedem Fall um eine leitende Verbindung handelt.

Ein anderes Beispiel ist das Durchfließen von Wasser durch ein festes Substrat. Dabei können Mineralien herausgelöst werden; die ausgewaschenen Partikel nennt man Perkolate. Perkolation wird auch als Sickerlaugung bezeichnet. Gebräuchlich ist dieser Begriff vor allem in der Hydrologie, wo er das den Boden durchsickernde Wasser umfasst. Es kommt erst zur Perkolation, wenn die maximale Wasserspei-

cherkapazität des Bodens gegenüber der Schwerkraft überschritten wird. Geschieht dies auf der gesamten Mächtigkeit der ungesättigten Zone, so wird neues Grundwasser gebildet.

In der pharmazeutischen Technik wird die Perkolation zur Gewinnung pflanzlicher Wirkstoffe benutzt. Hierbei wird ein zumeist warmes Lösungsmittel, beispielsweise Wasser oder Alkohol, durch die Pflanzen oder Pflanzenteile hindurch geleitet. Ein bekanntes Beispiel für die Perkolation ist auch die Zubereitung von Filterkaffee. In der Spirituosenfabrikation versteht man unter Perkolation eine Extraktion von Drogen durch **Verdrängung**. Man benutzt dafür zylindrische Gefäße aus Kupfer oder Steinzeug, die sich nach unten hin zum Auslauf verjüngen. Die zerkleinerten Drogen werden im Perkulator zwischen zwei Siebe platziert und der 30 – 60 - prozentige Ansatzsprit darauf gegeben. Der Sprit diffundiert und reichert sich innerhalb einiger Tage mit den alkohollöslichen Gehaltsstoffen der Drogen an. Nach drei bis sechs Tagen wird das Vorperkolat langsam abgelassen. Über den Zufluss wird dann gleichzeitig so lange weiterhin Ansatzsprit eingetrichtert, bis die Gehaltsstoffe der Pflanzenteile im Wesentlichen ausgezogen sind. Es folgt die Nachperkolation, bei der durch Zugabe von Wasser versucht wird, den verbliebenen Sprit herauszudrücken.

In vielen Ländern der Erde gibt es Regenzeit und Trockenzeit. Viele Gewässer trocknen völlig aus. Wenn die Regenzeit beginnt, entstehen erst kleinere, voneinander getrennte Pfützen. Die verbinden sich, wenn es weiter regnet; hält der Regen an, dann beginnt das Wasser zu fließen, bis hin zu mächtigen Überschwemmungen.

Aktuell sind gegenwärtig die Methoden zur Modellierung von Epidemien. Man kann sich für den Erreger empfängliche Individuen als Knoten und Kontakt zwischen den Individuen als Kanten vorstellen. Ab einer bestimmten Populationsdichte, wenn es genügend Kontakte zwischen den Individuen gibt, würde die Perkolationsschwelle überschritten, sich also große, zusammenhängende Cluster bilden, die zu einer Ausbreitung des Erregers auf größere Bereiche der Population führen.

Im täglichen Leben kommen viele perkolationsartige Phasenübergänge vor, z. B. das „Puddingproblem“ (Gel-Bildung), das „Sahnesteif-Problem“ und das Problem der „Verklumpung“. In allen Fällen geht die Wirkung erst bei Überschreiten eines kritischen Wertes des ursächlichen Parameters gegen das erwünschte oder unerwünschte Maximum, und zwar meist nach einem Potenzgesetz mit einem kritischen Exponenten, wobei die maximale Wirkung bei Überschreiten des kritischen Wertes zunächst sehr rasch ansteigt. Durch chemische Zusätze, etwa Pudding- oder „Sahnesteif“-Pulver, kann man den kritischen Wert herabsetzen, ohne allerdings das Prinzip zu ändern.

Bei der Analyse von Satellitenfotos dreier verheerender Waldbrände in Europa fiel



Abbildung 2.24: Waldbrände in Sibirien

Guido Caldarelli von der La-Sapienza-Universität in Rom deren Ähnlichkeit mit Fraktalen auf. In der Tat scheinen viele Waldbrände das typische Fransenmuster von so genannten Perkulations-Clustern an ihren Rändern aufzuweisen. Einen weiteren Hinweis auf das fraktale Wachstum der Brände ist der Umstand, dass meist nicht die gesamte von der Feuerfront eingeschlossene Fläche brennt, sondern nicht brennende Inseln übrig bleiben. Caldarelli nahm diese beiden für Perkulationen typischen Merkmale zum Anlass, Waldbrände in Simulationen auf fraktales Wachstum zu analysieren. Seine generierten Bilder weisen eine hohe Übereinstimmung mit Satellitenaufnahmen auf. Daher scheint die Ausbreitung von Waldbränden in der Tat einem fraktalen Wachstum zu gehorchen. Löschkaktionen sollten sich daher auf die Randbereiche eines Brandes konzentrieren – dort kann das fraktale Wachstum, sprich die Ausbreitung des Feuers, effektiver gestoppt werden als durch einen Löschversuch des Herdes (Abb. 2.24).

Abb. 2.25 zeigt, dass für Fraktale keine Grenze existiert. Es wird eine interstellare Gaswolke gezeigt, die zur Bildung neuer Sterne führen kann.

Eine andere Form der Percolation ist die Invasionspercolation. Sie liefert ein Modell für die wechselseitige Verdrängung von Flüssigkeiten, die beispielsweise beim Zusammentreffen von Erdöl und Wasser zu beobachten ist.

Der Jerusalem-Würfel ist ein solides Fraktal, das von Eric Baird entdeckt wurde. Seine Konstruktion ähnelt der des Menger-Schwammes. Jede Iteration erzeugt zwei



Abbildung 2.25: Eine interstellare Gaswolke

selbstähnliche Elemente vom Rang $n + 1$ und $n + 2$. Er wird so genannt wegen der Ähnlichkeit mit dem Jerusalem-Kreuz.

Seine Hausdorff-Dimension ist $D \approx 2,529$. Wieder sieht man das interessante Übergangsverhalten. Die Dimension drückt aus, dass es sich hier um mehr als eine Fläche handelt, dass es aber zu einem Körper nicht ganz reicht. Das kann man nun mit vielen weiteren Körpern ausprobieren; sie sind nach steigender Hausdorff-Dimension angeordnet.

- Fraktales Ikosaeder: $d = \frac{\log(12)}{\log(1+\phi)} \approx 2,5819$.
- Fraktales Oktaeder: $d = \log_2 6 \approx 2,5849$.
- Kochfläche: Jedes gleichseitige Dreieck wird durch 6 kleinere Dreiecke ersetzt, $d = \log_2 6 \approx 2,5849$.
- Griechisches Kreuz: $d = \log_2 6 \approx 2,5849$.

Erweiterungen der Hilbert-, Lebesgue-, Moore- und Mandelbrot-Kurven besitzen die Hausdorff-Dimension 3.

Bei Drachenkurven kann man sehr schön sehen, wie durch kurze geradlinige Bewegungen eine ebene Fläche langsam ausgefüllt wird. Die folgende Zeichnung zeigt, welche Schrittfolge zu einer Drachenkurve führt.

Die Kurve entsteht entsprechend dieser Zeichnung durch eine Folge von jeweils zwei



Abbildung 2.26: Verschmutzung einer Umgebung durch ausgelaufenes Öl

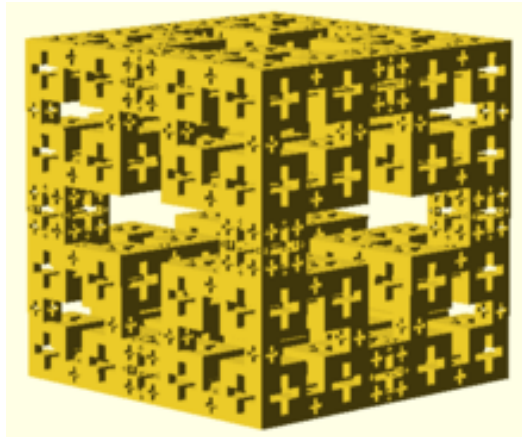


Abbildung 2.27: Der Jerusalem-Würfel

Schritten. wobei man sich nach einem Schritt nach rechts oder links wendet. Hier sind unterschiedliche Möglichkeiten vorhanden, die Anzahl der Schritte miteinander zu kombinieren, etwa einmal nach rechts, zweimal nach links, dreimal nach rechts usw. Wendet man beispielsweise ständig die Folge „nach oben, dann rechts“, „nach unten- dann links“ verschwindet die Kurve schnell im Unendlichen, man findet auch leicht Folgen, bei denen nur Quadrate oder Rechtecke entstehen usw. Ganz leicht lassen sich Fliesenmuster erzeugen.

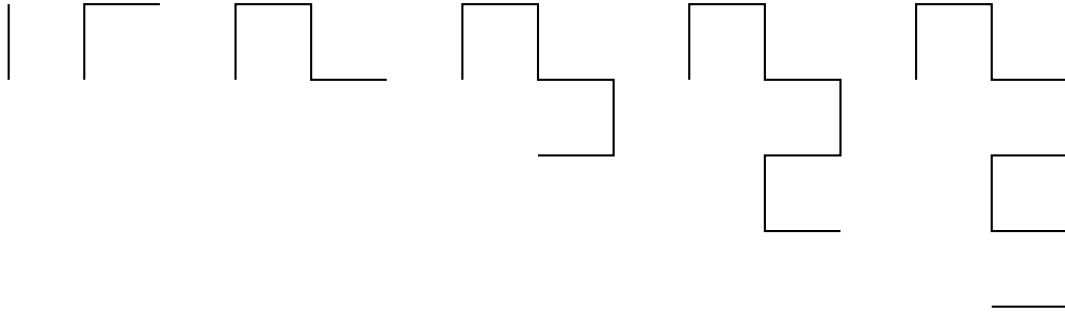


Abbildung 2.28: Drachenkurve

2.9 Das Wachstum von Fraktalen

Das *Wachstum von Aggregaten* ist einer der interessantesten Punkte, die man bei Fraktalen untersuchen kann. Im Alltag findet man viele Beispiele dafür [47]:

Das einfachste Wachstumsmodell ist das Eden-Modell. Schimmel an einer feuchten Wand ist ein Beispiel dafür. Moose und Flechten entwickeln sich nach dem gleichen Prinzip.

Es arbeitet nach folgendem Algorithmus.

- Das erste Teilchen des Aggregats wird im Koordinatenursprung eines gegebenen zwei- oder dreidimensionalen Gitters platziert und stellt den Keim des Aggregats dar.
- In den folgenden Schritten werden weitere Partikel an *zufällig ausgewählten Plätzen*, die an das im vorausgehenden Schritt gebildete Aggregat angrenzen, platziert. Man weiß also nicht, in welche Richtung sich die Struktur entwickelt. Wenn man die Punkte nach allen Richtungen mit gleicher Wahrscheinlichkeit auswählt, entsteht eine kreisförmige Struktur. Man kann aber auch gewisse Richtungen bevorzugen, wenn man die Wahrscheinlichkeitswerte entsprechend auswählt.
- Eden Cluster besitzen eine unregelmäßige Oberfläche.
- Zunächst vorhandene Hohlräume im Innern des Clusters werden im Laufe der Zeit aufgefüllt; die Struktur ist innen kompakt:
- Es liegt keine Skalenvarianz vor; die Struktur ist nur selbstaffin.

Die diffusions-limitierte Aggregation (DLA) ([48]) kann man zur Beschreibung

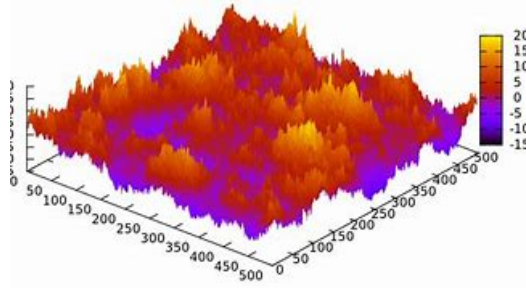


Abbildung 2.29: Ein Eden-Fraktal

der elektrolytischen Ausfällungen, der Verdrängung von Flüssigkeiten, elektrischer Durchbrüche (Blitze) und des Wachstums von Bakterienkolonien verwenden.

Die Gitterversion des DLA Modells arbeitet nach folgendem Algorithmus.

- Das erste Teilchen des Aggregats wird im Koordinatenursprung eines zwei- oder dreidimensionalen platziert und stellt den Keim des Aggregats dar.
- Weitere Teilchen werden an zufälligen Orten, die auf einem Kreis bzw. auf einer Kugelfläche liegen, freigesetzt. r_1 ist der Radius der Kugel.
- Die freigesetzten Teilchen führen anschließend stochastische Sprünge auf den Gitterplätzen aus und diffundieren durch das Gitter.
- Wenn diese Partikel während der Diffusion auf an das Aggregat angrenzende Gitterplätze gelangen, bleiben sie an dem Aggregat haften
- Die Anlagerung ist irreversibel, es handelt sich um eine starke Bindung.
- Es werden solange neue Teilchen freigesetzt, bis das Aggregat die gewünschte Größe erreicht hat.

DLA-Cluster sind selbstähnlich. Die fraktale Dimension des DLA-Clusters ist gleich 1,66.

Es gibt eine Reihe von Varianten dieses DLA-Modells.

1. Das DLA-Modell wird kombiniert mit einer reaktions-kontrollierten Anlagerung. Teilchen bleiben mit einer Anlagerungswahrscheinlichkeit $W < 100$ % am Aggregat haften. Man kann oft mit sehr geringen Wahrscheinlichkeiten operieren: $W = 10$ %, $W = 5$ %, $W = 1$ % ergeben ganz unterschiedliche Bilder. Bei unterschiedlichen Anlagerungswahrscheinlichkeiten W gewachsene

DLA-Cluster werden kompakter bei sinkender Anlagerungswahrscheinlichkeiten W . Teilchen können tiefer ins Zentrum des Clusters diffundieren, bevor sie endgültig eingefangen werden. Fraktale Dimension ist für $W > 0$ größer als 2.

2. Andere Randbedingungen: man betrachtet beispielsweise Teilchen, die von einer Seite einer Platte zu Aggregaten diffundieren, die auf der gegenüberliegenden, absorbierenden Plattenseite wachsen.

Die Agglomeration von Goldkolloiden wird verwendet, wenn sich Goldteilchen in einer geeigneten Lösung befinden. Die Agglomeration der Goldpartikel geschieht durch van der Waals-Wechselwirkungen. Die van der Waals-Kräfte sind schwache Wechselwirkungen, die zwischen verschiedenen Atomen bzw. Molekülen auftreten. Es sind zwischenmolekulare Kräfte. Bei der *van der Waals - Wechselwirkung* handelt es sich um keine echte Bindung. Sie besteht zwischen fast allen Atomen und Molekülen, wird aber hauptsächlich durch stärkere Bindungen überdeckt.

Zufallswege treten auf, wenn überhaupt keine Steuerung vorliegt. Für einen Prozess in der Ebene geht man in folgender Weise vor:

- Das Teilchen wird im Koordinatenursprung frei gesetzt.
- Das Teilchen wandert ziellos in willkürliche Richtungen auf dem Gitter zum nächsten Gitterpunkt.
- Am Gitterpunkt angelangt, wandert das Teilchen wiederum in einer willkürlich gewählten Richtung zum nächsten Gitterpunkt.
- Der letzte Schritt wird beliebig oft wiederholt.

Weitere Beispiele sind die Pollenkörner in der Luft, die im Frühjahr bei vielen Leuten zu Beschwerden führen, Rauchpartikel oder Staub in der Luft, Tusche im Wasser u.v.a.m.

Ganz typische Wachstumsprozesse, die durch Flüssigkeiten und leicht lösliche Mineralien hervorgerufen werden, sind **Stalaktiten** und **Stalagmiten** in unterirdischen Höhlen. Ein Stalaktit ist der von der Decke einer Höhle hängende Tropfstein, sein Gegenstück ist der vom Boden emporwachsende **Stalagmit**.

Wie alle Tropfsteinformationen entsteht der Stalaktit, wenn kohlensäurehaltiges Wasser in das Gestein eindringt und, bedingt durch die Oberflächenspannung an

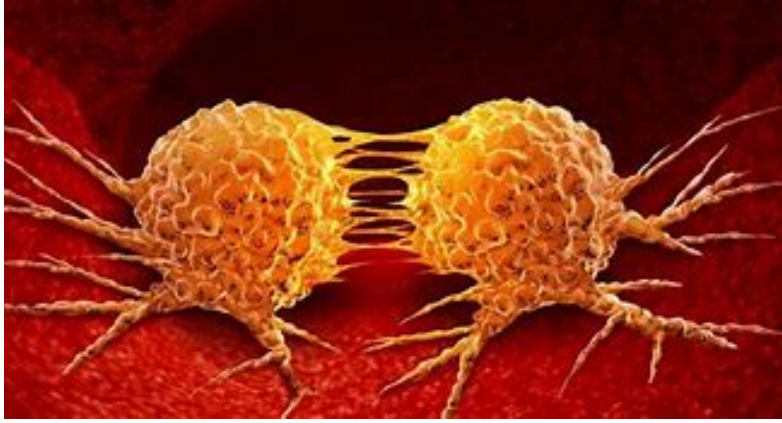


Abbildung 2.30: Das Wachstum von Krebszellen

der Decke eines Hohlraums, Calcit abgelagert. Das entstehende Material wird **Sinter** genannt. Tritt ein Tropfen auf Grund der Oberflächengestalt immer an exakt derselben Stelle aus, kann die Ablagerung die Form eines Ringes mit der Tropfengröße als Durchmesser bilden. Diese Formation kann zu einem Sinterröhrchen wachsen.



Abbildung 2.31: Der größte Stalaktit der Welt

Ein Stalaktit entsteht dadurch, dass Tropfen nicht mehr im Inneren des Röhrchens, sondern auf seiner Außenseite ablaufen und Kalzit abgelagern; das bewirkt Wachstum in die Breite. Bei der Untersuchung eines Stalaktiten kann das Sinterröhrchen im Inneren oft noch nachgewiesen werden. Es gibt auch seltenere Formen, bei denen ein Stalaktit ohne vorheriges Sinterröhrchen wächst, hier spielt die Oberflächengestalt der Höhlendecke die entscheidende Rolle, damit das Dickenwachstum ohne Leitstruktur beginnt. Die Größe eines Stalaktits ist durch sein Eigengewicht begrenzt, das ihn irgendwann von der Decke reißt. Ein Stalaktit ist immer deutlich schlanker als ein Stalagmit am Boden. Wächst der Stalaktit schließlich mit dem Stalagmit zusammen, bildet sich eine Tropfsteinsäule, Stalagnat genannt. Der mit mehr als sieben Metern größte bisher entdeckte Stalaktit Europas befindet sich in der Doolin Cave bei Doolin im County Clare, Irland. Er wurde in den 1950er-Jahren

von zwei Abenteuerlustigen entdeckt und im Jahr 2006 der Öffentlichkeit zugänglich gemacht. Stalaktiten und andere Tropfsteinformen können sich auch an älteren Bauwerken bilden, wenn Calciumhydroxid aus Zement oder Beton gelöst wird und dann mit dem Kohlendioxid der Luft reagiert.

Ein Stalagmit ist der vom Boden einer Höhle emporwachsende Tropfstein, sein Gegenstück ist der von der Decke hängende Stalaktit (Eselsbrücken siehe Tropfstein). Von Stalagnat spricht man, wenn beide Typen zusammengewachsen sind. Ein Stalagmit ist ein Speläothem, bei dem durch auftropfendes kohlensäurehaltiges Wasser, das Calcit ablagert, ein Tropfstein entsteht, der verschiedene Formen annehmen kann. Die Intensität der Tropfstelle, aber auch die Fallhöhe des Tropfwassers und die Bodenbeschaffenheit haben Einfluss auf die Form des Stalagmits. So unterscheidet man die gleichmäßig schlanken Kerzenstalagmiten und die kegelförmigen Palmstammstalagmiten. Die Kerzenstalagmiten entstehen durch eine gleichmäßige Lösungszufuhr und können bei geringem Durchmesser mehrere Meter Höhe erreichen. Die kegelförmigen Stalagmiten entstehen durch eine sehr starke Sickerwasserzufuhr und können an der Basis mehrere Meter Durchmesser aufweisen, wie zum Beispiel der Millionär in der Sophienhöhle.

Die Fallhöhe wiederum hat einen Einfluss auf das obere Ende des Tropfsteins. Ein gerundetes Ende entsteht bei geringer Fallhöhe des Wassers; bei zunehmender Fallhöhe wird das Ende flacher und kann im Extremfall nach innen gewölbt sein. Weiterhin gibt es Tropfsteinformen, die von diesen Grundwerten abweichen, wie etwa ein Stalagmit im Schulerloch, dessen Spitze als Wasserbecken ausgeformt ist, oder der Vesuv in der Eberstadter Tropfsteinhöhle, der durch kohlensäurehaltiges Wasser gebildet wurde, aber jetzt durch kohlensäurefreies Wasser langsam wieder abgetragen wird.



Abbildung 2.32: Die Feengrotten bei Saalfeld

Ein ganz anderes Beispiel für Anlagerungen sind Puzzles. Hier geht es darum,

bestimmte kleine Steine zu einem Bild zusammenzusetzen. Die Selbstähnlichkeit angefangener Bilder liegt auf der Hand. Nach dem Setzen eines Steines sieht das Bild so wie vorher aus, man sieht selten einen markanten Unterschied. Die vier Ecksteine eines Bildes findet man mit hundertprozentiger Sicherheit. Danach sucht man alle Randsteine, was auch ein lösbares Problem; danach beginnt dann die Suche nach geeigneten Anlagerungen am Rand, eventuell kann man auch gleichfarbige Inseln bilden, die man irgendwo anhängen kann usw.

2.10 Brownsche Bewegungen

Die **Brownsche Bewegung** ist die vom schottischen Botaniker Robert Brown (1773 - 1858) im Jahr 1827 unter dem Mikroskop entdeckte unregelmäßige und ruckartige Wärmebewegung kleiner Teilchen in Flüssigkeiten und Gasen. Nach der 1905 von Albert Einstein und 1906 von Marian Smoluchowski gegebenen Erklärung wird die im Mikroskop sichtbare Verschiebung der Teilchen dadurch bewirkt, dass die unsichtbaren Moleküle der Umgebung eine Wärmebewegung ausführen und auf Grund dieser ungeordneten Wärmebewegung ständig und aus allen Richtungen in großer Zahl gegen die beobachteten Teilchen stoßen und dabei rein zufällig manchmal die eine Richtung, manchmal die andere Richtung stärker zum Tragen kommt. Diese Vorstellung wurde in den folgenden Jahren durch die Experimente und Messungen von Jean Baptiste Perrin (1870 - 1942) quantitativ bestätigt. Die erfolgreiche Erklärung der Brownschen Bewegung gilt als Meilenstein auf dem Weg zum wissenschaftlichen Nachweis der Existenz der Moleküle und damit der Atome.

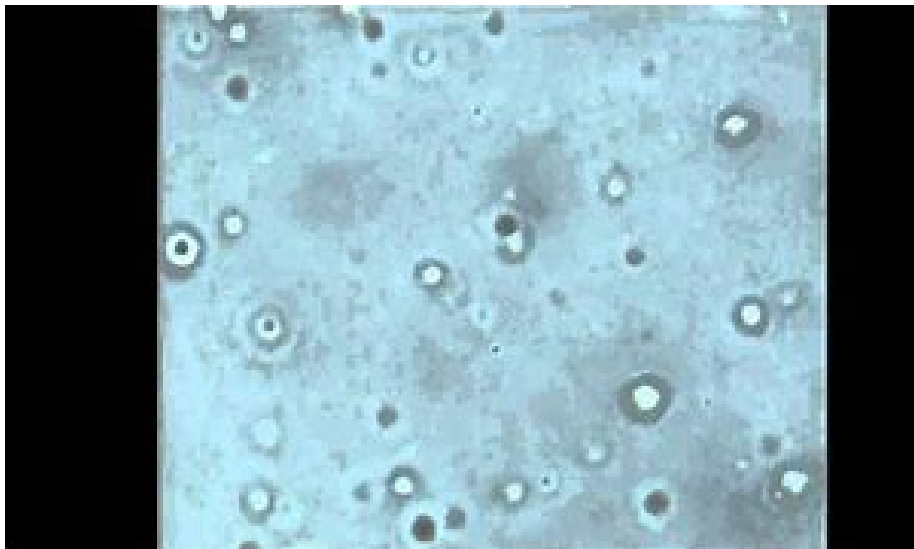


Abbildung 2.33: Milch auf kaltem Wasser

Die in DLA-Prozessen gebildeten Cluster werden als **Brownsche Bäume** bezeichnet.

In 2D weisen diese Fraktale eine Dimension von etwa 1,71 für freie Teilchen auf, die nicht durch ein Gitter eingeschränkt sind. Die Computersimulation von DLA ist ein wichtiges Mittel zur Untersuchung dieses Modells. Hierfür stehen mehrere Methoden zur Verfügung. Simulationen können auf einem Gitter mit beliebiger Geometrie und Einbettungs-Dimension durchgeführt werden (dies wurde in bis zu 8 Dimensionen durchgeführt), oder die Simulation kann eher nach dem Muster einer Standard-Molekulardynamik-Simulation durchgeführt werden, bei der ein Teilchen sich frei bewegen kann, bis es in einen bestimmten kritischen Bereich kommt, woraufhin es auf den Cluster gezogen wird. Von entscheidender Bedeutung ist, dass die Anzahl der Teilchen, die sich im System in Brownscher Bewegung befinden, sehr gering gehalten wird, damit nur die diffusive Natur des Systems zum Tragen kommt.

Ein Brownscher Baum wird in folgenden Schritten aufgebaut: Zunächst wird ein Samen irgendwo auf dem Bildschirm platziert. Dann wird ein Teilchen an einer zufälligen Stelle des Bildschirms platziert und nach dem Zufallsprinzip bewegt, bis es gegen den Keim stößt. Das Teilchen wird dort belassen, und ein weiteres Teilchen wird an einer zufälligen Position platziert und so lange bewegt, bis es gegen den Keim oder ein vorheriges Teilchen stößt, und so weiter. Der entstehende Baum kann viele verschiedene Formen haben, die im Wesentlichen von drei Faktoren abhängen:

- von der Position des Keimlings;
- von der Anfangsposition des Partikels (irgendwo auf dem Bildschirm, in einem Kreis um den Keim herum, am oberen Rand des Bildschirms, ...);
- vom Bewegungsalgorithmus;

Die Farbe der Partikel kann sich zwischen den Iterationen ändern, was zu interessanten Effekten führt. Man findet im Internet etliche Programme, mit denen diese Vorgehensweise simuliert wird, wobei man eine ganze Reihe von Parametern selbst einstellen kann (beispielsweise <https://www.leifiphysik.de/atomphysik/atomaufbau/>).

In der Medizin spielt die Verengung oder der Verschluss von Gefäßen und Organen eine große Rolle. Hier kann man sehr deutlich beobachten, wie das Anlagern von einzelnen negativen Elementen (beispielsweise von Kohlenstoffpartikeln, die beim Rauchen eingeatmet werden) zu einer zunehmenden Einengung von wichtigen Gefäßen führen kann, die letztendlich immer lebensgefährlich sind.

Besonders anfällig für derartige Verstopfungen sind die unterirdischen Abwassersysteme, wenn sie nicht ständig gereinigt und kontrolliert werden (Abb. 2.35).

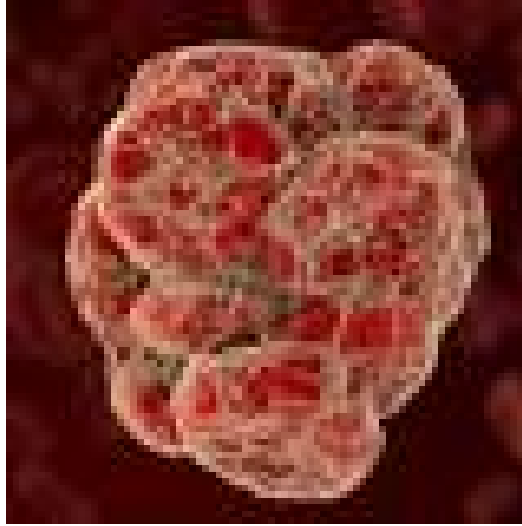


Abbildung 2.34: Verklumpung im Inneren eines Körpers



Abbildung 2.35: Der Fettberg in einem Londoner Abwasserkanal

Einer der größten Fettberge aller Zeiten wurde 2017 in Londons Untergrund gefunden. Er wog 130 Tonnen (Abb. 2.35).

Ähnlich geht die Anlagerung von Rußteilchen vor sich, bekannt aus allen Gegenden, wo Braunkohle-Kraftwerke arbeiteten: Teilchen lagern sich an den Wänden eines Kamins, an Häuserwänden, Fensterscheiben und Hemdkragen an. Andere Beispiele sind

- Ausfällungen in elektrolytischen Lösungen, z. B. Kupfersulfatlösung, die man durch geeignete Kathoden zu reinem Kupfer ausfällen lässt,
- baumartige Strukturen in der Biologie, wie Bildung der Fellzeichnungen bei Zebras, Tigern, Leoparden oder Tapiren.
- Lichtenberg-Figuren (Abb. 2.36) entstehen typischerweise durch die elektrostatische Entladung bzw. Umverteilung von auf der Oberfläche von Isolatorplatten befindlichen elektrischen Ladungen. Die der Bildung der Lichtenberg-Figuren zugrunde liegenden physikalischen Prinzipien sind dieselben, auf die

sich auch die moderne Elektrofotografie (Xerografie) gründet, die in allen heute gängigen Kopiergeräten (Fotokopierer, Laserdrucker etc.) eingesetzt wird (Abb. 2.37).



Abbildung 2.36: Eine Lichtenberg - Figur



Abbildung 2.37: Ladungen zwischen zwei Leitern

2.11 Risse und Bruchstellen

Das Brechen von festen Körpern wird meistens durch eine äußere Belastung erzeugt (Abb. 2.38). Oft sind Alterungsprozesse oder eine fortgeschrittene Korrosion daran schuld. Dafür kann aber auch eine Belastung über eine zulässige oder verträgliche Grenze hinaus verantwortlich sein. Relativ hilflos ist der Mensch gegen die Kräfte, die bei Erdbeben auftreten. Für Hochhäuser in durch Erdbeben gefährdeten Gebieten hat man schon viele technologische Verbesserungen gefunden, die eventuelle Schäden verringern oder vermeiden. Zerbrochene Glasscheiben gehören beinahe schon zum Alltag. Oft treten auch Einsturzgefahren auf, wenn in einem bestimmten Gebiet lange Zeit Untertage-Bergbau betrieben wurde.

Beim Fracking wird durch eine Bohrung, unter hohem Druck von typischerweise mehreren hundert Bar, eine Flüssigkeit (Fracfluid) in den geologischen Horizont, aus dem gefördert werden soll, gepresst. Als Fracking-Flüssigkeit dient Wasser, das



Abbildung 2.38: Eine zerbrochene Eisdecke

zumeist mit Stützmitteln, wie z. B. Quarzsand, und Verdickungsmitteln versetzt ist. Üblicherweise werden zunächst im Zielhorizont mehrere abgelenkte Bohrungen (Laterale) mittels Richtbohren angelegt, wobei der Bohrmeißel schichtparallel geführt wird. Dadurch ist die zur Verfügung stehende Bohrlochlänge in der Lagerstätte wesentlich größer, was generell die Ausbeute der Förderung erhöht. Zum Einsatz kommen beim Hochvolumen-Hydrofracking große Flüssigkeitsmengen mit mehr als $1000m^3$ pro Frackphase bzw. insgesamt mehr als $10.000m^3$ pro Bohrloch.

Seit Ende der 1940er Jahre wird Fracking vor allem bei der Erdöl- und Erdgasförderung sowie bei der Erschließung tiefer Grundwasserleiter für die Wassergewinnung und der Verbesserung des Wärmetransportes bei der tiefen Geothermie eingesetzt. In den letztgenannten Anwendungsfällen werden keine Stützmittel oder chemischen Zusätze benötigt. Seit Anfang der 1990er Jahre und insbesondere in den USA ab etwa dem Jahr 2000 fokussiert sich die Förderung mittels Fracking auf sogenanntes unkonventionelles Erdöl und Erdgas. Der dortige Fracking-Boom veränderte den US-Energiemarkt erheblich und sorgte für einen Preisverfall. Dies führte zu einer Debatte über die Rentabilität des Verfahrens. Die US-Regierung unterstützt daher seit etwa 2013 Bestrebungen zum verstärkten Export von Flüssigerdgas nach Europa und Japan, unter anderem mit beschleunigten Genehmigungsverfahren.

Während einige Stimmen diese geostrategische Komponente durch die Veränderung der internationalen Abhängigkeiten betonen, führen die Umweltrisiken und mögliche Gesundheitsgefahren des „Fracking-Booms“ vor allem in Europa zu einer kontrovers geführten und noch andauernden fachlichen, politischen und gesellschaftlichen Debatte. Einige Länder und Regionen haben Erdgas-Fracking auf ihrem Gebiet gesetzlich verboten.

Risse und Brüche kann man auf unterschiedlichen Ebenen untersuchen. Auf dem Niveau von Elektronen oder der Verschiebung von Körner-Positionen ($\leq 10^{-6}$ cm) sind die Materialwissenschaft bzw. die Werkstoffwissenschaft zuständig. Die Begriffe Materialwissenschaft und Werkstoffkunde sind eng miteinander verknüpft:

Die Materialwissenschaft mit einer eher naturwissenschaftlich geprägten Herangehensweise beschäftigt sich mit der Herstellung von Materialien und deren Charakterisierung von Struktur und Eigenschaften, während die Werkstofftechnik die ingenieurwissenschaftlich orientierte Werkstoffentwicklung sowie die entsprechenden Verarbeitungsverfahren und das Betriebsverhalten von Bauteilen im Einsatz beinhaltet. Beide Teilgebiete umfassen Forschungsaktivitäten der verschiedensten Materialklassen und Werkstoffentwicklungsketten.

Ein wesentliches Merkmal der Materialwissenschaft und Werkstofftechnik ist die Berücksichtigung des strukturellen Aufbaus der Werkstoffe und der davon abhängigen mechanischen, physikalischen und chemischen Eigenschaften. Dies umfasst die Charakterisierung, Entwicklung, Herstellung und Verarbeitung von Konstruktionswerkstoffen und Funktionsmaterialien.

Das Fachgebiet setzt sich aus der erkenntnis-orientierten Grundlagenforschung zu Materialien und der ingenieurwissenschaftlichen Werkstoffentwicklung mit Anwendungsbezug zusammen. Es entfaltet dabei eine starke Hebelwirkung im Sinne einer Umsetzung von Forschungsergebnissen in marktrelevante Innovationen. Gleichzeitig haben Materialwissenschaft und Werkstofftechnik als interdisziplinäre Wissenschaft eine weitreichende fachliche Integrationswirkung, in dem sie Erkenntnisse aus benachbarten Fachgebieten aufgreifen und mit ihnen in wechselseitiger Beziehung stehen. Für die Materialwissenschaft sind hier insbesondere die Verknüpfungen mit der Chemie, der Physik und den Lebenswissenschaften zu nennen, während für die Werkstofftechnik die Gebiete Mechanik, Konstruktionstechnik, Produktionstechnik und Verfahrenstechnik relevant sind.

Sind die betrachteten Größenordnungen $> 10^{-1}$ cm, dann ist die Vermeidung von Brüchen und ähnlichen Schäden ein Gegenstand der Ingenieurwissenschaften. In den dazwischen liegenden Größenordnungen spielt die **Angewandte Mechanik** eine Hauptrolle (Abb. 2.40).

Brüche, Risse, Zersplitterungen, Zusammenstöße u.a. weisen meistens fraktale Strukturen auf.

Risse in der Straße werden kontinuierlich durch eintretendes Oberflächenwasser ausgewaschen. Die Straße verliert an Substanz, der zusätzliche Wechsel zwischen Frost und getautem Boden führt zu einer Auflockerung des gesamten Gefüges. Daher führen offene Risse unweigerlich zu größeren Schäden (Abb. 2.41).

Die geologischen Prozesse auf der Erde liefern viele fraktale Bilder: der Zusammenstoß der Kontinentalplatten, die aufgeworfenen Gebirge, die Vulkane entlang der Stoßkante haben alle eine fraktale Struktur (Abb. 2.42 und Abb. ??).



Abbildung 2.39: Ein Riss in einer alten Mauer



Abbildung 2.40: Ein Riss in einem Autoreifen

2.12 Fraktale Dimensionen

Bei Wikipedia findet man eine Tabelle vorhandener Hausdorff-Dimensionen, die hier auszugsweise wiedergegeben werden soll.

1. Boxcounting-Dimension

Man überdeckt die Menge mit einem Gitter der Breite ϵ . Benötigt man $N(\epsilon)$ Kästchen, so ist

$$\lim_{n \rightarrow \infty} \frac{\log N(\epsilon)}{\log \frac{1}{\epsilon}} \quad (2.15)$$

die Box-Dimension der gegebenen Menge.

Ist die horizontale Entfernung der Gitterpunkte gleich ϵ , dann hat das überdeckende Quadrat einen Flächeninhalt von $2\epsilon^2$ (Abb. 2.43). Die Verwendungen Kugeln mit dem Radius ϵ funktioniert im dreidimensionalen Fall in gleicher Weise.



Abbildung 2.41: Risse in einer Straße

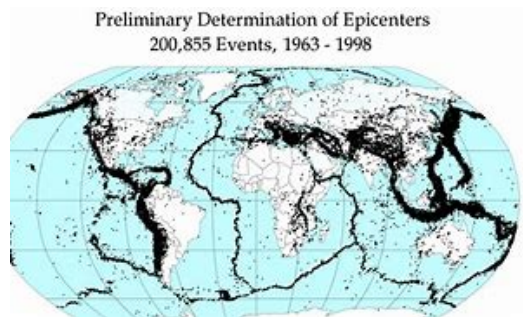


Abbildung 2.42: Der Zusammenstoß von Kontinentalplatten

2. Yardstick-Methode Diese Methode eignet sich nur für Kurven. Man wählt einen kleinen Wert l und einen Startpunkt. Diesen Startpunkt wählt man als Mittelpunkt eines Kreises mit dem Radius r . Dieser Kreis schneidet die Kurve, und man verbindet diesen Schnittpunkt mit dem vorhergehenden Punkt.

Mengen auf einem Bildschirm sind immer endlich, dort muss man immer die Hausdorff-Dimension verwenden. Ein Grenzwert würde ja stets den Wert 0 ergeben.

Andererseits kann man die Hausdorff-Dimension durch die computerisierte Bildverarbeitung noch genauer bestimmen. Setzt man die Pixelmenge des Bildschirms als Fläche auf den Wert 2, dann kann mit Hilfe eines Programmes die Anzahl der Pixel auf und über der Kurve genau zählen und ins Verhältnis zur Gesamtmenge

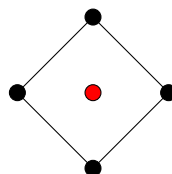


Abbildung 2.43: Die Überdeckung eines Punktes

setzen. Der entstehende Wert liegt immer zwischen 1 und 2 und gibt Auskunft über die Eigenschaften des fotografierten Objektes.

3 Fraktale in der Medizin

Die Gehirne kleiner Säugetiere sind relativ glatt, die von Menschen dagegen höchst faltenreich. Eine fraktale Dimension zwischen 2,73 und 2,79 scheint für das menschliche Gehirn typisch zu sein. Auch in den Membranen von Leberzellen findet man fraktale Strukturen. Die Nasenknochen von Hirschen und Polarfüchsen sorgen für maximale Geruchsempfindlichkeit, indem sie die größtmögliche Oberfläche in ein kleines Volumen packen. Daraus ergibt sich eine fraktale Struktur mit konstanter gebrochener Dimension. Die Blutversorgung des Menschen verzweigt sich wenigstens achtmal, die maximale Anzahl liegt etwa bei 30. Damit werden alle Körperzellen erreicht; die fraktale Dimension ist gleich drei. Fraktale Selbstähnlichkeit durchzieht die Körper der Organismen, aber es ist nicht die platte homunculus-artige Selbstähnlichkeit, die sich die frühere Wissenschaft vorgestellt hatte. Der Körper ist eine Vernetzung von lauter selbstähnlichen Systemen wie den Lungen, den Gefäßsystemen, den Nervensystemen.

3.1 Enzephalographie

Auch Gehirnströme folgen fraktalen Mustern (Abb. 3.1). Erhöhte EEG-Synchronität führt zur Verringerung der fraktalen Dimension. Dies wurde bei epileptischen Anfällen, die vermehrte synchrone Oszillationen und eine Verringerung der arrhythmischen Aktivität aufweisen, und unter Narkose beobachtet. Im Gegensatz dazu wird in manischen Phasen einer bipolaren Störung und in den Frontallappen von Schizophrenie-Patienten eine erhöhte Komplexität beobachtet. Sowohl die bipolare Störung als auch die Schizophrenie zeichnen sich durch eine gestörte synchrone Aktivität aus.

Epilepsie ist eine der häufigsten Erkrankungen des Gehirns. In der Klinik werden zur Bestätigung eines epilepsieähnlichen Symptoms und zur Lokalisierung des Anfalls die Daten des Elektroenzephalogramms (EEG) häufig von einem Arzt visuell untersucht, um das Vorhandensein epileptischer Entladungen zu erkennen. Epileptische Entladungen sind flüchtige Wellenformen, die mehrere zehn bis hundert Millisekunden andauern und hauptsächlich in sieben Typen unterteilt werden [5].

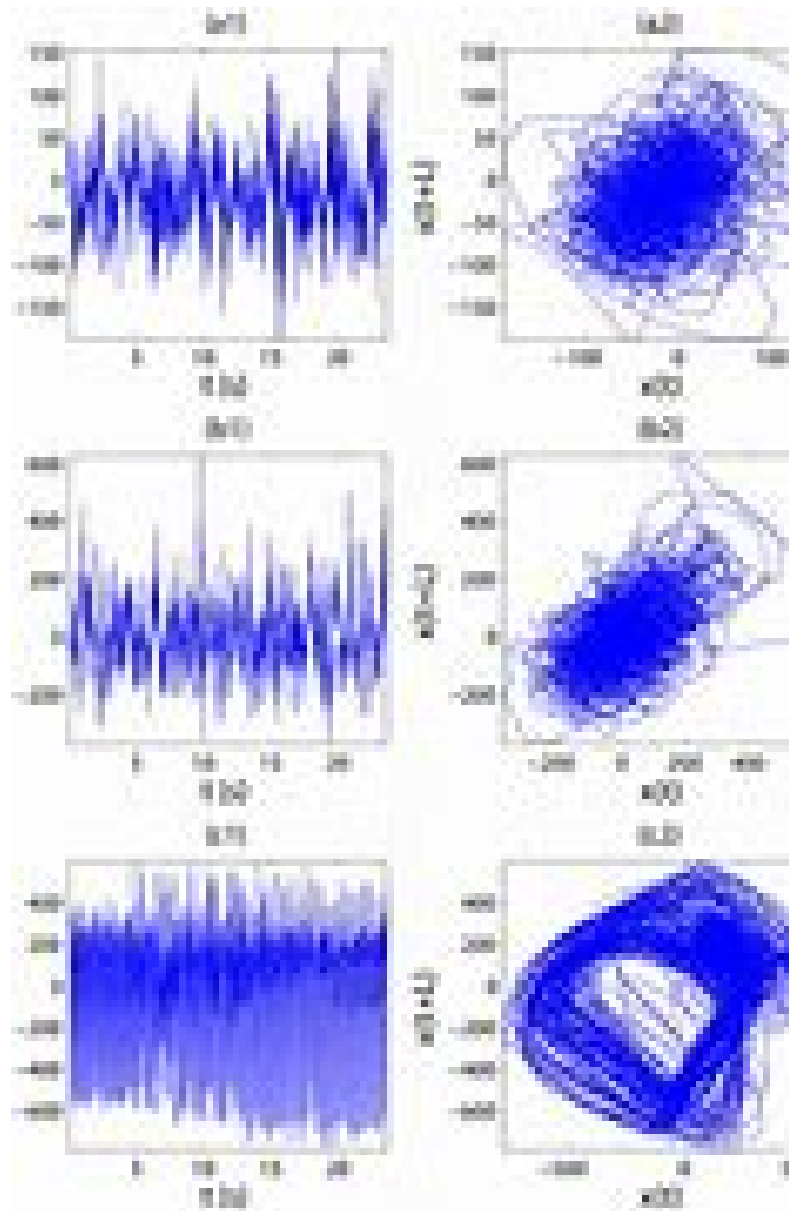


Abbildung 3.1: Aufzeichnungen eines EEG-Gerätes

Es ist wichtig, systematische Ansätze zu entwickeln, um diese Wellenformen genau von normalen Kontrollwellen zu unterscheiden. Dies ist eine schwierige Aufgabe, da klinisch verwendete Kopfhaut-EEGs in der Regel viel Rauschen und Artefakte enthalten. Um dieses Problem zu lösen, wurden 640 mehrkanalige EEG-Segmente mit einer Länge von jeweils 4 Sekunden analysiert. Von diesen Segmenten sind 540 kurze epileptische Entladungen, und 100 stammen von gesunden Kontrollpersonen. Es wurden zwei Ansätze vorgeschlagen, um epileptische Entladungen von normalen EEGs zu unterscheiden. Die erste Methode basiert auf der Signalreichweite und den Korrelationseigenschaften der EEGs. Die zweite Methode basiert auf Netzwerken, die aus drei Aspekten der Kopfhaut-EEG-Signale konstruiert werden: der Signalreichweite, der Energie der Alphawellenkomponente und den Langzeitkorrelationseigenschaften des EEG. Die Netzwerke werden mit Hilfe der Singulärwertzerlegung

(SVD) weiter analysiert.

Die in dieser Studie analysierten EEG-Daten stammen aus dem First Affiliated Hospital der Medizinischen Universität Guangxi. Die Studien mit menschlichen Teilnehmern wurden von der Ethikkommission des First Affiliated Hospital der Medizinischen Universität Guangxi geprüft und genehmigt. Die Teilnehmer gaben ihre schriftliche Einwilligung zur Teilnahme an dieser Studie. Neunundfünfzig Epilepsiepatienten unterzogen sich einer dreistündigen Video-EEG-Überwachung mit 19-Kanal-EEG-Aufzeichnung mit auf der Kopfhaut angebrachten Elektroden nach dem internationalen 10-20-System bei einer Abtastrate von 256 Hz. Die Elektrodenimpedanzen wurden unter 10 Kil Ω gehalten. Die 19 elektroenzephalographischen Kopfhautelektroden waren entsprechend den Bezeichnungen Fp1, Fp2, F7, F3, Fz, F4, F8, T3, C3, Cz, C4, T4, T5, P3, Pz, P4, T6, O1 und O2 angeordnet.

Alle epileptischen Entladungen wurden von einem erfahrenen klinischen Neurophysiologen auf der Grundlage der durchschnittlichen Montage mit einer analogen Bandbreite von 0,1 bis 70 Hz und einem Kerbfilter von 50 Hz kommentiert. Die EEG-Signale wurden in 4s-Epochen segmentiert und für jeden Teilnehmer mit Zufallszahlen versehen. Die gesammelten Epochen wurden für die weitere Analyse in das Europäische Datenformat (EDF) umgewandelt. Insgesamt wurden von allen Teilnehmern 532 EEG-Aufzeichnungen epileptiformer Entladungen und 100 gesunde Kontrollpersonen mit einer Dauer von jeweils 4 Sekunden erfasst. Unter den 532 kurzen epileptischen Entladungen befanden sich 69 Spikes, 82 Sharps, 174 Spike- und Slow-Wave-Komplexe, 72 scharfe und Slow-Wave-Komplexe, 64 Polyspike-Komplexe, 77 Polyspike- und Slow-Wave-Komplexe und zwei rhythmische Spike-Entladungen.

- **Spike:** Die Spikes sind die grundlegendste EEG-Aktivität mit einer Dauer von 20 bis 70 ms. Die Amplitude variiert und ist typischerweise $>50 \mu V$.
- **Sharp:** Eine scharfe Welle ähnelt einem Spike und hat eine Dauer von 70 bis 200 ms (5 bis 14 Hz). Ihre Amplitude variiert zwischen 100 und 200 μV , und die Phase ist normalerweise negativ.
- **Spike und langsamer Wellenkomplex:** Ein epileptisches Muster, das aus einem Spike und einer damit verbundenen langsamen Welle nach dem Spike besteht, die sich deutlich von der Hintergrundaktivität abhebt; es kann einfach oder mehrfach auftreten.
- **Komplex aus scharfen und langsamen Wellen:** Ein epileptisches Muster, das aus einer scharfen Welle und einer damit verbundenen langsamen Welle

besteht, die auf die scharfe Welle folgt und sich deutlich von der Hintergrundaktivität abhebt; sie kann einmal oder mehrfach auftreten.

- **Polyspike-Komplex:** Eine Sequenz von zwei oder mehr Spikes.
- **Polyspike and slow wave complex:** Ein epileptisches Muster, das aus zwei oder mehr Spikes besteht, die mit einer oder mehreren langsamen Wellen verbunden sind.
- **Spike-Rhythmus:** Bezieht sich auf einen weit verbreiteten Ausbruch eines Spike-Rhythmus mit 10 bis 25 Hz, mit einer Amplitude von 100 bis 200 μV und der höchsten Spannung im frontalen Bereich, der mehr als 1 s andauert.

Oft hängen epileptische Entladungen mit einer Amplitude zusammen, die größer ist als die eines normalen Kontroll-EEGs. Der Signalbereich wird wie folgt berechnet:

$$\max x(t), t \in [t_1, t_2] - \min x(t), t \in [t_1, t_2], \quad (3.1)$$

Dieses Verfahren wird auf jedes der 19 EEG-Signale in Bezug auf die Ohrläppchen angewandt (d. h. auf die Differenz der an den 19 Elektroden und den Ohrläppchen gemessenen EEG-Signale) oder auf die Differenz der EEG-Signale gemäß der Netzwerkkonstruktion. Im ersten Fall wird der endgültige Signalbereich als Durchschnitt der 10 größten Signalbereiche geschätzt, die aus den 19 EEG-Signalen ermittelt wurden.

In klinischen Anwendungen werden die Gehirnwellen häufig in fünf Bänder eingeteilt: Delta (0,5 bis 3 Hz), Theta (4 bis 7 Hz), Alpha (8 bis 13 Hz), Beta (14 bis 30 Hz) und Gamma (> 30 Hz). Die Alphawelle ist am deutlichsten sichtbar, wenn der Mensch entspannt ist und die Augen geschlossen hat. Es wurde festgestellt, dass die Alphawellenkomponente im Bereich des Hinterkopfes bei epileptischen Entladungen oft größer ist. Zur Berechnung dieser Komponente verwendet man eine Fourier-Transformation des EEG-Signals; daraus berechnet man die Leistungsspektraldichte (PSD) und integriert schließlich die PSD-Kurve über das Alphawellenband.

Die Adaptive Fractale Analyse (AFA) benutzt einen adaptiven Detrending-Algorithmus, um global glatte Trendsignale aus den Daten für eine bestimmte Zeitskala zu extrahieren, und analysiert dann die Skalierung der Residuen der Anpassung als Funktion der Zeitskala [6],[7]. Die Hauptschritte der AFA zur Schätzung von H sind wie folgt:

Es wird angenommen, dass ausgehend von einem stationären inkrementellen Prozess $x(1), x(2), x(3), \dots$ ein Random Walk durch die folgende Gleichung konstruiert

wird:

$$u(n) = \sum_{k=1}^n (x(k) - \bar{x}), n = 1, 2, 3, \dots N. \quad (3.2)$$

wobei $\bar{x}$ der Mittelwert des Prozesses ist.

Ausgehend von diesem Random Walk $u(n)$ soll ein globaler Trend $v(i), i = 1, 2, \dots N$ für eine beliebige Zeitskala w ermittelt werden, wobei N die Länge der ursprünglichen Zeitreihe ist. Dies wird erreicht, indem der obige Random-Walk-Prozess in überlappende Fenster unterteilt wird, wobei die Größe jedes Fensters w eine ungerade Anzahl von Stichproben enthält und benachbarte Fenster sich um $(w + 1)/2$ Stichproben überlappen. Der Random-Walk-Prozess in jedem Fenster wird durch ein Polynom der Ordnung M angepasst, und die Polynome in überlappenden Regionen werden kombiniert, um einen einzigen globalen Trend zu erhalten. Typischerweise sollte M gleich 1 oder 2 sein, eine lineare oder quadratische Funktion. Die lokale Anpassung stellt sicher, dass der globale Trend optimal oder nahe am Optimum ist, da eine lokale Taylorreihen-Entwicklung verwendet wird.

Nachdem man den globalen Trend $v(i)$ von $u(i)$ durch die obige Methode erhalten hat, kann das Residuum $u(i) - v(i)$ die Schwankung um den globalen Trend beschreiben. Für fraktale Prozesse kann der Hurst-Exponent H durch die folgende Gleichung berechnet werden:

$$F^{(2)}(w) = \frac{1}{N} \sqrt{\sum_{i=1}^N (u(i) - v(i))^2} \approx w^H. \quad (3.3)$$

Die obige Gleichung bedeutet, dass man durch Berechnung der Varianz des Residuums zwischen dem ursprünglichen Random-Walk-Prozess und dem angepassten globalen Trend unter einem variierenden Fenster w eine lineare (oder mehrfach lineare) Beziehung zwischen $\log_2 F(w)$ und $\log_2 w$ erhalten kann.

Auf diese Ausgangssituation wird eine ganze Reihe von Algorithmen angewendet, im Ergebnis dessen erhält man eine ganze Reihe von Parametern, die eine genauere Einschätzung der Situation erlaube. Für diese Überlegungen sei auf den Artikel [5] verwiesen.

3.2 Das Herz

Schon Leonardo da Vinci beobachtete, dass die Wände der menschlichen Herzkammern innen mit einem komplexen Geflecht aus Herzmuskelfasern ausgekleidet

sind. Anders als bei den großen Blutgefäßen ist die Wand der Ventrikel dadurch nicht glatt, sondern eher schwammartig. Die Trabekel bilden eine vielfach verzweigte Übergangsschicht zwischen dem massiven Herzmuskel und dem Innenraum der Kammern. Sie spielen beim frühen Embryo eine wichtige Rolle für die Versorgung des Herzmuskels. Bevor die Herzkranzgefäße voll ausgebildet und durchblutet sind, ermöglicht die große Oberfläche der Trabekel die Diffusion von Nährstoffen und Sauerstoff aus dem Blut in den Herzmuskel.

Die Muskelfaser-Netzwerke an den Wänden der Herzkammern sind fraktal aufgebaut, sie beeinflussen die Herzleistung. Je komplexer die Verzweigungen der Trabekel sind, desto besser pumpt das Herz. Eine höhere fraktale Dimension ist mit einem größeren Herzschlagvolumen, einer stärkeren Pumpleistung und einem besseren Herzzustand verknüpft.

Der Herzschlag des Menschen folgt einem fraktalen Rhythmus. Werden Herzschlag und Atemrhythmus zu regelmäßig, kann das zu Herzversagen durch Stauung führen. Wird der Rhythmus allerdings zu unregelmäßig, verursacht dies das Flimmern eines Herzanfalles. In der Praxis schwankt der Rhythmus also ein Leben lang im Grenzbereich zwischen Chaos und Ordnung hin und her.

Man identifizierte 16 Gen - Orte, die eng mit den kardialen Muskelfaserstrukturen und ihrer Form verknüpft sind. Zehn dieser Genorte kontrollieren zusätzlich noch andere Aspekte der Herzfunktion, wie die Pulsrate, die Struktur der linken Herzkammer und die Dauer eines Ausschlagskomplexes im EKG [6].



Abbildung 3.2: Die fraktale Struktur des menschlichen Herzens

3.3 Das Auge

Netzhaut-Implantate mit fraktaler Geometrie könnten geeignet sein, Patienten zu helfen, die aufgrund von Netzhaut-Erkrankungen ihre Sehkraft verlieren. Das zeigen Forscher um den Physikprofessor Richard Taylor von der University of Oregon. Demnach würden fraktale Implantate gesunde Neuronen in der Netzhaut wesentlich besser stimulieren als jene mit normalen, flachen Elektroden. So sollte es den Forschern zufolge erstmals möglich werden, die Sehkraft soweit wiederherzustellen, dass sich Betroffene in Räumen und auch auf der Straße zurechtfinden.



Abbildung 3.3: Netzhautgefäße - fraktal strukturiert

Netzhauterkrankungen wie die Makuladegeneration bedrohen das Augenlicht von vielen Menschen. Abhilfe könnte die Nutzung von Implantaten bringen, die noch gesunde Nervenzellen stimulieren. Das Team um Physikprofessor Richard Taylor zielt auf einen neuen Ansatz: die Veränderung der Geometrie der Implantate, von einfachen euklidischen Formen wie Quadraten hin zu fraktalen Strukturen, die letztlich auch jenen der natürlichen Neuronen besser entsprechen.

Computersimulationen deuten nun darauf hin, dass sich damit wirklich drastische Verbesserungen erzielen lassen. Unter idealen Bedingungen in der Simulation konnte man zeigen, dass eine einzelne fraktale Elektrode alle Zielneuronen stimuliert, während die euklidische Elektrode sich nur mit zehn Prozent verbindet. Gleichzeitig kommt die fraktale Elektrode mit einer weniger als halb so hohen Spannung aus.

So vielversprechend die Simulations-Ergebnisse auch sind, bis womöglich tatsächlich fraktale Implantate Blinden eine Wahrnehmung zurückgeben könnten, ist es noch ein weiter Weg. Die in den Simulationen gewonnen Daten werden Taylor zufolge helfen, zunächst Miniaturversionen realer Implantate zu entwickeln. Diese sind für Tests an Mäusen bestimmt. Bis zu klinischen Tests an Menschen oder gar einer allgemeinen Zulassung kann es also noch lange dauern.

3.4 Einige Schlussfolgerungen

Krankheiten stellen für den Körper einen chaotischen Zustand dar. Durch fraktale Bedingungen, die dieses Chaos wieder in einen stabilen Zustand bringen, kann eine Gesundung erfolgen. Dies geschieht beispielsweise durch die Anwendung von **fraktalen Schwingungsmustern**. Ein instabil gewordenes System kann auf diese Weise zum gewünschten Stabilitätsbereich zurückgeführt werden.

Welches Schwingungsmuster für die Gesundung jeweils erforderlich ist, kann mithilfe einer speziellen mathematischen Formel bestimmt werden. Im Allgemeinen beschreiben diese mathematischen Formeln stabile Bahnbereiche als Attraktoren für mechanische Körperbewegungen; hierbei handelt es sich um Bahnen, die sich im dreidimensionalen geometrischen Raum befinden. In Analogie dazu können für biologische Systeme stabile Zustände oder Bereiche definiert werden, die nicht für den geometrischen Raum gelten, sondern bspw. für die Dimensionen **Frequenz**, **Phase** und **Zeit**.

Wird beispielsweise die Einnahme von Medikamenten nach diesem fraktalen Formalismus sowohl zeitlich als auch mengenmäßig untersucht, zeigt sich, dass über die Zeit hinweg die Wirkung steigt, sodass die Dosis reduziert werden kann.

Bei der Behandlung mit Schwingungsmustern erscheint es ratsam, in Abhängigkeit von der erwünschten therapeutischen Wirkung zwischen Frequenzmustern der auf- und abbauenden Wirkungen zu unterscheiden. Aufbauende Wirkungen sind z.B. Zusammenhalt und Koordination von Körpersystemen und im Körper befindlichen Kommunikationssystemen. Abbauende Wirkungen sind bspw. die Ausleitung von Toxinen oder Mikroorganismen. Die Therapieergebnisse können dadurch wesentlich verbessert werden.

Wird die Behandlung mit Schwingungsmustern mittels einer Licht- und Farbtherapie durchgeführt, sind für den Therapieerfolg folgende Fragen wesentlich:

- Wie schnell sollen die farbigen Leuchtmittel blinken?
- Wie lang sollen die Ein- und Auszeiten der Leuchtmittel sein?
- In welchen Rhythmen sollen die unterschiedlichen Farben nacheinander ablaufen?
- Mit welcher Intensität sollen die Farben leuchten?

Wie wichtig die Antworten auf diese und ähnliche Fragen sind, wurde bereits vor fast 80 Jahren durch Alexander Gurwitsch, dem frühen Pionier der Biophotonen, nachgewiesen. In einem Experiment hat er zwei Kulturen von Hefezellen durch eine undurchsichtige Scheibe getrennt, die mit Löchern versehen war. Wenn die Kulturen einander kontinuierlich sehen konnten (ein Loch befand sich genau zwischen den Kulturen), war eine beidseitige Erhöhung des prozentualen Anteils der sich teilenden Hefezellen nach sechs bis acht Minuten festzustellen. Wenn die Scheibe mit 50 Hertz rotierte, war bereits nach 30 Sekunden diese Erhöhung vorhanden. Die Zeit konnte auf 12 Sekunden reduziert werden, wenn die Scheibe mit 100 bis 800 Hertz rotierte. Dies zeigte nicht nur, dass eine gegenseitige Beeinflussung durch Biophotonen stattfindet, sondern auch, dass diese viel effektiver wird, wenn der Strom der Biophotonen mit einer bestimmten Frequenz unterbrochen wird (Pulsung). Korrekt gepulstes Licht bzw. gepulste Farben haben eine größere Wirkung auf Organismen, als eine gleichmäßige Bestrahlung.

Es werden bereits kommerziell verfügbare Geräte angeboten, zum Beispiel mit dem Namen **OptiSanPro**. Als spezialisierte Anwendung wird das Gerät von Therapeuten erfolgreich zur Basisbehandlung und weiterführenden Behandlung eingesetzt. Bei allen Detoxmaßnahmen kann die Ausleitung von Metallen, petrochemischen Substanzen, Phtalaten, Herbiziden, Pestiziden, Insektiziden, Umweltgiften und Impfbelastungen unterstützt und beschleunigt werden. Auch bei Allergenen zeigen sich Erfolge. Faszien, die lange unterschätzten Umhüllungen aus Bindegewebe, können oft die Quelle und Ursache unerklärlicher Zustände und Störungen bei chronischen Krankheiten sein. Auch hier berichten Therapeuten von guten und schnellen Erfolgen.

Für die Behandlung wurden unterschiedliche Programme für bestimmte Zustände oder Bereiche konzipiert, bspw. Rumpf unten, Rumpf oben, Herz, Wundbehandlung, Haut, Gewebe, Hormon/Membran und Blockade. Besondere Aufmerksamkeit und Sorgfalt wurden dabei der Berechnung der fraktalen Komponenten wie der Lichtintensität und Rhythmen (Impulsdauer und Pausenzeiten) der einzelnen farblichen Leuchtmittel und deren Kombinationen gewidmet. Die entsprechenden Programme sind gespeichert und können einfach gestartet werden. Ein Mikroprozessor steuert die eingespeicherten Programme, die zum Teil aus Tausenden verschiedenen Veränderungen an den Leuchtmitteln bestehen. Die dafür eingesetzten LEDs reagieren zeitgerecht auf die unterschiedlichsten Kombinationen von Lichtimpulsen, die durch entsprechende Pausenzeiten unterbrochen sind.

Das Gerät besitzt acht verschiedene LEDs und eine weitere, die ausschließlich im infraroten Bereich wirkt. Neben der Ansteuerung der Einzelfarben, z.B. Türkis oder Gelb, steht durch unterschiedliche Kombinationen der Intensitäten von Rot, Grün und Blau eine beliebige Anzahl von verschiedenen Farbnuancen zur Verfügung,

die das Auge wahrnehmen kann. Diese große Farbenvielfalt ist gleichbedeutend mit einer großen Photonenvielfalt, die mit unterschiedlichen Rhythmen effektiv zur Anwendung kommt. Alle LEDs besitzen zudem die Möglichkeit, Farben mit erhöhter Intensität abzustrahlen, um bestimmte physiologische Wirkungen gezielt zu verstärken. Die behandelten Personen empfinden die Art dieser fraktalen Impulse als sehr angenehm und wohltuend. (Dr. Siegfried Kiontke, Physiker und Autor)

3.5 Die Lunge

Krankheitsrisiken von Rauchern sind anhand des Verzweigungsgrades der Bronchien abschätzbar. Das Ausmaß einer Lungenschädigung durch Rauchen lässt sich zeitig erkennen, wenn man den Verlust an Verzweigungskomplexität im gesamten Bronchialsystem quantifiziert. Dies wird ermöglicht durch die fraktale Analyse von CT-Aufnahmen, die die Tatsache nutzt, dass die Lunge wie ein Fraktal nach dem Prinzip der Selbstähnlichkeit aus sich immer weiter verkleinernden Kopien ihrer bronchialen Verzweigungen aufgebaut ist [7].



Abbildung 3.4: Ein gesunder Lungenflügel und sein Aussehen nach zehnjährigem Rauchen

Kennzeichen einer beginnenden chronischen Raucherbronchitis (COPD) sind verdickte Wände und dadurch verengte Durchmesser der Atemwege sowie ein zunehmender Verlust von bronchialen Verzweigungen in der Peripherie der Lunge, bei den kleinsten Atemwegen. Bei einigen Rauchern, die bei einer Untersuchung mit Computer - Tomographen (CT) noch keine nachweisbare Atemwegsverengung aufweisen, lässt sich trotzdem ein erhöhtes Krankheits- und Sterberisiko feststellen, mit Hilfe der fraktalen Analyse von CT-Aufnahmen zur Bestimmung der fraktalen Dimension der Atemwege (airway fractal dimension = AFD). „Ein großer Vorteil der fraktalen Analyse von CT-Aufnahmen ist, dass sich das Ausmaß der Lungenschädigung durch das Rauchen früher erkennen lässt, indem man den Verlust an Verzweigungskomplexität im gesamten Bronchialsystem quantifiziert. Dadurch kann man

das Risiko des Patienten auf eine chronisch obstruktive Lungenerkrankung (COPD) besser abschätzen“, erläutert Dr. med. Thomas Voshaar, Vorstandsvorsitzender des Verbands Pneumologischer Kliniken (VPK) und Chefarzt des Lungenzentrums am Krankenhaus Bethanien in Moers unter Berufung auf aktuelle Studienergebnisse aus den USA. Forscher und Kliniker suchen schon seit einiger Zeit nach einer Methode zur früheren Erfassung von Veränderungen im Bronchialsystem, die bei einer COPD immer als erstes in der Peripherie auftreten und dann meistens lange Zeit unbemerkt bleiben. Erst wenn Symptome wie Kurzatmigkeit bei körperlicher Belastung auftreten, lassen sich auch bei einer Lungenfunktionsprüfung Einschränkungen nachweisen.

Grundlage der fraktalen Analyse ist der Aufbau der Lunge, die in ihrer Struktur einem Fraktal entspricht. Es beginnt mit der Luftröhre, die sich in zwei Hauptäste aufteilt, die jeweils einen der beiden Lungenflügel mit Sauerstoff versorgen. Bewegt man sich längs eines Astes, sieht man, dass die an den Enden der Hauptäste sitzenden Lappenbronchien sich weiter verzweigen zu den Segmentbronchien und dann in immer kleinere Äste. So gelangt man nach etwa 20 – 25 Ebenen zu den kleinsten Verzweigungen der Bronchien, den so genannten Bronchiolen, die einen Innendurchmesser von weniger als 1 mm haben und sich dann noch einmal in mikroskopisch feinste Ästchen (Bronchioli respiratorii und Alveolargänge) teilen. Diese führen schließlich in das eigentliche, atmende Lungengewebe mit insgesamt rund 300 Millionen Lungenbläschen (Alveolen). Die Lungenbläschen, die als Miniatur-Luftballone einen Durchmesser von 0,1-0,2 mm haben, geben der Lunge ihr schwammartiges Aussehen und sind traubenförmig dicht gepackt den feinsten röhrenartigen Ästchen (Alveolargänge und Bronchioli respiratorii) angelagert. Ihre hauchdünnen Wände sind von einem Netz kleinster Blutgefäße (Kapillaren) durchzogen, die einen schnellen Austausch der Atemgase ermöglichen. Jedes Lungenbläschen ist von etwa 1000 Kapillaren umgeben.

In der aktuellen Studie mit über 8000 Rauchern [8] haben Forscher der University of Alabama in Birmingham einen deutlichen Zusammenhang zwischen der fraktalen Dimension der Atemwege und den folgenden Parametern der Lungengesundheit festgestellt:

- Je niedriger die Verzweigungskomplexität umso höher der Grad der Atemwegsverengung, die Erkrankungshäufigkeit der Atemwege und der Verlust der Lungenfunktion hinsichtlich der Einsekundenkapazität FEV1 (die das maximale Volumen Luft angibt, das der Patient innerhalb einer Sekunde ausatmen kann).
- Weniger deutlich, aber tendenziell ebenfalls ausgeprägt hängt die AFD auch mit der Lebensqualität der Patienten und ihrer körperlichen Leistungsfähig-

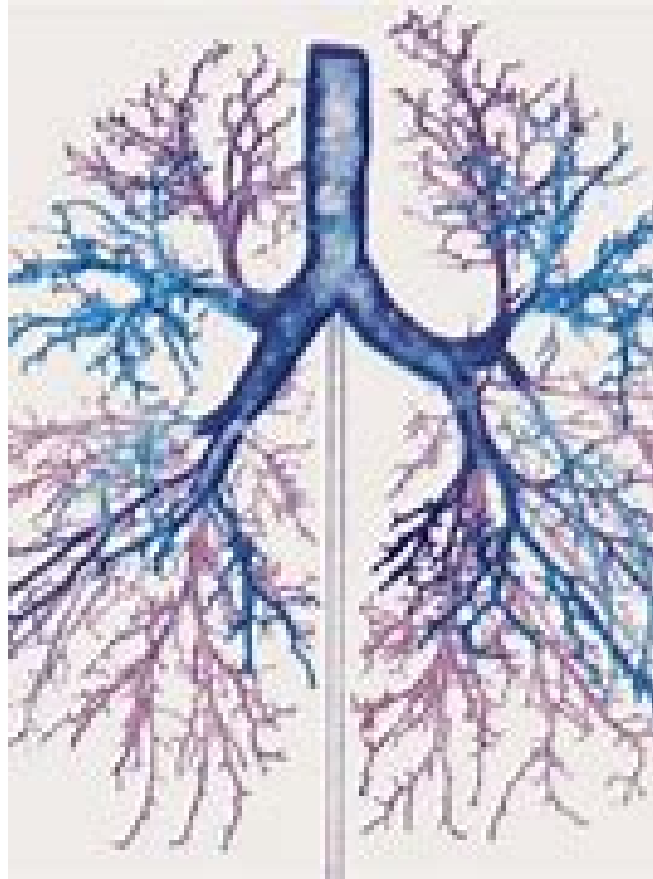


Abbildung 3.5: Die fraktale Struktur der Blutgefäße in der Lunge

keit (6-Minuten-Gehstrecke) zusammen, sowie - in inverser Beziehung - mit der Häufigkeit von Verschlechterungen (Exazerbationen).

„Künftig wäre es sehr aufschlussreich, bei CT-Untersuchungen von Rauchern auch die fraktale Dimension der Atemwege zu bestimmen, da ein Verlust an Komplexität der bronchialen Verzweigungen einen guten Hinweis darauf liefern kann, wie stark geschädigt und damit krankheitsanfällig die Lunge des Patienten bereits tatsächlich ist“.

3.6 Krebsdiagnose

Forscher um Joachim Spatz, Direktor am Max-Planck-Institut für Intelligente Systeme in Stuttgart und Professor an der Universität Heidelberg, haben festgestellt, dass sich Tumorzellen und gesunde Zellen anhand ihrer fraktalen Geometrie unterscheiden. Die Wissenschaftler vergrößerten die Ränder von Zellen der Bauchspeicheldrüse und analysierten deren Unregelmäßigkeiten. Indem sie diese mathema-

tisch erfassten, bestimmten sie die fraktale Dimension des Zellrandes, die ein Maß für die statistische Verteilung der Unregelmäßigkeiten ist. Krebszellen weisen einen höheren Fraktalisierungsgrad auf als gesunde Zelle, da sich beim unkontrollierten Tumorstadium sehr unregelmäßige Ausbuchtungen verschiedener Größe auf der Zelloberfläche bilden. Die Forscher erkannten an der fraktalen Dimension aber nicht nur, ob eine Tumorzelle vorlag, sondern bestimmten mit 97-prozentiger Sicherheit auch, um welchen von zwei unterschiedlich bösartigen Bauchspeicheldrüsentumoren es sich handelt. „Auf diese Weise lassen sich Krebszellen sehr viel genauer und schneller unterscheiden als mit der bisher gängigen Methode“.

Krebszellen zu identifizieren, ist bisher unsicher und zeitaufwendig. Bisher werden Krebszellen und ihr Entstehungsort im Körper identifiziert, indem eine Zellprobe mit bestimmten Antikörpern und Biomarkern eingefärbt wird. Die Färbemethode hat jedoch Nachteile: Sie erfordert zahlreiche Einzelschritte mit kostspieligen Antikörpern und ist daher zeitaufwendig und teuer. Außerdem können die derzeit gebräuchlichen Farbstoffe feinste Unterschiede von Zellen nicht immer sichtbar machen. Daher lässt sich ein Krebs mit dieser Methode nur in 85 Prozent der Proben korrekt diagnostizieren. Mithilfe der fraktalen Geometrie erkennt das Team von Joachim Spatz Krebszellen nicht nur zuverlässiger, sondern auch deutlich schneller. Denn die Zellen lassen sich unter einem Mikroskop untersuchen, ohne dafür besonders präpariert werden zu müssen. Um die Details der Zellränder erfassen zu können, verwendet das Team von Joachim Spatz ein Reflexionskontrastmikroskop. Statt die Probe wie ein herkömmliches Lichtmikroskop von unten zu durchleuchten, misst das Mikroskop der Stuttgarter Forscher die Reflexion des Lichtstrahls an der Zelloberfläche. Diese ist unterschiedlich, je nachdem, ob das Licht direkt auf eine Zelle oder zuerst auf wässriges Zellkulturmedium und dann auf eine Zelle trifft. Anhand des reflektierten Lichtes lassen sich selbst winzige Strukturen am Zellrand untersuchen.

„Die fraktale Geometrie der Zelloberflächen zu analysieren, birgt ein sehr großes Potenzial für die klinische Diagnostik“. Nun erforschen die Wissenschaftler, wie sich ihre Methode in der Praxis anwenden lässt. Dazu untersuchen sie unterschiedliche bösartige Zell-Linien und primäre Zellen, also solche Zellen, die aus menschlichen Organen gewonnen werden und die im Gegensatz zu Tumorzell-Linien nur über eine bestimmte Zeitspanne kultiviert werden können. „Der nächste Schritt für uns werden konkrete Kooperationen mit Kliniken sein, um die Methode direkt an relevanten Gewebeproben zu testen“. [17]



Abbildung 3.6: Ein weit fortgeschrittener Tumor

3.7 Die Haut

Elektronische Bauteile im Gehirn können heute schon Menschen mit schweren Erkrankungen an Nerven und Gehirn helfen – etwa Parkinson-Patienten: Hier lindern die Stromimpulse deren typische Lähmungen und Zuckungen. Auch bei Menschen mit Epilepsien oder Depressionen setzen Hirnchirurgen solche Elektroden ins Gehirn ein.

US-Forscher haben erstmals ein dreidimensionales Modell des menschlichen Erbguts gebaut. Damit fanden sie heraus, wie die Zelle es schafft, ihre zwei Meter lange Erbsubstanz im Zellkern zu verstauen. Zum einen organisiert die Zelle zwei Abteilungen, eine Art Ruheraum und eine Fabrik, in die sie jeweils die aktiven und die inaktiven Gene sortiert. Zum anderen formt sich die Erbsubstanz DNA zu einer perlenkettenartigen Struktur, die sich wiederum wie ein Wollknäuel verdrillt. Die

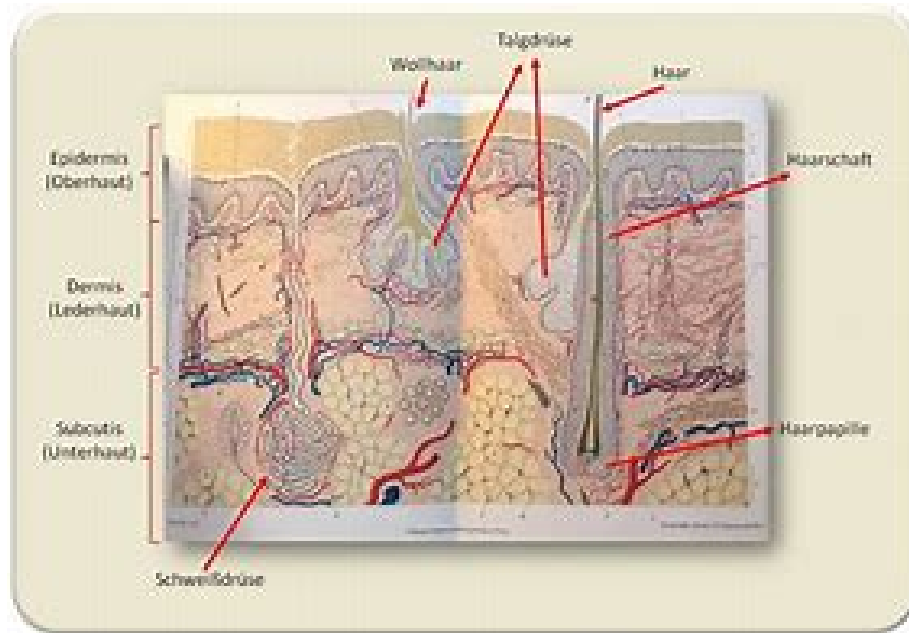


Abbildung 3.7: Auch die Haut besitzt eine Fraktale Struktur

Informationsdichte dieser Kugelstruktur ist viel höher als die eines Computerchips. Die Organisation verhindert gleichzeitig, dass sich das Erbgut verheddert und die Zelle ihr eigenes Genom nicht mehr lesen kann.

Bisher war es nicht klar, wie sich die Erbgut-Stränge mit ihren über drei Milliarden DNA-Basenpaaren so organisieren, dass sie in einen menschlichen Zellkern mit einem Durchmesser von einem Hundertstel Millimeter passen. Um die räumliche Struktur des Genoms sichtbar zu machen, verklebten die Wissenschaftler jetzt eng im Zellkern beieinander liegende DNA-Stränge und identifizierten sie anschließend. Das Genom wird so in Millionen Stücke zerlegt und dann wieder zu einer räumlichen Karte zusammengestellt, die nachbarschaftliche Beziehungen zeigt. Das 3-D-Modell zeigt, dass die Zellen ihr Erbgut auf zwei Abteilungen verteilen, berichtet Job Dekker, Systembiologe an der Harvard Medical School. In der ersten befinden sich die aktiven Gene, leicht erreichbar für Proteine und andere Steuerelemente. Im zweiten Séparée steckt die passive DNA, die dafür sehr eng zusammengepackt wird. Die einzelnen DNA-Moleküle wechseln je nach Anforderung zwischen der Fabrik und dem Ruheraum, wobei sich die jeweils aktiven Regionen annähern.

Für die Speicherung von Informationen hat die Natur eine zusätzliche außerordentlich dichte knotenfreie Struktur geschaffen: Die DNA ballt sich in der sogenannten Fraktal-Kugel zusammen, ohne dass sie bei der Entfaltung für die Zellteilung behindert wird. Die Fraktal-Kugel-Architektur war als theoretische Möglichkeit schon vor über 20 Jahren diskutiert worden. Erst die Entwicklung des neuen Verfahrens, das die nachbarschaftlichen Beziehungen einzelner Gene offenlegt, hat nun Klarheit

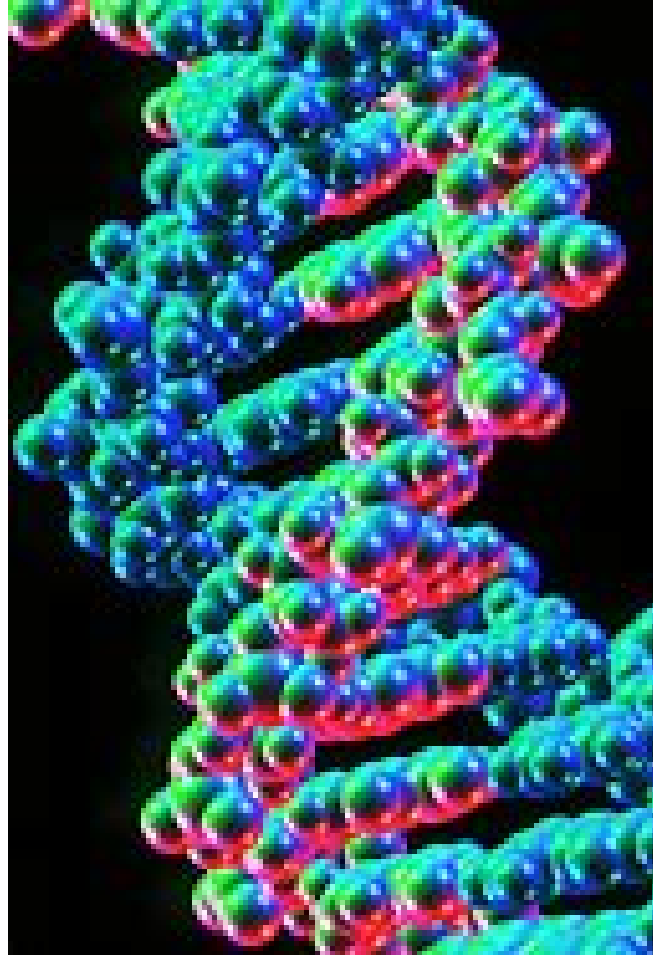


Abbildung 3.8: Die fraktale Struktur der DNA

geschaffen.[18]

Die DNA ist der Träger des genetischen Codes, und aus ihr werden alle Gene gebildet, die für den Aufbau unseres Körpers verantwortlich sind. Lange Zeit glaubte man, dies geschehe ausschließlich auf biochemischem Wege. Die DNA bildet einen riesigen Doppelstrang aus Basen, in denen die Erbinformation codiert ist, mit deren Hilfe dann im Innern der Zelle Eiweißkörper hergestellt werden können.

Russische Wissenschaftler haben aber herausgefunden, dass die DNA noch viel mehr kann. Fast 90 Prozent dieses Moleküls werden nämlich überhaupt nicht zur Eiweißsynthese benötigt, sondern dienen zur Kommunikation und als Informationsspeicher. Durch die charakteristische Form der Doppelhelix stellt die DNA eine geradezu ideale elektromagnetische Antenne dar. Einesteils ist sie langgestreckt und damit eine Stabantenne, die sehr gut elektrische Impulse aufnehmen kann. Andererseits ist sie, von oben gesehen, ringförmig und damit eine sehr gute magnetische Antenne. Auf diese Weise kann unsere DNA elektromagnetische Strahlung (Licht) aus der Umwelt aufnehmen. Die aufgenommenen Energie wird ganz einfach

in der DNA gespeichert, indem das Molekül in Schwingung versetzt wird, mit einer Eigenfrequenz von 150 Megahertz.

Nach den Forschungsergebnissen von Pjotr Garjajev und seinem Team ist die DNA nicht nur Sender und Empfänger elektromagnetischer Energie, sondern nimmt auch die in der Strahlung enthaltene Information auf und interpretiert sie weiter. Die DNA ist also ein höchst komplexer interaktiver Biochip auf Lichtbasis mit 3 Gigabits Speicherfähigkeit, der noch dazu in der Lage ist, die menschliche Sprache zu verstehen.

3.8 Fraktale in der Anatomie

Viele Organe und Strukturen des menschlichen Körpers, zum Beispiel die Netzhautgefäße, der Atemtrakt, die Nierenarterien, die das Herz versorgenden Blutgefäße, die Leber und die Dendriten von Nervenzellen können als fraktale Objekte betrachtet werden. Die fraktale Geometrie kommt in der Kardiologie zur Berechnung der Herzfrequenz, in der Neurologie zur Analyse sich ändernder Muster von Elektroenzephalogrammen (EEG) und in der Radiologie zur Analyse von Knochenheilungen, Brustläsionen und Tomographien zur Anwendung. Die fraktale Geometrie konnte auch in der Histopathologie und Zytologie angewandt werden. Dort wird sie zur Berechnung mehrerer Neoplasien wie Gallenblasen-, Lungen-, Gebärmutter-, Brust-, Mundhöhlen- oder Kehlkopfkrebs herangezogen. Seit kurzer Zeit werden fraktale Methoden auch angewandt, um die Komplexität der Oberflächentopografie zu beschreiben. Die fraktale Geometrie hat sich auf mehreren Ebenen als ein geeignetes Mittel zur Messung der Unregelmäßigkeit und Komplexität von Geweben erwiesen: von der Organisation von Zellen auf der Organebene bis hin zur Oberfläche von Zellen, der Verteilung von Protein-Aggregaten in der Membran, der Organisation des Zytoplasmas, des Zellkerns und der Struktur einzelner Proteine.

Die Anatomie-Ausbildung wird gegenwärtig an vielen Stellen auf mit 3D - Druckern erzeugte Modelle umgestellt. Der Bedarf an realen toten Körpern geht zurück (Abb. 3.9).

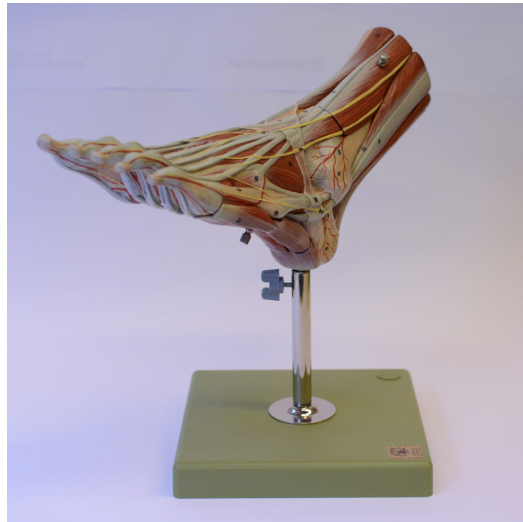


Abbildung 3.9: Modell eines Fußes (für die Anatomie-Ausbildung)

3.9 Epidemien

Eine **Epidemie** ist ein zeitlich und örtlich begrenztes vermehrtes Auftreten von Krankheitsfällen einheitlicher Ursache innerhalb einer menschlichen Population und entspricht damit einem großen Ausbruch einer Krankheit. Der Begriff ist nicht auf Infektionskrankheiten beschränkt. Die Ausbreitungsmuster einer Epidemie haben ebenfalls ein fraktales Aussehen. Sie ähneln stark der Ausbreitung von Waldbränden. Sie beginnen punktförmig an einer oder an mehreren Stellen und breiten sich in kleinen Schritten nach allen Seiten aus.

In der **Epidemiologie** wird von einer Epidemie gesprochen, wenn die Anzahl an neuen Erkrankungsfällen (**Inzidenzen**) über einen gewissen Zeitraum in einer bestimmten Region zunimmt. Nach der Geschwindigkeit der Zunahme der Erkrankungsfälle werden **Explosiv-** und **Tardiv-Epidemien** unterschieden. Bei einer auf Länder und Kontinente übergreifenden Ausbreitung wird von einer **Pandemie** gesprochen. Ein Rückgang der Erkrankungshäufigkeit wird als **Regression** bezeichnet. Als eine **Endemie** wird demgegenüber das andauernd gehäufte Auftreten einer Krankheit in einer umschriebenen Population bezeichnet; hierbei bleibt die Inzidenz annähernd gleich, ist aber gegenüber nichtendemischen Gebieten erhöht. Im Unterschied zu einer Endemie – bei der eine Krankheit innerhalb einer Population fortwährend mit etwa gleicher Fallzahl auftritt ($\text{Reproduktionsrate} = 1$) – verbreitet sich eine Epidemie mit einer Reproduktionsrate größer als 1. Dies bedeutet bei einer Infektionskrankheit, dass die Zahl an Infizierten zunimmt und die Zahl an Neuinfektionen ansteigt. Für die Ausbreitung bedeutsam ist die Rate, mit der durch Kontakt mit Infizierten neue Infizierte auftreten; sie entspricht zu Beginn der **Basis-Reproduktionszahl** R_0 . Anfangs erhöht sich die Zahl an neuen Infektions-

fällen pro Zeitintervall im Vergleich zum vorigen um einen ungefähr gleichen Anteil und wächst exponentiell. Der Anstieg neuer Infektionsfälle in absoluten Zahlen fällt daher zunächst eher gering aus und wächst mit fortschreitendem Geschehen stärker an.

Diese dynamische Entwicklung kann gedämpft werden, wenn die Zahl an infektiösen Kontakten eingeschränkt wird – beispielsweise durch Quarantäne oder ein verändertes Sozialverhalten mit Distanzierung und geeigneten Hygiene-Maßnahmen – und die Zahl der pro Fall übertragenen Zweitinfektionen absinkt. Bei einer Netto-Reproduktionszahl $R_t \leq 1$ nimmt die Zahl an neu auftretenden Krankheitsfällen nicht mehr zu. Kann eine Epidemie während des Verlaufs nicht eingedämmt werden, kommt es hierzu erst, nachdem die Krankheit sich in der Bevölkerung soweit ausgebreitet hat, dass der Anteil anfälliger, noch nicht infizierter Individuen stark reduziert ist. Dadurch sinkt die Zahl der Neuinfektionen nach einiger Zeit immer weiter ab, bis die Krankheit einen endemischen Status erreicht oder in der Population ausstirbt (Populationsdynamik).

Das vermehrte Auftreten neuer Krankheitsfälle möglichst früh zu erfassen, ist für den Schutz der Bevölkerung wesentlich. Viele Betroffene suchen im Internet nach Information zu Krankheiten. Die Auswertung der Daten von Suchmaschinen kann daher Hinweise geben, um Epidemien frühzeitig zu erkennen. Auch die Auswertung von persönlichen Nachrichtendiensten im Internet kann für diese Bewertung herangezogen werden. Allerdings ist eine gehäufte Suche nach einer Krankheit oder deren Erwähnung im Internet nicht unbedingt immer Folge einer erhöhten Prävalenz oder Inzidenz dieser Krankheit. Daher können überhöhte Prognosen gestellt werden, wenn nicht andere zusätzliche Datenquellen in die Bewertung einfließen.

Zu den epidemisch auftretenden Krankheiten gehören verschiedene Tropenkrankheiten wie das Dengue-Fieber, aber auch Cholera, Grippe, Typhus und Poliomyelitis (Kinderlähmung). Früher traten Milzbrand-Epidemien öfter im Abstromgebiet von Gerbereien auf. Die wohl verheerendsten Epidemien der Menschheitsgeschichte wurden von der Pest ausgelöst. Als **Schwarzer Tod** wird eine der verheerendsten Pandemien der Weltgeschichte bezeichnet, die in Europa zwischen 1346 und 1353 geschätzt 25 Millionen Todesopfer – ein Drittel der damaligen Bevölkerung – forderte. Als Ursache gilt die durch das Bakterium *Yersinia pestis* hervorgerufene Pest. Das Wort „Pest“ leitet sich vom lateinischen Wort *pestis* für Seuche ab und wird daher auch ohne direkten Bezug auf die Krankheit verwendet. Die Pandemie trat nach heutigem Wissensstand zuerst in Zentralasien auf und gelangte über die Handelsrouten (unter anderem über die Seidenstraße) nach Europa. Aus dem östlichen Mittelmeerraum verbreitete sich die Krankheit wahrscheinlich über Rattenflöhe in das restliche Europa, jedoch blieben einige Landstriche relativ verschont. Für das Gebiet des heutigen Deutschland wird geschätzt, dass jeder zehnte Ein-

wohner infolge des Schwarzen Todes sein Leben verlor. Bremen, Hamburg, Köln und Nürnberg zählten dabei zu den Städten, in denen ein sehr hoher Bevölkerungsanteil starb. Sehr viel geringer war dagegen die Anzahl der Todesopfer im östlichen Gebiet des heutigen Deutschland.



Abbildung 3.10: Eine an der Pest erkrankte Frau

Der Zeitzeuge Boccaccio hat in seinem Werk *Decamerone* eindrucksvoll geschildert, wie nach dem Ausbruch der Pandemie viele Einwohner von Florenz ihren sozialen Verpflichtungen nicht mehr nachkamen:

„Wir wollen darüber schweigen, dass ein Bürger den anderen mied, dass fast kein Nachbar für den anderen sorgte und sich selbst Verwandte gar nicht oder nur selten und dann nur von weitem sahen. Die fürchterliche Heimsuchung hatte eine solche Verwirrung in den Herzen der Männer und Frauen gestiftet, dass ein Bruder den anderen, der Onkel den Neffen, die Schwester den Bruder und oft die Frau den Ehemann verließ; ja, was noch merkwürdiger und schier unglaublich scheint: Vater und Mutter scheuten sich, nach ihren Kindern zu sehen und sie zu pflegen – als ob sie nicht die ihren wären. Viele starben, die, wenn man sich um sie gekümmert hätte, wohl wieder genesen wären. Aber wegen des Fehlens an ordentlicher, für den Kranken nötiger Pflege und wegen der Macht der Pest war die Zahl derer, die Tag und Nacht starben, so groß, dass es Schaudern erregte, davon zu hören, geschweige denn es mitzuerleben.“

Viele der Menschen, welche den Schwarzen Tod als Gottesstrafe ansahen, fanden zu dieser Zeit den Trost in der Religion. Religiöse Bewegungen entstanden spontan im Gefolge oder in Erwartung der Seuche – viele davon forderten das Monopol der Kirche auf geistliche Lenkung heraus. Bittgottesdienste und Prozessionen kennzeichneten den Alltag. Flagellanten zogen in „Geißlerzügen“ durch die Städte. Der „Pestheilige“ St. Rochus wurde intensiv verehrt, Pilgerfahrten nahmen zu. An vielen

Orten zeugen Kirchen und andere Monumente wie so genannte Pestsäulen von der Angst der Menschen und ihrem Wunsch nach Erlösung von der Seuche. Verschiedene Menschen versuchten jede Minute ihres Lebens noch auszukosten, und mit Tanz und Musik versuchte man, dem Schwarzen Tod zu entgehen. Der italienische Chronist Matteo Villani schrieb:

„Die Menschen, in der Erkenntnis, dass sie wenige und durch Erbschaften und Weitergabe irdischer Dinge reich geworden waren, und der Vergangenheit vergessend, als wäre sie nie gewesen, trieben es zügelloser und erbärmlicher als jemals zuvor. Sie ergaben sich dem Müßiggang, und ihre Zerrüttung führte sie in die Sünde der Völlerei, in Gelage, in Wirtshäuser, zu köstlichen Speisen und zum Glücksspiel. Bedenkenlos warfen sie sich der Lust in die Arme.“ Eine funktionierende Wirtschaft konnte unter dem Eindruck einer Pandemie nicht mehr aufrechterhalten werden. Arbeitskräfte starben, flohen und nahmen ihre Aufgaben nicht mehr wahr. Vielen schien es sinnlos, die Felder zu bestellen, wenn der Tod sie doch bald ereilen würde.

Die kirchliche und weltliche Macht verlor angesichts der Hilflosigkeit, mit der sie der Pandemie begegnete, rapide an Autorität. Der Dichter Boccaccio vermerkte in seinem Decamerone: „In solchem Jammer und in solcher Betrübnis der Stadt war auch das ehrwürdige Ansehen der göttlichen und menschlichen Gesetze fast gesunken und zerstört; denn ihre Diener und Vollstrecker waren gleich den übrigen Einwohnern alle krank oder tot oder hatten so wenig Gehilfen behalten, dass sie keine Amtshandlungen mehr vornehmen konnten. Darum konnte sich jeder erlauben, was er immer wollte.“

Unter dem Autoritätsverlust der weltlichen und kirchlichen Macht litten diejenigen Menschen am meisten, die zu den kulturellen Randgruppen der mittelalterlichen Gesellschaften zählten. So kam es im Zuge der Pandemie zu schweren Judenpogromen, die von den geistlichen und weltlichen Herrschern nicht mehr unterbunden werden konnten und die zur Folge hatten, dass nach 1353 nur noch wenige Juden in Deutschland und den Niederlanden lebten. Die Pogrome brachen aus, da das aufgebrachte Volk in den Juden die Schuldigen für die Katastrophe auszumachen glaubte. Ein Breslauer Arzt namens Dr. Heinrich Rybbinus stellte angesichts der Pest-Epidemie fest, dass die Menschen in ihrer Ratlosigkeit darauf verfallen seien, Juden zu verbrennen und sich gegenseitig zu bekriegen. Erste Übergriffe gegen Juden begannen in Toulon am Palmsonntag 1348: Kurz nachdem es dort zu ersten Pesttoten gekommen war, griffen Teile der Stadtbevölkerung das jüdische Viertel an und töteten 40 Personen. Wenige Tage später gab es Übergriffe auch in Avignon, Grasse sowie anderen Städte in der Provence und dann in Katalonien.

Das Gerücht, dass bestimmte Personenkreise Gift in Brunnen und Quellen träufelten, zirkulierte sehr häufig in Notzeiten und wurde beispielsweise 1321 nach dem

Hirtenkreuzzug von 1320 Leprakranken vorgeworfen. Sehr schnell nach den ersten Pesttoten wurde dies auch den jüdischen Mitbürgern vorgeworfen: In Savoyen hatten jüdische Angeklagte sich unter der Folter solcher Vergehen für schuldig bekannt. Ihr Geständnis fand in ganz Europa rasch Verbreitung und war die Basis für eine Welle von Übergriffen in der Schweiz und in Deutschland – vor allem im Elsass und entlang des Rheins. Am 9. Januar 1349 wurde in Basel ein Teil der jüdischen Einwohnerschaft ermordet – die Basler Stadträte hatten zuvor zwar die schlimmsten Hetzer aus der Stadt verbannt, mussten unter dem Drängen der Stadtbevölkerung diesen Bann jedoch wieder aufheben und stattdessen die Juden vertreiben. Ein Teil der Vertriebenen wurde festgesetzt und in einem eigens für sie gebauten Haus auf einer Rheininsel verbrannt. In Straßburg versuchte die Stadtregierung ebenfalls, die ansässigen Juden zu schützen, wurde jedoch mit den Stimmen der Zünfte ihres Amtes enthoben. Die neue Straßburger Stadtregierung duldete das anschließende Massaker, dem im Februar 1349 – also zu einem Zeitpunkt, zu dem der Schwarze Tod die Stadt noch nicht erreicht hatte – 900 von 1.884 in Straßburg lebende Juden zum Opfer fielen. Im März 1349 verbrannten sich vierhundert Mitglieder der jüdischen Gemeinde von Worms in ihren Häusern, um der Zwangstaufe zu entgehen; im Juli 1349 beging auch die jüdische Gemeinde von Frankfurt auf diese Weise Selbstmord. In Mainz griffen Juden zur Selbstverteidigung und töteten 200 angreifende Stadtbürger. Selbst die in Mainz lebende jüdische Gemeinde – damals die größte in Europa – beging letztlich Selbstmord durch Anzünden der eigenen Häuser. Die Pogrome setzten sich bis Ende des Jahres 1349 fort. Die letzten fanden in Antwerpen und Brüssel statt. Für Städte wie Freiburg im Breisgau, Köln, Augsburg, Nürnberg, Königsberg und Regensburg wird angenommen, dass noch vor dem lokalen Ausbruch der Seuche Flagellanten Teile der Bevölkerung aufhetzten, die jüdische Bevölkerung als Brunnenvergifter zu ermorden. Die neuere Forschung geht jedoch davon aus, dass das Abwälzen der Schuld auf die Geißler zumeist lediglich als bequemer Rechtfertigungsversuch der Geschichtsschreibung des 14. Jahrhunderts für die Morde zu werten ist. Neben der Suche nach einem Sündenbock und einer seit dem 12. Jahrhundert angestiegenen Intoleranz der Kirche gegenüber Andersgläubigen war auch Habgier ein wesentliches Motiv für den Mord an jüdischen Mitbürgern. Die Bedeutung der Juden als Geldverleiher war zwar nicht mehr so groß wie noch im 12. und 13. Jahrhundert, doch offenbar sah ein großer Teil der Bevölkerung im Mord an den Juden auch die Möglichkeit, sich ihrer Gläubiger zu entledigen. So war der Augsburger Bürgermeister Heinrich Portner bei jüdischen Geldleihern hoch verschuldet und ließ den Mord an den Juden bereitwillig geschehen.

Es fehlte nicht an Personen, die auf das Unrecht dieser Morde aufmerksam machten. Bereits am 4. Juli 1348 wandte sich der in Avignon lebende Papst Klemens VI. in einer Bulle gegen die Verfolgung von Juden. Die päpstliche Bulle wirkte nur in Avignon und trug ansonsten verhältnismäßig wenig zum Schutz der Juden bei. Daher folgte am 26. September 1348 eine zweite päpstliche Bulle mit dem Titel

Quamvis perfidiam. Die Anschuldigung, die Juden würden durch das Vergiften von Brunnen die Seuche verbreiten, bezeichnete er darin als unvorstellbar, da sie in Gegenden der Erde wüte, wo keine Juden lebten, und dort, wo sie lebten, sie selbst Opfer der Seuche würden. Er forderte die Geistlichkeit auf, die Juden unter ihren Schutz zu stellen. Klemens VI. – der selbst hebräische Manuskripte sammelte – untersagte außerdem, Juden ohne Gerichtsverfahren zu töten oder sie auszuplündern. Er drohte den Verfolgern die Strafe der Exkommunikation an. Die Geißlerbanden, die sich bei den Judenpogromen besonders hervorgetan hatten, erklärte er zu Häretikern. Ähnlich wirkungslos blieben auch die Maßnahmen, die Königin Johanna I. von Neapel ergriff und die im Mai 1348 die Steuerlast der in ihrem provenzalischen Herrschaftsgebiet lebenden Juden um die Hälfte reduzierte, um den Plünderungen Rechnung zu tragen. Im Juni desselben Jahres wurden ihre Beamten aus den provenzalischen Städten vertrieben, was die Schutzlosigkeit der Juden aufgrund des fortschreitenden Autoritätsverlusts der Herrschenden illustriert. Ebenso wie Papst Klemens waren Peter IV. von Aragon, Albrecht II. von Österreich und Kasimir III. von Polen entschiedene Beschützer ihrer jüdischen Einwohner. Wenn sie auch Gewalttaten nicht gänzlich unterbinden konnten, blieben solche Massaker wie in Brüssel und Basel aus. Kasimir III. bot darüber hinaus den Juden an, sich in seinem Herrschaftsgebiet anzusiedeln. Es setzte eine Emigration vor allem von deutschen Juden nach Polen ein, die bis ins 16. Jahrhundert anhielt. In der Ansiedlung jüdischer Bürger sah Kasimir III. die Möglichkeit, die Größe der durch die Mongolenraubzüge dezimierten Bevölkerung zu erhöhen und damit sein Land wirtschaftlich weiterzuentwickeln.

Auf der anderen Seite fehlte es nicht an weltlichen Herrschern, die sich die sogenannten Pestpogrome zu Nutze machten. Der römisch-deutsche König Karl IV. machte sich mindestens der Mitwisserschaft schuldig: Um seine Schulden zu tilgen, verpfändete Karl das königliche Judenregal, unter anderem an Frankfurt am Main. Es wurde gar geregelt, was mit dem Besitz von Juden zu geschehen habe, falls „die Juden daselbst nächstens erschlagen“ würden (Frankfurter Urkunden vom 23., 25., 27. und 28. Juni 1349, bezogen auf Nürnberg, Rothenburg ob der Tauber und Frankfurt am Main). Obwohl er in seinem Herrschaftsbereich die Juden effektiv schützen konnte, wirft dieses Ereignis viele Fragen bezüglich Karls Charakter auf, besonders da Karl sonst immer bestrebt war, das Bild eines gerechten christlichen Herrschers zu vermitteln. Dabei verstieß nämlich die Duldung der Morde auch gegen das damalige Rechtsverständnis, da die Juden unter dem direkten Schutz des Königs standen und dafür auch Zahlungen leisteten. Noch weiter ging der Markgraf von Meißen, der zu Beginn des Jahres 1349 die Stadtbevölkerung von Meißen aufforderte, Juden zu attackieren, und ihnen zusicherte, dass solchen Übergriffen keine Sanktionen folgen würden.

Langfristig bewirkte und beschleunigte die Seuche einen tiefgreifenden Wandel in

der mittelalterlichen Gesellschaft Europas. Wie David Herlihy zeigt, konnten die Generationen nach 1348 nicht einfach die sozialen und kulturellen Muster des 13. Jahrhunderts beibehalten. Der massive Bevölkerungseinbruch bewirkte eine Umstrukturierung der Gesellschaft, die sich langfristig positiv bemerkbar gemacht habe. So bezeichnete Herlihy die Pandemie als „die Stunde der neuen Männer“: Die Entvölkerung ermöglichte einem größeren Prozentsatz der Bevölkerung den Zugang zu Bauernhöfen und lohnenden Arbeitsplätzen. Unrentabel gewordene Grenzböden wurden aufgegeben, was in manchen Regionen dazu führte, dass Dörfer verlassen oder nicht mehr wiederbesiedelt wurden (sogenannte Wüstungen), die im Hochmittelalter im Zuge des Landesausbaus abgeholzten Wälder breiteten sich wieder aus. Die Zünfte ließen nun auch Mitglieder zu, denen man zuvor die Aufnahme verweigert hatte. Während der Markt für landwirtschaftliche Pachten zusammenbrach, stiegen die Löhne in den Städten deutlich an. Damit konnte sich eine größere Anzahl von Menschen einen höheren Lebensstandard leisten als jemals zuvor; allerdings kam es teilweise auch zu Nahrungsmittelknappheit, weil viele Felder nicht mehr bewirtschaftet wurden, so z. B. in England, wo die Löhne für Landarbeiter stark anstiegen. Obwohl die Adeligen 1349 im Parlament das **Statute of Labourers** durchsetzten, das die Löhne für Feldarbeit begrenzte, wurden die Landarbeiter zusätzlich mit Naturalien bezahlt. Die Lohnkonflikte führten schließlich zum großen Bauernaufstand von 1381, in dessen Folge England als erstes Land Europas die Leibeigenschaft abschaffte. Freie Bauern wurden in der Folge durch Pächter ersetzt, die weniger arbeitsintensive Schafzucht verdrängte den Ackerbau.

Ein Tiefstand der Bevölkerung wurde in Europa um 1400 erreicht. Der deutliche Anstieg der Arbeitskosten sorgte dafür, dass manuelle Arbeit zunehmend mechanisiert wurde. Damit wurde das Spätmittelalter zu einer Zeit eindrucksvoller technischer Errungenschaften. David Herlihy nennt als Beispiel den Buchdruck: Solange die Löhne von Schreibern niedrig waren, war das handschriftliche Kopieren von Büchern eine zufriedenstellende Reproduktionsmethode. Mit dem Anstieg der Löhne setzten umfangreiche technische Experimente ein, die letztlich zur Erfindung des Buchdrucks mit beweglichen Lettern durch Johannes Gutenberg führten.

In Italien, insbesondere in der Toskana, schwächte die Seuche die Adelsherrschaft und führte zum Aufstieg der Handwerker und kleinen Händler, die oft Schuldner der Juden waren. Diese verloren oft den Schutz der Fürsten. Grund und Boden wurden relativ entwertet, da sie nun im Überfluss vorhanden waren, während Kapital an Bedeutung gewann und leichter akkumuliert werden konnte.

Die Kirche – von zahlreichen Seuchenopfern als Erbe eingesetzt – ging reicher, aber unpopulärer aus der Zeit des Schwarzen Todes hervor. Weder hatte sie eine zufriedenstellende Antwort auf die Frage gefunden, warum Gott der Menschheit eine solche Prüfung auferlegt hatte, noch hatte sie geistlichen Beistand geleistet, als das

Bedürfnis der Menschen danach am größten war. Die Bewegung der Flagellanten hatte die Autorität der Kirche auf die Probe gestellt. Auch nach dem Abklingen dieser Bewegung suchten viele Gott bei mystischen Sekten und in Reformbewegungen, was letztlich die katholische Glaubenseinheit auseinanderbrechen ließ.

Insbesondere der österreichische Kulturhistoriker Egon Friedell vertrat in seinem Werk *Kulturgeschichte der Neuzeit* die Auffassung, dass die Seuche der Jahre 1348/49 die Krise des mittelalterlichen Welt- und Menschenbildes verursacht und bis dahin bestehende Glaubensgewissheiten erschüttert habe. Er sieht eine direkte, kausale Verbindung zwischen der Katastrophe des Schwarzen Todes und der Renaissance.

Die erste große Pandemiewelle, die als **Schwarzer Tod** in die Geschichtsbücher einging, endete 1353. Sie flackerte in den Folgejahren immer wieder in einzelnen Regionen Europas auf, da sich die Seuche endemisierte: In lokalen und regionalen Epidemien trat sie die nächsten drei Jahrhunderte in nahezu regelmäßigen Abständen in verschiedenen Gegenden Europas auf, so im Jahr 1400 als zweitschlimmste Epidemie des ausgehenden Mittelalters bzw. der jungen Neuzeit. Große Pest von London 1665/1666, bei der in Südengland etwa 100.000 Menschen starben (davon alleine 70.000 in London) oder die Große Pest von 1708 bis 1714 in Nord- und Osteuropa mit etwa einer Million Toten. Ende des 19. Jahrhunderts begann in China die dritte Pest-Pandemie.



Abbildung 3.11: Eine mittelalterliche Gesichtsmaske zum Schutz gegen Infektionen

Die meisten Kunstwerke, die die Auswirkungen des Schwarzen Todes thematisieren, entstanden erst nach den Pandemie Jahren 1347 bis 1353. Eine Ausnahme stellt „Il Decamerone“ von Giovanni Boccaccio dar, das nach heutigem Wissensstand zwischen 1350 und 1353 geschrieben wurde. Ort der Rahmenhandlung ist ein Landhaus

in den Hügeln von Florenz, zwei Meilen vom damaligen Stadtkern von Florenz entfernt. In dieses Landhaus sind sieben Mädchen und drei junge Männer vor dem Schwarzen Tod geflüchtet, der im Frühjahr und Sommer des Jahres 1348 Florenz heimsuchte. Die Einleitung des Buches ist eine der detailliertesten mittelalterlichen Quellen über die Auswirkung des Schwarzen Todes in einer Stadt.

Der Schwarze Tod wurde auch in der Kunst des ausgehenden Mittelalters zu einem wichtigen Thema. Künstler wie der Lübecker Maler und Bildschnitzer Bernt Notke stellten das Geschehen in Form des Totentanzes eindrucksvoll dar, das auch in der Musik verarbeitet wurde. Der Schwarze Tod fand auch bei dem Bauernkriegspanorama von Werner Tübke Verwendung. Er wurde dort symbolisiert durch einen großen offenen Sarg mit den Todkranken in der Szene *Die Pestkranken*. Die Bewältigung dieser Probleme wurde bis ins 20. Jahrhundert fortgesetzt. Sehr eindrucksvoll ist der Roman „Die Pest“ des französischen Schriftstellers Albert Camus aus dem Jahr 1947.

Die Ebolafieber-Epidemie 2014 bis 2016 in Westafrika und die Ebolafieber-Epidemie 2018 bis 2020 im Osten der Demokratischen Republik Kongo sind nach Fallzahlen und zeitlichem Verlauf ebenfalls Beispiele für epidemische Ausbrüche des Ebolafiebers. Im Falle der Grippe spricht man von einer *Grippewelle*, wenn während einer Saison in verschiedenen Regionen erhebliche Anteile der Bevölkerung infiziert sind (Abb. 3.12)



Abbildung 3.12: Eine Ebola-Behandlungsstation in Westafrika

Epidemien lassen sich nach räumlichen und zeitlichen Merkmalen des Geschehens sowie nach den Bedingungen des Auftretens und Ausbreitens kennzeichnen. So können folgende Arten von Epidemien unterschieden werden:

- **Epidemie mit Punktquelle:** Eine Epidemie, deren Krankheitserreger sich kurzzeitig und gleichzeitig von einer Punktquelle ausgebreitet haben.
- **Kleinraumepidemie:** Eine Häufung von Inzidenzen in einem räumlich begrenzten Milieu, z. B. in einem Heim, einer Kindereinrichtungen oder einer Schule.
- **Streuepidemie:** Vermehrtes Auftreten von Infektionen an verschiedenen Orten, denen eine gemeinsame Ursache zugrunde liegt, die z. B. durch Bevölkerungsbewegungen oder Lebensmitteltransporte gestreut wird.
- **Explosivepidemie:** Bei der Explosivepidemie handelt es sich um eine Epidemie mit schlagartigem Anstieg der Erkrankungszahlen. Häufig sind Epidemien dieser Art mit bestimmten Übertragungsfaktoren assoziiert, beispielsweise als Infektionen, die über Lebensmittel oder Trinkwasser übertragen werden.
- **Mischepidemie:** Eine Mischung aus einer **Explosivepidemie** und einer **Tardiv-Epidemie**, bei der das Infektionsgeschehen zunächst explosiv ist und sich im Verlauf eine Tardiv-Epidemie entwickelt.
- **Tardivepidemie** heißt eine Epidemie mit langsam, aber stetig ansteigenden Erkrankungszahlen. Die Tardiv-Epidemie ist neben der Explosivepidemie eine der beiden klassischen Grundtypen der Epidemie. Im Unterschied zum Begriff der Kontaktepidemie soll mit dem Begriff der Tardiv-Epidemie, ähnlich wie bei der Explosivepidemie, der zeitliche Ablauf einer Epidemie charakterisiert werden, der Übertragungsmodus ist hierbei weitgehend unerheblich. Tardiv-Epidemien können u. a. durch lange Inkubationszeiten begründet sein, durch Infektionswege, die lediglich zur Ansteckung eines einzelnen Menschen führen (sexuell übertragbare Erkrankungen), durch eine geringe Zahl der Überträger, durch eine stark ausgebildete Immunität (etwa infolge vorheriger latenter oder manifester Infektion) oder durch das Zusammenwirken verschiedener Faktoren. Erkrankungen, die zu einer typischen Tardiv-Epidemie führen können sind beispielsweise Pest, Pocken, Grippe, HIV.
- **Kontaktepidemie:** Bei dieser Art von Epidemie nehmen Infektionen und Erkrankungen durch direkte Mensch-zu-Mensch-Kontakte zu, z. B. durch Tröpfcheninfektionen, Kontaktinfektionen, oder sexuell übertragbare Krankheiten.
- **Versandepidemie:** Aufgrund des Versands kontaminierter Lebensmittel treten an verschiedenen Orten vermehrt lebensmittelbedingte Infektionen auf.
- **Epidemie über ein allgemein zugängliches Medium:** Eine Epidemie, deren

Krankheitsauslöser sich über ein allgemein zugängliches Medium wie z. B. Luft, Trinkwasser oder Lebensmittel ausgebreitet haben.

- **Pfropfepidemie:** Eine Epidemie, die aus einer Endemie hervorgeht.
- **Provokationsepidemie:** Eine Epidemie, entstanden nach Aktivierung latenter Infektionen infolge einer Resistenzsenkung in der Bevölkerung.
- **Summenepidemie:** Eine Epidemie, die aus einer endemischen Situation heraus durch eine Summation von Infektionen (Verdichtungswelle, Attraktionswelle) entsteht, weil sich empfängliche Individuen angesammelt haben und sich ein Krankheitserreger mit hoher Kontagiosität ausbreitet.
- **Komplexe Epidemie:** Eine Epidemie, die durch mehrere Krankheitsauslöser induziert wird.
- Als **Pseudoepidemie** wird ein örtlich vermehrtes Auftreten von Fällen einer Infektionskrankheit bezeichnet, das auf eine vermehrte Manifestation von Infektionen durch eine plötzliche Zunahme der Empfänglichkeit (Suszeptibilität) in der Population zurückgeht oder durch erhöhte diagnostische Aktivitäten zustande kommt und nicht durch eine echte Zunahme von Neuinfektionen ausgelöst ist.

Mit Epidemien wurde in der Geschichte auf unterschiedliche Art umgegangen. Auf die **Lepra** im Mittelalter hat die jeweilige Regierungsmacht mit Verbannung, Aussonderung und Ausschließung der Kranken reagiert; angesichts der **Pest** in der frühen Neuzeit hat sie Strategien der Überwachung und Einschließung, dann auch Disziplinarmechanismen, Kontrollnetze und eine minutiöse Beobachtung von Individuen entwickelt; auf die **Pocken** ab Ende des 18. Jahrhunderts hat sie mit Impfmaßnahmen, Immunitätsstrategien, statistischen Erhebungen und Risikoabschätzungen reagiert. Angesichts der **AIDS-Epidemie** habe man zunächst Homosexuelle verfolgt und Risikogruppen denunziert, später hat sich der Umgang mit der Epidemie stärker auf Bereiche außerhalb des geschlossenen Bereichs der medizinischen Beobachtung verlagert.

Die Pocken waren in Amerika völlig unbekannt, die Weißen hatten sie aus Europa mitgebracht, die einheimische Bevölkerung war völlig wehrlos. Es starben etwa 56 Millionen Indianer an dieser Krankheit.

Die **Spanische Grippe** (40–50 Million Todesfälle) trat am Ende des 1. Weltkrieges auf und hatte ebenfalls sehr schwere Folgen. Die US-amerikanischen Zentren für Krankheitskontrolle und -prävention bezeichnen die Grippepandemie von 1918



Abbildung 3.13: Eine an Pocken erkrankte indianische Frau

als die schwerste Pandemie in der jüngeren Geschichte. Die ersten Fälle wurden im März 1918 auf einer Militärbasis in Fort Riley, Kansas, gemeldet, die Krankheit wurde jedoch als Spanische Grippe bekannt, da sie erstmals in der spanischen Presse beschrieben wurde. Der Mangel an wirksamen Impfstoffen oder Behandlungen (die erste Gripeschutzimpfung wurde erst in den 1940er Jahren entwickelt), der Erste Weltkrieg (reisende Soldaten brachten die Krankheit mit) und ein Mangel an Gesundheitspersonal erschwerten den Umgang mit der Pandemie. In den USA ähnelten die lokalen Reaktionen den heutigen COVID-19-Kontrollmaßnahmen: „Beamte verhängten Quarantänen, befahlen den Bürgern, Masken zu tragen und schlossen öffentliche Plätze, einschließlich Schulen, Kirchen und Theater. Den Menschen wurde geraten, Händeschütteln zu vermeiden und zu Hause zu bleiben. Die Bibliotheken stellten die Ausleihe von Büchern ein, und es wurden Vorschriften erlassen, die das Spucken verboten.“

1981 waren die Mediziner ratlos, als bei jungen homosexuellen Männern in den USA seltene Formen von Lungenentzündung und Hautkrebs auftraten. Im Laufe des darauffolgenden Jahres traten diese Erkrankungen auch bei Drogenkonsumenten und Empfängern von Bluttransfusionen auf. Französische Wissenschaftler waren die ersten, die das Humane Immundefizienz-Virus (HIV) isolierten, das das erworbene Immunschwächesyndrom (AIDS) verursacht. Es wird angenommen, dass der Virus im frühen 20. Jahrhundert in Zentralafrika von Schimpansen auf den Men-

schen übertragen wurde und sich schließlich auf alle Kontinente ausbreitete. Es gibt immer noch keine Heilung für AIDS. Obwohl die Entwicklung antiretroviraler Medikamente ab Mitte der neunziger Jahre AIDS für viele Patienten, insbesondere in wohlhabenderen Ländern, mittlerweile zu einer beherrschbaren chronischen Erkrankung machen, starben 2019 immer noch fast 700.000 Menschen an den Folgen einer AIDS-Erkrankung.

Die dritte Pestpandemie von 1855 verursachte 12 Millionen Todesfälle. Sie hatte ihren Ursprung 1855 in China und breitete sich später auf Indien und Hongkong aus. Zumindest zu Beginn wurde sie hauptsächlich von Flöhen verbreitet. Diese Pandemie galt bis 1960 als aktiv, also so lange, bis die Fallzahlen unter ein paar hundert Infizierten lag. Noch heute ist die Beulenpest in einigen Ländern eine Bedrohung für die öffentliche Gesundheit, insbesondere in abgelegenen Gebieten Madagaskars, wo es häufig erneut zu kleinen Ausbrüchen kommt. Die Krankheit, die einst Millionen Menschen das Leben kostete und das Überleben von Nationen bedrohte, kann jetzt leicht mit Antibiotika behandelt werden.

Im Jahr 1545 tauchte unter den Indigenen in Mexiko eine mysteriöse Krankheit namens *Cocoliztli* auf. Ihr werden 5 bis 15 Millionen Todesfälle zugeschrieben. Die Ursache der Krankheit wurde nie eindeutig identifiziert, obwohl DNA, die aus den Zähnen der Opfer entnommen wurde, einen überraschenden neuen möglichen Erreger enthüllt hat. *Zahnpulpa* – das weiche, lebende Gewebe in den Zähnen – ist voll von Blutgefäßen und damit auch mit allen Krankheitserregern, die einst im Blut zirkulierten. Der äußere harte Zahnschmelz schützt die DNA dieser Krankheitserreger für Jahrhunderte. Wissenschaftler kamen zu dem Schluss, dass *Salmonella enterica*, ein Bakterium, das sich durch kontaminierte Lebensmittel oder Wasser ausbreitet, in vielen der Zähne vorhanden war, was stark vermuten lässt, dass dieses Bakterium für die Erkrankung verantwortlich war, obwohl es keineswegs die einzig mögliche Ursache der Krankheit darstellt. Eine ähnliche Epidemie traf Mexiko im Jahr 1576 (Abb. 3.14).

Die Antoninische Pest war eine Pockenepidemie, die um 165 n. Chr. einen Großteil des Römischen Reiches heimsuchte und bis zu 10 Prozent der Bevölkerung das Leben kostete. Laut *Smithsonian Magazine* starben auf dem Höhepunkt der Pest im Jahr 189 n. Chr. in Rom täglich bis zu 2.000 Menschen, wodurch das Imperium bis ins Mark erschüttert wurde: „Die Pest wütete so verheerend innerhalb der Berufsarmeen des Imperiums, dass Angriffe abgebrochen wurden. Sie dezimierte die Aristokratie, und verlassene Bauernhöfe und entvölkerte Städte prägten die Landschaft von Ägypten bis Deutschland.“ Die Römer fanden jedoch geniale Wege, um die Strukturen der Gesellschaft aufrecht zu erhalten, unter anderem kamen Söhne befreiter Sklaven für die Qualifizierung eines lokalen Amtes infrage und das Reich blieb weitere zwei Jahrhunderte bestehen. Der römische Historiker Cassius

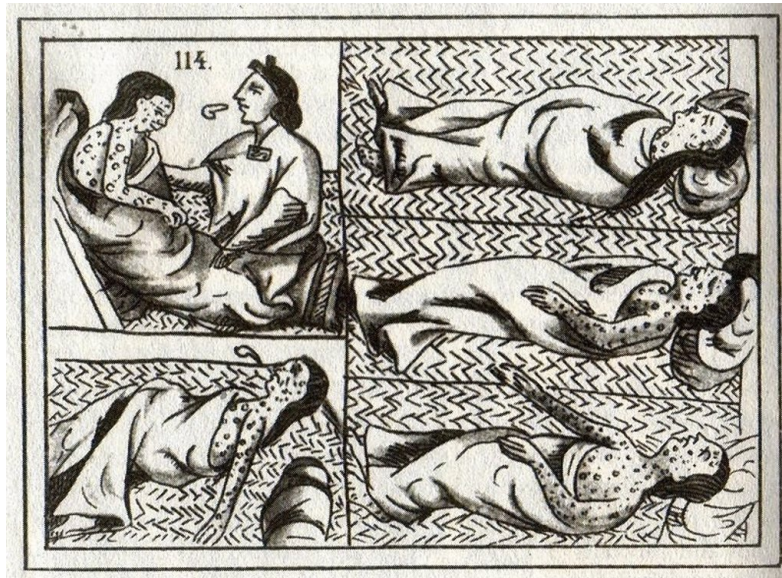


Abbildung 3.14: Opfer der Cocoliztli-Epidemie in Mexico

Dio glaubte, „das Trauma durch das Erleben der Pest kann überwunden werden, wenn eine gut regierte Gesellschaft zusammenarbeitet, um sich zu erholen und sich wieder aufzubauen, und die Gesellschaft, die aus diesen Bemühungen hervorgeht, kann stärker werden als je zuvor.“

Die Spanische Grippe und der Krieg waren nicht die einzigen potenziell tödlichen Bedrohungen, denen Soldaten im Ersten Weltkrieg ausgesetzt waren. Zwischen 1918 und 1922 wurde Typhus, eine durch Kleiderläuse übertragene bakterielle Krankheit, von Serbien durch Russland und die umliegenden Länder – die Ostfront des Ersten Weltkriegs – verbreitet und führte schätzungsweise zu 30.000.000 Infizierten und 3.000.000 Todesfällen. Soldaten lebten auf engstem Raum, drängten sich zusammen, um sich zu wärmen, und machten es den Läusen leicht, sich von Mann zu Mann zu bewegen und bequem in den Nähten der Uniformen zu leben und Blutmahlzeiten zu genießen.



Abbildung 3.15: Typhus-Erreger

Aktuell bewegt natürlich COVID-19 (mindestens 2 Million Todesfälle) die Welt, obwohl die bisherigen Zahlen zeigen, dass das bisherige Ausmaß durchaus längst nicht an die Probleme heranreicht, die in der Vergangenheit bewältigt werden mussten. Am 31. Dezember 2019 informierten die Gesundheitsbehörden in der chinesischen Stadt Wuhan über Dutzende Fälle von Lungenentzündung aufgrund einer unbekannten Ursache. Eine Woche später wurde die Ursache des Ausbruchs als neuer **Coronavirus** identifiziert. Am 20. Januar 2020 wurden die ersten Fälle außerhalb Chinas bestätigt. Bis März hatte sich die Krankheit auf der ganzen Welt verbreitet, internationale Großereignisse wurden abgesagt, und die Länder schlossen ihre Grenzen für Nichtstaatsangehörige, und mehrere Nationen verhängten Sperrungen. Die Weltgesundheitsorganisation schätzt, dass 1,6 Milliarden Menschen aufgrund der **wirtschaftlichen Auswirkungen** der Pandemie Gefahr laufen, ihre Lebensgrundlage zu verlieren. Bis zum jetzigen Zeitpunkt hat die Weltgesundheitsorganisation weltweit 2 Millionen Todesfälle durch den Coronavirus gemeldet, über 400.000 allein in den USA. Die Entwicklung mehrerer Impfstoffe hat die Hoffnung geweckt, dass die Pandemie in einigen Ländern bis Ende dieses Jahres unter Kontrolle sein könnte, aber die Art und Weise, wie wir reisen, arbeiten und Kontakte knüpfen, wird möglicherweise nie wieder dieselbe sein.

Weitere Epidemien größeren Ausmaßes:

- die Russische Grippe von 1889–90 (1 Million Todesfälle)

Im Herbst 1889 brach laut dem History Channel eine Influenza-Pandemie aus, die in der russischen Hauptstadt St. Petersburg mit aller Macht die Hälfte der Bevölkerung infizierte. Innerhalb von wenigen Monaten breitete sie sich in weiten Teilen der Welt über Hauptstraßen, Flüsse und Eisenbahnlinien aus. Sie wird oft als die erste moderne Pandemie angesehen – die erste im Zeitalter des Schienen- und Dampfschiffverkehrs. Solange sich die Menschen problemlos von Stadt zu Stadt und von Land zu Land bewegen konnten, war es nahezu unmöglich, die Ausbreitung zu stoppen.

- Die Hongkong-Grippe von 1968 (1–4 Million Todesfälle)

Die Grippepandemie von 1968 wurde durch einen hoch ansteckenden Grippestamm verursacht, der im Sommer 1968 erstmals in Hongkong auftrat. Innerhalb von zwei Wochen nach dem Auftreten des Virusstammes wurden bereits fast 500.000 Fälle gemeldet. Aus Vietnam zurückkehrende Soldaten verbreiteten die Krankheit in den Vereinigten Staaten. Die höchsten Sterblichkeitsraten gab es bei den infektionsanfälligen Gruppen, nämlich Säuglingen und älteren Menschen. Obwohl ein Impfstoff gegen das Virus entwickelt

wurde, war er erst verfügbar, nachdem die Pandemie in vielen Ländern ihren Höhepunkt bereits erreicht hatte. Der Virusstamm, der die Pandemie von 1968 verursachte, ist noch heute im Umlauf.

- Japanische Pockenepidemie (1 Million Todesfälle)

Pocken tauchten erstmals 735 n. Chr. in Japan auf. Es wird angenommen, dass diese Epidemie das Leben von bis zu einem Drittel der Bevölkerung des Landes forderte. Wie die Pockenepidemie, die Jahrhunderte später Mittelamerika heimsuchte, war auch diese Epidemie eine, die eine Bevölkerung heimsuchte, die zuvor mit Pocken nicht in Berührung gekommen war. Nach der ersten Epidemie kam es mehrere Jahrhunderte lang alle 10 bis 15 Jahre erneut zu Ausbrüchen. Im Durchschnitt starben drei von zehn infizierten Menschen. Pocken wurden schließlich im größten Teil des Landes zu einer endemischen Krankheit. Massenimpfkampagnen, die Ende der 1950er Jahre begannen, führten zur Ausrottung der Krankheit, 1975 in Asien und 1980 auf der ganzen Welt.

- Encephalitis lethargica (bis zu 400.000 Todesfälle)

Eine mysteriöse Pandemie der **Schlafkrankheit**, auch bekannt als Enzephalitis lethargica oder Enzephalitis epidemica, führte zwischen 1916 und 1926 weltweit bis zu einer Million Krankheitsfällen. Die Krankheit hatte in den kalten, feuchten Ebenen und Gräben von Nordfrankreich und Belgien, den Schlachtfeldern des Ersten Weltkriegs, ihren Ursprung. Die Krankheit äußerte sich durch grippeähnliche Symptome, unkontrollierte Augenbewegungen und später durch ein unwiderstehliches Schlafbedürfnis. Diese Lethargie dauert Wochen, in einigen Fällen sogar ein Jahr oder länger. Die Sterblichkeit in dieser akuten Phase der Krankheit lag bei 30 bis 40 Prozent. Viele der Überlebenden erkrankten anschließend an Parkinson. Die Epidemie verschwand ungefähr 10 Jahre später, mit der gleichen Geschwindigkeit, mit der sie ausbrach. Wissenschaftler waren nicht in der Lage, eine einzige Ursache für diese Krankheit ausfindig zu machen, somit bleibt der auslösende Erreger weiterhin unbekannt und existiert möglicherweise heute noch.

- Gelbfieber-Epidemie (100.000—150.000 Todesfälle) **Gelbfieber** ist ein von Mücken übertragener Virus mit einer Sterblichkeitsrate von bis zu 50 Prozent. Eine Reihe von Epidemien zwischen 1793 und 1905 im Süden und Osten der Vereinigten Staaten forderte Tausende Menschenleben. Allem Anschein nach war der Tod durch Gelbfieber äußerst unschön: Die Opfer hatten eine Vielzahl unangenehmer Symptome: Gelbsucht, Schüttelfrost, Übelkeit, Kopf-



Abbildung 3.16: Der Erreger der Schlafkrankheit

schmerzen, Fieber, Krämpfe, Delirium. In den späteren Stadien der Krankheit bluteten die Opfer durch Augen, Nasen und Ohren. Menschen, die die Krankheit überlebten, waren immun. Im New Orleans des 19. Jahrhunderts wurden Überlebende laut NPR als akklimatisiert bezeichnet, und nicht akklimatisierte Menschen hatten Schwierigkeiten, Arbeit und Unterkunft zu finden. Ab dem frühen 20. Jahrhundert wurde eine Reihe von Impfstoffen entwickelt, mehrere Impfstoffforscher infizierten sich jedoch dabei versehentlich mit der Krankheit und kamen ums Leben. Gelbfieber fordert immer noch etwa 30.000 Menschenleben pro Jahr, hauptsächlich in Afrika. Für internationale Reisen in viele afrikanische Länder ist ein Impfnachweis erforderlich.

Man sieht an diesen historischen Fakten, dass die gegenwärtige Corvid-19-Situation sehr unangenehm ist, aber durchaus nichts Neues darstellt. Im Sinne von Fraktalen ist die Ausbreitung immer damit zu begründen, dass ein Punkt, an dem Krankheitserreger existieren, Kontakt zu einem Nachbarpunkt hat und der Erreger mit einer gewissen Wahrscheinlichkeit, die sehr hoch sein kann, übertragen wird. Deshalb ist die maximale Reduktion von Kontakten zu anderen Infektionskandidaten das Wichtigste. In der heutigen Sprache: Test und Quarantäne sind die wichtigsten Elemente der Bekämpfung, sie verhindern die Übertragung auf andere Personen und ermöglichen die auf eine Person reduzierte Behandlung der Krankheit. Die Impfung schützt die geimpfte Person, kann aber die Existenz einer Infektion verschleiern und muss durch Tests ergänzt werden. Unbeschränkte Kontakte geimpfter Personen mit beliebigen anderen Personen sind ohne Test durchaus riskant.

Die Problematik von Epidemien ist nicht auf Menschen beschränkt. Bekannt ist vor allem die Vogelgrippe, die zum ersten Mal 1959 in Schottland auftrat und sich dann in Europa, Asien und Teilen von Afrika auftrat. Für Tiere gibt es seit

langem eine Reihe wirksamer Totimpfstoffe gegen Influenzaviren, ferner wurden seit 2005 mehrere Lebendimpfstoffe zur Marktreife entwickelt. In Deutschland ist die Impfung gegen die Geflügelpest allerdings nur in Ausnahmefällen erlaubt. Für Hühner gibt es auch einen zugelassenen DNA-Impfstoff.

Auch die Weltgesundheitsorganisation warnte wiederholt nachdrücklich vor Impfungen, da geimpfte Tiere nicht mehr von virentragenden Tieren unterschieden werden könnten. Überdies könnten geimpfte, infizierte Vögel zu Überträgern der Viren werden, ohne selbst Symptome zu zeigen. Auch bestehe die Gefahr, dass die Viren in unerkannt infizierten Tieren mutieren und dass diese Erbgut-Veränderungen sich leichter ausbreiten könnten als in nicht-geimpften Beständen, da diese ja nach jedem Ausbruch getötet würden; eine derartige Entwicklung wurde nach einem H5N2-Ausbruch in Mexiko beobachtet.

Ein anderes Beispiel aus dem Tierreich ist die **Schweinepest**. Die Klassische Schweinepest (KSP) ist seit 1833 als Infektionskrankheit bekannt. Diese Virusinfektion tritt mit Ausnahme Nordamerikas, Australiens, Neuseelands und Teilen von Europa und Südamerikas weltweit auf. Sie zählt zu den gefährlichsten Schweinekrankheiten überhaupt und ist bis heute schwer kontrollierbar und nicht getilgt. Die klassische Schweinepest ist von der Afrikanischen Schweinepest (ASP) abzugrenzen. Trotz der ähnlichen Symptome sind ASP- und KSP-Erreger nicht näher verwandt.

Der Schweinepesterreger ist das Pestivirus C, auch Klassischer Schweinepest-Virus genannt. Obwohl Verwandtschaften zu Erregern anderer Krankheiten bestehen, ist jedoch keine andere Tierart empfänglich.

Als ständiges Erregerreservoir kann das Wildschwein gelten. Neben dem Kontakt mit Wildschweinen stellt häufig der Zukauf von bereits infizierten, aber nicht sichtbar kranken Schweinen eine Ansteckungsquelle dar. Daneben kann der Virus aber auch durch sogenannte Vektoren (verunreinigte Fahrzeuge und Gerätschaften, Kleidung oder Speiseabfälle) übertragen werden. Die Ansteckung innerhalb eines Bestandes erfolgt dann direkt von Tier zu Tier hauptsächlich peroral oder über die Atemwege. Die Inkubationszeit hängt von der Virulenz des jeweiligen Erregers ab. Sie kann zwischen 2 Tagen bis über 5 Wochen betragen. Das Virus vermehrt sich zunächst in den Mandeln und den Lymphknoten des Rachenraums. Bereits nach 24 Stunden befindet sich der Erreger im Blutkreislauf und erreicht innerhalb von einer Woche seine maximale Konzentration. Sofern der Erreger sich im Blut befindet, wird er über Harn, Speichel, Kot, Augen- und Nasensekret ständig ausgeschieden. Dieses ist auch der Grund für die seuchenhafte Ausbreitung der Krankheit; große Schweinebestände mit mehreren tausend Tieren können innerhalb von einer Woche vollständig infiziert sein.

Die **Afrikanische Schweinepest** ist eine gesonderte Entwicklung, die mit der europäischen Schweinepest nicht verwechselt werden darf.

Auch hier ist eine vollständige Trennung der Populationen das sicherste Mittel zur Verhinderung weiterer Ansteckungen. Man denke etwa an die Sperrzäune, die im Osten Deutschlands gegen die Einwanderung von Wildschweinen aus Polen errichtet wurden.

3.10 Digitalisierung und Künstliche Intelligenz

Alle im Verlaufe von Untersuchungen gewonnenen Informationen liegen gegenwärtig in digitaler Form vor. Sie können in beliebiger Größe angesehen werden, und das nicht nur einmal, sondern sehr oft. Sie können auch weitergeleitet werden, so dass kritische Probleme von mehreren Spezialisten beurteilt werden können. Die Bildverarbeitung bietet viele Möglichkeiten, die Bilder zu vergrößern, zu drehen usw. Außerdem kann man sie mit anderen Bildern vergleichen. Die Übereinstimmung mit Bildern von erkrankten oder gesunden Personen kann die Diagnose wesentlich unterstützen.

Dazu kommt die zeitliche Entwicklung. Oft hat man von einer Person die Entwicklung über einen längeren Zeitraum zur Verfügung, beispielsweise wenn eine Person regelmäßig zu Kontrolluntersuchungen geht. Erkrankt eine Person zu einem bestimmten Zeitpunkt, so kann man die vorhandenen Bilder überprüfen und untersuchen, ab wann Veränderungen sichtbar sind. Diese Bilder kann man dann speziell speichern und mit Bildern anderer Personen vergleichen, was wiederum eine zeitige Diagnose und die entsprechenden Maßnahmen unterstützt.

In den letzten zehn Jahren hat man beim Trainieren von Neuronalen Netzen große Erfolge erreicht. Der Begriff „Neuronales Netz“ beschreibt Strukturen im menschlichen Hirn. Nervenzellen (Neuronen) sind mittels Synapsen miteinander verbunden und spannen dadurch Nervennetze (neuronale Netze) auf. Die Neuronen bilden dabei die Knotenpunkte des Netzes. Neuronen nehmen Reize von außerhalb des Netzes oder Reize von anderen Neuronen auf. Reize können etwa Sinneswahrnehmungen sein.

Ein Neuron hat üblicherweise einen Eingang und verzweigt sich in mehrere Ausgänge. Übersteigt die Stärke des eingehenden Reizes ein bestimmtes Maß, „feuert“ das Neuron, das heißt, es gibt den Reiz über Synapsen an andere Neuronen weiter. Die Übertragung kann dabei durch elektrische Impulse erfolgen oder als chemische

Reaktion mit Hilfe von Botenstoffen. So werden Reize und Informationen im Gehirn übermittelt und verarbeitet. Das ist für die Verarbeitung von Sinneswahrnehmung genauso wichtig wie für Lernprozesse, neuronale Netze können komplizierte Muster erlernen. Dabei ist es nicht nötig, dass sie die abstrakten Regeln hinter den Mustern verstehen.

Diese in der Natur vorhandene Struktur wird nun in einem Modell durch ein künstliches Neuronales Netz nachgebildet. Die Information wird durch die Input-Neuronen aufgenommen und durch die Output-Neuronen ausgegeben. Die Hidden-Neuronen liegen dazwischen und bilden innere Informationsmuster ab. Die Neuronen sind miteinander über sogenannte Kanten verbunden. Je stärker die Verbindung ist, desto größer die Einflussnahme auf das andere Neuron.

Es wird die folgende Struktur verwendet:

- **Eingabeschicht:** Die Eingangsschicht versorgt das neuronale Netz mit den notwendigen Informationen. Die Input-Neuronen verarbeiten die eingegebenen Daten und führen diese gewichtet an die nächste Schicht weiter.
- **Verborgene Schicht:** Die verborgene Schicht befindet sich zwischen der Eingabeschicht und der Ausgabeschicht. Während die Ein- und Ausgabeschicht lediglich aus einer Ebene bestehen, können beliebig viele Ebenen an Neuronen in der verborgenen Schicht vorhanden sein. Hier werden die empfangenen Informationen erneut gewichtet und von Neuron zu Neuron bis zur Ausgabeschicht weitergereicht. Die Gewichtung findet in jeder Ebene der verborgenen Schicht statt. Die genaue Prozessierung der Informationen ist jedoch nicht sichtbar. Daher stammt auch der Name, verborgene Schicht. Während in der Ein- und Ausgabeschicht die eingehenden und ausgehenden Daten sichtbar sind, ist der innere Bereich des Neuronalen Netzes im Prinzip eine Black Box.
- **Ausgabeschicht:** Die Ausgabeschicht ist die letzte Schicht und schließt unmittelbar an die letzte Ebene der verborgenen Schicht an. Die Output-Neuronen beinhalten die resultierende Entscheidung, die als Informationsfluss hervorgeht. Tiefes Lernen ist eine Hauptfunktion eines Künstlichen Neuronalen Netzes und funktioniert wie folgt: Bei einer vorhandenen Netzstruktur bekommt jedes Neuron ein zufälliges Anfangsgewicht zugeteilt. Dann werden die Eingangsdaten in das Netz gegeben und von jedem Neuron mit seinem individuellen Gewicht gewichtet. Das Ergebnis dieser Berechnung wird an die nächsten Neuronen der nächsten Schicht oder des nächsten Layers weitergegeben, man spricht auch von einer „Aktivierung der Neuronen“. Eine Berechnung des Gesamtergebnisses geschieht in der Ausgabeschicht.

Es gibt viele verschiedene Architekturen von neuronalen Netzwerken.

- **Perceptron:** Das ist das einfachste und älteste neuronale Netz. Es nimmt die Eingabeparameter, addiert diese, wendet die Aktivierungsfunktion an und schickt das Ergebnis an die Ausgabeschicht. Das Ergebnis ist binär, also entweder 0 oder 1 und damit vergleichbar mit einer Ja- oder Nein-Entscheidung. Die Entscheidung erfolgt, indem man den Wert der Aktivierungsfunktion mit einem Schwellwert vergleicht.

Bei Überschreitung des Schwellwertes, wird dem Ergebnis eine 1 zugeordnet, hingegen 0 wenn der Schwellwert unterschritten wird. Darauf aufbauend wurden weitere Neuronale Netzwerke und Aktivierungsfunktionen entwickelt, die es auch ermöglichen, mehrere Ausgaben mit Werten zwischen 0 und 1 zu erhalten. Am bekanntesten ist die Sigmoid-Funktion, in dem Fall spricht man auch von Sigmoid-Neuronen.

- **Feed forward neural networks:** Der Ursprung dieser neuronalen Netze liegt in den 1950 Jahren. Sie zeichnen sich dadurch aus, dass die Schichten lediglich mit der nächst höheren Schicht verbunden sind. Es gibt keine rückwärts gerichteten Kanten. Der Trainingsprozess eines Feed Forward Neural Network (FF) läuft dann in der Regel so ab:

1. alle Knoten sind verbunden;

2. die Aktivierung läuft von der Eingangs- zur Ausgangsschicht. Zwischen der Eingangs- und Ausgangsschicht liegt mindestens eine weitere Schicht. Wenn besonders viele Schichten zwischen Eingangs- und Ausgangsschicht sind, spricht man von Deep Feed Forward Neural Networks,,

- **Faltende Neuronale Netze** die auch **Convolutional Neural Networks (CNN)** genannt werden, sind Künstliche Neuronale Netzwerke, die besonders effizient mit 2D- oder 3D-Eingabedaten arbeiten können. Für die Objektdetektion in Bildern verwendet man insbesondere CNNs.

Der große Unterschied zu den klassischen Neuronalen Netzen besteht in der Architektur von CNNs und damit kann auch der Name „Convolution“ oder „Faltung“ erklärt werden. Die verborgene Schicht basiert bei CNNs auf einer Abfolge von Faltungs- und Poolingoperationen. Bei der Faltung wird ein so genannter Kernel über die Daten geschoben und währenddessen eine Faltung gerechnet, was vergleichbar mit einer Multiplikation ist. Die Neuronen werden aktualisiert. Die anschließende Einführung eines Poolinglayers sorgt dafür, dass die Ergebnisse vereinfacht werden. Nur die wichtigen Informatio-

nen bleiben anschließend erhalten.

Dies sorgt zudem dafür, dass sich die Zahl der 2D- oder 3D-Eingangsdaten verringert. Wird dies immer weiter fortgeführt, ergibt sich am Ende in der Ausgabeschicht ein Vektor der „fully connected layer“. Dieser hat vor allem in der Klassifikation eine besondere Bedeutung, da er so viele Neuronen wie Klassen enthält und die entsprechende Zuordnung über eine Wahrscheinlichkeit bewertet.

- **Recurrent Neural Networks (RNN)** fügen den KNN wiederkehrende Zellen hinzu, wodurch neuronale Netze ein Gedächtnis erhalten. Das erste künstliche, neuronale Netzwerk dieser Art war das Jordan-Netzwerk, bei dem jede versteckte Zelle ihre eigene Ausgabe mit fester Verzögerung – eine oder mehrere Iterationen – erhielt. Ansonsten ist es vergleichbar mit den klassischen Feed Forward Netzen.

Natürlich gibt es viele Variationen – wie die Übergabe des Zustands an Eingangsknoten, variable Verzögerungen usw., aber die Grundidee bleibt die gleiche. Diese Art von NNs wird insbesondere dann verwendet, wenn der Kontext wichtig ist. Denn dann beeinflussen Entscheidungen aus vergangenen Iterationen oder Proben maßgeblich die aktuellen. Da rekurrente Netze jedoch den entscheidenden Nachteil haben, dass sie über die Zeit instabil werden, ist es inzwischen üblich, sogenannte Long Short-Term Memory Units (kurz: LSTMs) zu verwenden. Diese stabilisieren das RNN auch für Abhängigkeiten, die über einen längeren Zeitraum bestehen. Das häufigste Beispiel für solche Zusammenhänge ist die Textverarbeitung – ein Wort kann nur im Zusammenhang mit früheren Wörtern oder Sätzen analysiert werden. Ein weiteres Beispiel ist die Verarbeitung von Videos, beispielsweise beim autonomen Fahren. Objekte, innerhalb von Bildsequenzen, werden erkannt und über die Zeit verfolgt.

Als Schwellwert-Funktion wird häufig die Sigmoid-Funktion

$$f(x) = \frac{1}{1 + e^{-x}} \quad (3.4)$$

verwendet. Für $x \rightarrow +\infty$ strebt sie gegen $+1$, für $x \rightarrow -\infty$ strebt sie gegen -1 . Für $x = 0$ hat sie den Wert $\frac{1}{2}$. Diese Funktion kann man in der Weise zur Entscheidung zwischen 0 und 1 verwenden, dass man für $f(x) < \frac{1}{2}$ den Wert 0 ausgibt, für $f(x) \geq \frac{1}{2}$ den Wert 1.

Ein neuronales Netzwerk wird in drei Schritten trainiert.

- Schritt 1: Bestimmung der Ausgabe für bestimmte Datensätze. Hierfür wer-

den alle Trainingsdatensätze als Eingabe für das Modell verwendet und die berechneten Ergebnisse gesammelt.

- Schritt 2: Bestimmung der Abweichung bzw. des Fehlers der Ausgabe von der gewünschten Ausgabe. Mit Hilfe der Fehlerfunktion bestimmt man für jeden Trainingsdatensatz, wie groß die Abweichung des tatsächlichen Ergebnisses vom erwarteten Ergebnis ist. Der durchschnittliche Fehlerwert aller Trainingsdatensätze bildet nun die Grundlage für den dritten Schritt.
- Schritt 3: Rückwärtsrechnen des Fehlers und Anpassung der Gewichtungen des Netzes. Beim Rückwärtsrechnen (Backpropagation) wird nun der durchschnittliche Wert verwendet, um die Richtung des Parameters zu bestimmen. Dabei muss man herausfinden, inwieweit sich die Gewichtungen des Netzes verändern muss, um den Fehler zu minimieren.

Dafür bilden wir die Ableitung der Fehlerfunktion. Diese gibt nun an, welche Steigung die Fehlerfunktion für die aktuellen Parameter besitzt. Ist die Steigung der Fehlerfunktion für einen Parameter größer als 0, dann wissen wir, dass der Parameter kleiner werden muss. Wenn wiederum die Steigung der Fehlerfunktion kleiner als 0 ist, dann muss der Parameter größer werden. Um das gewünschte Ergebnis zu erhalten, muss der Vorgang nun für jeden Parameter des Models durchgeführt werden.

Durch die Entwicklung immer leistungsfähigerer Computer kann man die Größen natürlich stark erhöhen, für die Auswertung medizinischer Daten etwa auf die Größe eines Bildschirms. Bei einem 27-Zoll-Monitor (68,58 cm) muss man $1920 \times 1080 = 2.073.600$ Pixel berücksichtigen. Das bedeutet, dass die Eingangsschicht 2.073.600 Neuronen enthalten muss. Diese Größenordnung ermöglicht natürlich sehr feinsinnige Klassifizierungen.

4 Fraktale Strukturen in der menschlichen Gesellschaft

Wir gehen noch einmal zur Brownschen Bewegung zurück (Abschnitt 2.9) und verwenden dieses Konzept für Gruppierungen von Menschen. Man erkennt ziemlich schnell, dass man hier sehr viele fraktale Strukturen antrifft. Viele interessante Ausführungen zu diesen Problemen findet man in dem Buch von Elias Canetti: „Masse und Macht“.[18]

4.1 Menschengruppen und ihr Verhalten

Das Auftreten von Menschengruppen ist eine alltägliche Erscheinung, die jedem bekannt ist. Schulklassen, Reisegruppen, Mannschaften verschiedener Sportarten, Zuschauer bei Wettkämpfen - das sind im Sinne einer Brownschen Bewegung Teilchen, die sich zusammengefunden haben und zusammenbleiben, wobei hier ein gemeinsames Ziel existiert. Es handelt sich um eine offene Masse, alle Mitglieder sind freiwillig zusammengekommen und wollen auch ein Teil dieser Masse bleiben.

Das ändert sich ganz schnell, wenn eine echte oder vermeintliche Gefahr existiert. Meldungen aus der Welt des Sportes, vor allem des Fußballs, gehören fast schon zum Alltag.

„Der Fußball hat viele Tragödien auf den Tribünen erlebt. Das Desaster vom Brüsseler Heysel-Stadion 1985 wurde live in 77 Ländern übertragen. Oftmals wurden Menschen erdrückt, weil die Ausgänge verschlossen waren. Und in Moskau starben 340 Zuschauer bei einem Europapokal-Spiel. Die schlimmsten Stadionkatastrophen des Fußballs (Süddeutsche Zeitung):

- 24. Mai 1964: Im Nationalstadion der peruanischen Hauptstadt Lima brechen bei einem Olympiaqualifikationsspiel zwischen den Fußballmannschaften von Peru und Argentinien Krawalle aus, nachdem der Schiedsrichter kurz vor Schluss einen Treffer nicht gibt. Weil Peru damit die Qualifikation verpassen



Abbildung 4.1: Panik am Ausgang eines Fußball-Stadions

würde, kommt es zu Randalen. Die Polizei schießt mit Tränengas, woraufhin Zuschauer zu den Ausgängen drängen. Die Tore sind aber noch verschlossen und nur nach innen zu öffnen. Trauriges Ende: 318 Tote, mehr als 500 Verletzte.

- 17. September 1967: Es gibt 44 Tote und 600 Verletzte in Kayseri/Türkei, als nach einem umstrittenen Tor Zuschauer einander mit Pistolen, Messern und abgebrochenen Flaschen bekämpfen.
- 23. Juni 1968: 74 Menschen kommen in der argentinischen Hauptstadt Buenos Aires ums Leben, und mehr als 150 werden verletzt. Fans waren nach der Partie zwischen River Plate und Boca Juniors zu einem Ausgang geströmt, der verschlossen war. Ausgelöst wurde die Massenpanik durch die Boca-Fans. Diese hatten brennende Zeitungen in die Menge geworfen. Gerichtlich konnte der Vorfall nie ganz geklärt werden. Niemand ist wegen gewalttätigen Vorgehens verurteilt worden.
- 2. Januar 1971: 66 Tote im Glasgower Ibrox Park, als Menschenmassen nach dem Lokalderby Rangers - Celtic ein Gelände durchbrechen.
- 17. Februar 1974: Beim Spiel Zamalek Kairo gegen Dukla Prag versuchen Zuschauer, die Absperrungen zu einem Stadion in der ägyptischen Hauptstadt Kairo zu durchbrechen, 49 werden zu Tode getrampelt.
- 20. Oktober 1982: Bis zu 340 Menschen sterben bei der Europapokal-Partie

zwischen Spartak Moskau und dem HFC Haarlem aus den Niederlanden. Die Polizei soll Fans noch vor dem Abpfiff durch zu enge und vereiste Abgänge aus dem Moskauer Stadion gedrängt haben. Als ein spätes Tor fiel, versuchten die Fans ins Stadion zurückzugelangen, es kommt zur Massenpanik. Die russischen Behörden sprechen offiziell von 61 Toten und sehen keine Schuld bei der Polizei.

- 1. Mai 1985: Eine Zigarette setzt die Holztribüne in Bradford (England) in Brand, bei dem Großbrand kommen 56 Menschen um, Hunderte werden verletzt.
- 29. Mai 1985: Beim Endspiel um den Europapokal der Landesmeister zwischen Juventus Turin und dem FC Liverpool stürmen englische Fans im Brüsseler Heysel-Stadion einen Block mit italienischen Fans, Panik bricht aus, 39 Menschen sterben, 400 Besucher werden verletzt. Das Sterben wurde in 77 Ländern live übertragen. Die Barriere zwischen einem Block Liverpool- und Juve-Fans bestand aus einem Hühnerzaun und acht Polizisten. Nach dem Angriff der Engländer flüchtete die Italiener nach unten, wo eine Begrenzungsmauer barst. Die Schreckensbilder gingen um die Welt, dennoch wurde das Finale angepfiffen. Juventus gewann 1:0.
- 12. März 1988: In Kathmandu (Nepal) kommen mindestens 93 Menschen bei einer Massenpanik ums Leben. Sie hatten versucht, sich vor einem Hagelsturm in Sicherheit zu bringen, verschlossene Stadiontüren wurden ihnen zum Verhängnis.
- 15. April 1989: Bei der Partie Nottingham Forest gegen FC Liverpool im Hillsborough Stadion in Sheffield versuchen Fans, sich auf eine bereits überfüllte Tribüne zu drängen, Menschen werden gegen die Begrenzungszäune gequetscht, es gibt 95 Tote, mehr als 200 Verletzte.
- 13. Januar 1991: In Orkney (Südafrika) sterben 40 Menschen bei Ausschreitungen, die meisten werden totgetrampelt oder zerquetscht.
- 5. Mai 1992: 15 Tote und rund 1.300 Verletzte beim Einsturz einer Zuschauer-Hilfstribe kurz vor Beginn des französischen Pokal-Halbfinals SSC Bastia - Olympique Marseille.
- 16. Oktober 1996: Bei der WM-Qualifikations-Partie zwischen Guatemala und Costa Rica kommen in Guatemala-Stadt bei einer Massenpanik 78 Menschen ums Leben.

- 11. April 2001: Bei der Partie zwischen den Kaizer Chiefs und den Orlando Pirates kommen in Johannesburg in Südafrika 47 Menschen ums Leben. Fans hatten versucht, in das überfüllte Stadion zu gelangen und waren in Stacheldraht gedrückt worden.
- 9. Mai 2001: In Ghanas Hauptstadt Accra werfen frustrierte Fans nach dem Meisterschaftsspiel der Accra Hearts of Oak gegen Kumasi Ashanti Kotoko abmontierte Stadionsitze und Flaschen auf das Spielfeld. Die Polizei setzt Tränengas ein und schließt die Stadione. Daraufhin bricht unter den 40.000 Zuschauern Panik aus, 124 werden zu Tode getrampelt oder ersticken, mehr als 150 Personen werden verletzt.

Nur die beiden Brände stellen eine echte Gefahr da, die anderen Opfer sind auf das kopflose Verhalten der Zuschauer zurückzuführen.

Bei Demonstrationen kommt es ebenfalls oft zu gewaltsamen Auseinandersetzungen, entweder zwischen feindlichen Gruppen oder zwischen Demonstranten und der Polizei. Die einzelnen Partikel verlieren ihre Individualität und lassen sich in einer Herde treiben.

4.2 Das Kommunikationsverhalten von Gruppen

Die Anzahl von Mitgliedern einer Gruppe ist nicht sehr hoch, zehn bis fünfzehn dürften die normalen Werte sein. Wird eine Gruppe neu zusammengestellt, dann beginnen die Gruppenmitglieder zu kommunizieren.

Ungelöste Konflikte, eine falsche Erwartungshaltung oder aktuelle Schwierigkeiten fließen in die Kommunikation ein. Feste Gesprächsregeln tragen dazu bei, das eigene Verhalten zu reflektieren und eine selbstkritische Haltung einzunehmen. Die folgenden Regeln sollte man beachten:

- Erteile keine ungebetenen Ratschläge.

Je nach Situation und Person können ungebetene Ratschläge Trotzreaktionen hervorrufen. Man kann beste Absichten hegen, aber nicht wenige Menschen fühlen sich bevormundet und nicht genügend ernst genommen, wenn sie ungefragt Feedback bekommen. Eine Möglichkeit als Gesprächsregel wäre es, das Gegenüber zu fragen: Möchten Sie meine Meinung dazu hören? Mir fallen spontan ein paar Gedanken dazu ein – interessiert es Sie zu erfahren, wie

eine weitere Möglichkeit aussähe?

- Man soll nicht von sich auf die Allgemeinheit schließen.

Es geht zwar in Diskussionen nicht darum, den Egozentriker hervorzukehren, aber verstecken muss man sich auch nicht.

- Andere Menschen muss man ausreden lassen.

Die wohl wichtigste Gesprächsregel ist die, andere Menschen nicht sofort zu unterbrechen, sondern sie ausreden zu lassen. Das fällt nicht immer leicht, vor allem, wenn eine Aussage getätigt wird, bei der man sich direkt angegriffen fühlt. Dennoch ist es wichtig, dass andere Menschen die Chance bekommen, ihren Standpunkt darzulegen – nicht jeder kann sich unglaublich pointiert und geschliffen ausdrücken. Manche kritischen Passagen klären sich schon von allein, wenn der Gesprächspartner die Möglichkeit hatte, seinen Satz zu vollenden.

- Vermeide Monologe.

Manche Menschen neigen dazu zu monologisieren. Das ist nicht unbedingt Boshaftigkeit, sondern teilweise mangelnde Präzision. Es ist besser, indem Sie versuchen, sich auf die Kernaussage zu konzentrieren und sich nicht in Nebenschauplätzen verlieren. So geben Sie anderen Menschen die Gelegenheit, ebenfalls etwas zum Thema beizutragen oder Fragen zu stellen. Außerdem beugen Sie Ermüdungserscheinungen beim Gegenüber vor.

- Melden Sie Redebedarf an.

Auch wenn Sie nicht mehr in der Schule sind und nicht mit dem Finger aufzeigen müssen: Melden Sie an, wenn Sie Redebedarf haben. Vielleicht hat der Kollege im Meeting versehentlich auf die Folie vom Vorjahr zurückgegriffen und falsche Zahlen genannt? Bevor sich auf der Basis falscher Annahmen neue Irrtümer entwickeln, kann es sinnvoll sein, frühzeitig einzuschreiten. Tun Sie das trotzdem mit größtmöglichen Respekt: Herr Müller, wenn ich Sie kurz unterbrechen darf, mir scheint, dass hier auf Datenmaterial vom Vorjahr zurückgegriffen wurde. . .

- Treffen Sie persönliche Aussagen statt zu fragen.

Fragen, besonders Verständnisfragen, haben absolut ihre Berechtigung in Diskussionen. Eine Gesprächsregel sollte sein, Fragen nur dafür zu benutzen und

nicht versteckte Kritik zu äußern. Rhetorische Fragen oder verpackte Vorwürfe gehören nicht dazu, ein Negativbeispiel: Haben Sie mal überlegt, was für einen Aufwand das zur Folge hat? Solche Fragen führen den anderen vor und können herablassend wirken. Das wollen Sie vermeiden, daher sollten Sie lieber direkte inhaltliche Rückmeldungen geben: Ich könnte mir vorstellen, dass aufgrund unserer letztjährigen Erfahrung der Aufwand sehr groß wird. Bereits letztes Jahr waren wir personell zu knapp besetzt, daher wäre eine Möglichkeit ...

- Senden Sie Ich-Botschaften.

Ein Klassiker unter den Gesprächsregeln ist der, dass Sie sogenannte Ich-Botschaften senden: Ich konnte beobachten, dass häufiger vergessen wird, den Toner im Kopierer nachzufüllen. Ich-Botschaften haben den großen Vorteil, dass Ihr Gegenüber sich nicht sofort angegriffen fühlt, wenn Sie etwas Kritisches formulieren. Du-Botschaften wirken automatisch anklagend und führen leicht dazu, dass Ihr Gesprächspartner in die Verteidigung geht. Im schlimmsten Fall verhärten sich die Fronten. Die Krönung ist eine Kombination mit „immer“, „ständig“ oder „nie“ – so kann eine Unterhaltung nur in die Hose gehen. Eine weitere Gesprächsregel sollte hier sein, diese Worte nach Möglichkeit zu vermeiden.

- Sprechen Sie andere direkt an.

Sie sollten anwesende Kollegen ins Gespräch mit einbeziehen. Das gelingt, indem Sie mit ihnen sprechen und nicht über sie. Negativbeispiel: Kollege Schneider ist der Ansicht, das wird nichts vor Ende des Jahres. Besser: Herr Schneider, so wie ich Sie vorhin verstanden habe, sehen Sie den zeitlichen Rahmen auch. . .

- Sprechen Sie nur, wenn Sie das Wort haben.

Diese Gesprächsregel ist eigentlich eine Selbstverständlichkeit, aber manchmal geht es auch in Meetings wie im Kindergarten zu: Es sollte nach Möglichkeit nur eine Person gleichzeitig sprechen. Das hat mehrere Gründe: Zum einen ist es unhöflich demjenigen gegenüber, der das Wort hat. Zum anderen erschwert ein Durcheinander das Verständnis und die Konzentration auf den Gesprächspartner. Zum dritten: Was Sie oder eine andere Person zu sagen haben, könnte für den Rest der Runde ebenfalls wichtig sein. Wenn alle gleichzeitig reden, geht das aber unter.

- Üben Sie Rücksichtnahme.

Jede Gruppe hat Teilnehmer, die eher mehr beitragen wollen und andere, die deutlich zurückhaltender sind. Beobachten Sie Ihre Gesprächspartner: gibt es vielleicht Anzeichen, dass jemand etwas sagen möchte, sich aber nicht traut? Eine Gesprächsregel ist Rücksichtnahme zu zeigen, sich selbst an der einen oder anderen Stelle zurückzunehmen und andere zu ermutigen, damit sie ebenfalls Gelegenheit haben, sich zu äußern.

Eine weit verbreitete Erscheinung entspricht dem Verhalten des Cantorschen Staubes: man flüstert jemand geheimnisvoll und vertraulich eine falsche oder verzerrte Information ins Ohr. Diese Information verbreitet sich oft sehr schnell über viele Leute hinweg. Mit dem Gegenteil kann man eine zweite Meinung in Umlauf setzen, die der ersten Meinung diametral entgegengesetzt ist.

Ein Gerücht ist eine unverbürgte Nachricht, die stets von allgemeinem bzw. öffentlichem Interesse ist, sich diffus verbreitet und deren Inhalt mehr oder weniger starken Veränderungen unterliegt. Das wichtigste Merkmal eines Gerüchts stellt die Ungesicherheit der weitergegebenen Information dar. Gemäß der Klassifikation von Robert H. Knapp (1944) lassen sich Gerüchte in drei Kategorien einteilen:

- Wunschgerüchte (Hoffnung auf ein positives Ereignis)
- Aggressionsgerüchte (Feindseligkeit gegenüber anderen)
- Angstgerüchte (Furcht vor einem negativen Ereignis)

Außerdem lassen sich Gerüchte hinsichtlich ihres Inhalts klassifizieren, wie beispielsweise organisationale Gerüchte (DiFonzo et al., 1994) oder Produktgerüchte (Miller, 2013).

Als Flüsterpropaganda bezeichnet man einen Vorgang, bei dem meist durch die Politik geheim gehaltene Vorkommnisse weitererzählt werden und so langsam unter die Bevölkerung und damit in die Öffentlichkeit gelangen. Diese häufig in totalitären Staaten vorkommende Verbreitung von Nachrichten kann zu Gerüchten führen.

Latrinenparolen sind umgangssprachlich abwertend bezeichnete Gerüchte, die zu meist irreführend oder falsch sind und heimlich verbreitet werden. Das Wort stammt aus der Soldatensprache, da sich in Kasernen oder anderen Unterkünften an der dortigen Sickergrube oder auch Latrine alle Mannschaftsgrade zur gemeinsamen Entleerung trafen und wo auch Informationen ausgetauscht und dann weitergegeben worden sind. Synonyme sind Latrinengerücht oder derber Scheißhausparole.

Stammtischparolen bezeichnen stereotype Versatzstücke einer lokalen Meinungsbildung und umfassen ebenfalls Gerüchte.

Bei allen diesen Erscheinungen ist natürlich die Selbstähnlichkeit der Kommunikation leicht zu sehen.

Die Werbung liegt oft in der Nähe von Gerüchten und Parolen. Man vertraut darauf, dass sich eine endlose Wiederholung eines Namens für ein Produkt im Kopf eines Kunden so festsetzt, dass er beim nächsten Einkauf dieses Produkt bevorzugt. Oft ist aber das Gegenteil der Fall. Bei Amazon hat man festgestellt, dass ein Kunde sich nach drei bis vier Angeboten entscheidet oder die Suche beendet. „Viel hilft viel“ stimmt nicht immer. Noch schlimmer ist es, wenn man etwas online gekauft hat und anschließend immer wieder dieses Produkt angeboten wird.

Oft verwendet man in der Formulierung bei Werbeanzeigen Beschreibungen, die durch nichts bestätigt sind: „am billigsten“, „am schönsten“, „am schnellsten“ usw.

5 Fraktale in der Natur

Es hat sich in der Evolution von biologischen Systemen gezeigt, dass fraktale Strukturen optimal sind:

- a) Es gelingt, innerhalb einer möglichst kleinen Fläche eine möglichst lange Struktur unterzubringen.
- b) Es gelingt, innerhalb eines möglichst kleinen Volumens eine möglichst große Fläche unterzubringen.

5.1 Fraktale in der unbelebten Natur

Unregelmäßige Oberflächen gibt es in der realen Welt in Hülle und Fülle, vom Atom über das Gebirge bis zum Weltraum. Man geht meistens so vor:

- Beschreibung ihrer Geometrie,
- Erklärung ihrer Existenz,
- Verständnis ihrer Eigenschaften.

Die in Abb. 5.1 dargestellten Berge sind ein typisches Beispiel für fraktale Körper. Sie sind möglicherweise die Reste eines Vulkans, an dem der Zahn der Zeit genagt hat; durch Erosion, Wind, Wasser und unterschiedliche Temperaturen sind viele kleine Stücke abgesplittert, die härtesten Teile sind noch da.

Zum Standard-Beispiel für Fraktale sind Küstenlinien geworden.

Durch die Wellen des Ozeans werden kleinere, weichere Felsstücke einer Küste als erstes ausgewaschen, wodurch immer mehr Unregelmäßigkeiten entstehen. Diesen Vorgang nennt man Erosion. Es bilden sich kleine Einkerbungen, die dann durch die Kraft des Wassers immer größer werden. Auf diese Weise verlängert sich die Küste



Abbildung 5.1: Fraktale Berge

allmählich, und die Wucht der Wellen verteilt sich auf eine längere Küstenlinie. Fraktale Küsten sind deshalb weniger der Erosion ausgesetzt. Geographen haben diesen Entwicklungsverlauf der Küstenmorphologie modelliert. In ihrem Modell ist die Stärke der Wellen umgekehrt proportional zur Länge der Küste. Die Wissenschaftler fanden heraus, dass dieser Prozess einen stabilen Punkt erreicht, wenn die Küste bei einer fraktalen Dimension von genau $4/3 = 1,333$ angelangt ist.

Die Bestimmung der Länge der Küstenlinie von England war der Ausgangspunkt der Arbeiten von Mandelbrot. Auf einer Satellitenaufnahme kann man die fraktale Struktur sehr schön erkennen.

Für die weiteren Untersuchungen ist eine Beobachtung von Robert Brown zu besprechen. Er entdeckte 1827, dass kleinste Teilchen (Pollenkörner), die in einer Flüssigkeit schweben, ungeordnete Zickzackbewegungen vollführen. Diese Bewegung wird daher auch Brownsche Bewegung genannt. Wenn man die Bewegung eines Teilchens im Raum aufzeichnet und nur eine Dimension betrachtet, so erhält man eine Funktion über die Zeit. Dabei stellt die horizontale Achse die Zeit, die vertikale den Zustand der Bewegung dar. Die Analyse dieser Funktion ergab, dass die Brownsche Bewegung statistisch selbstähnlich ist und eine fraktale Dimension von 1,5 aufweist.

B. Mandelbrot und J. van Ness verallgemeinerten die Brownsche Molekular-Bewegung, in dem sie die fraktale Brownsche Bewegung definierten. Damit charakterisierten sie eine Familie von stochastischen Prozessen, die auf der Normalverteilung basieren. Es war nun möglich, die 1,5-dimensionale Brownsche Bewegung mit jeder

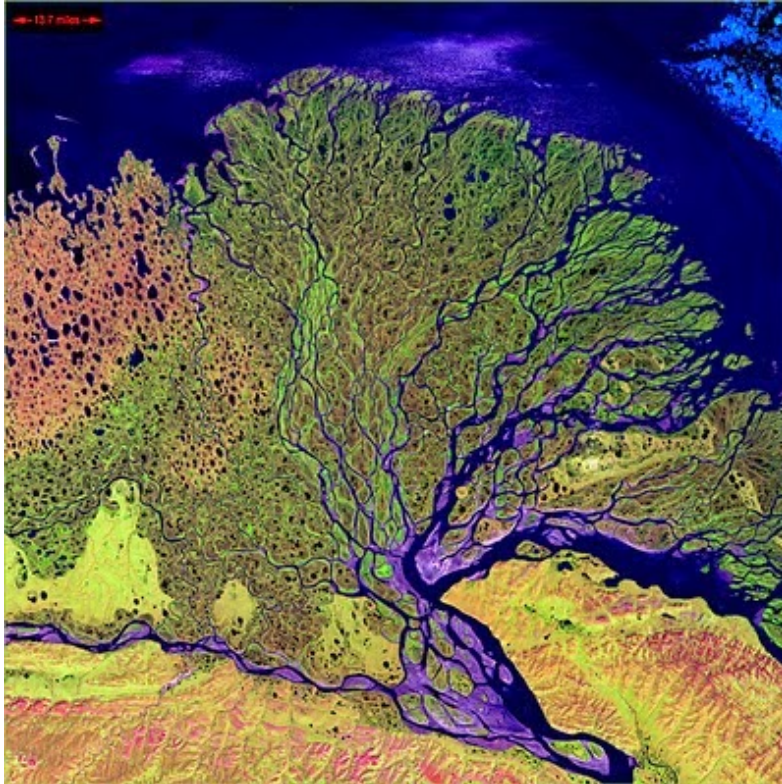


Abbildung 5.2: Das Delta der Lena

beliebigen Dimension zwischen 1 und 2 zu simulieren. Bei diesen Untersuchungen entdeckte B. Mandelbrot die Ähnlichkeit zwischen der generierten Funktion mit fraktaler Dimension 1,2 und der Skyline eines Berges. Es lag somit nahe, dass man die Brownsche Bewegung um eine zweite Dimension erweitert, um Terrains zu generieren.

Zu Beginn des 20. Jahrhunderts (1905) war es Albert Einstein, der diese Molekularbewegung mathematisch beschrieb. Schließlich gelang es Norbert Wiener 1923, die wahrscheinlichkeitstheoretische Existenz dieses von Einstein modellierten Prozesses nachzuweisen. Diese Kritik rief dann B. Mandelbrot und John W. van Ness auf den Plan. Mit ihrer bahnbrechenden Arbeit „Fractional Brownian motions, fractional noises and applications, SIAM Rev. 10 (1968), 422-437“ gelang es den beiden, den Wiener-Prozess zur fraktalen Brownschen Bewegung zu verallgemeinern.

Entsprechend den Erklärungen von Einstein geht das ungeordnete Gewimmel von Teilchen in Flüssigkeiten auf die Wärmebewegung der Flüssigkeitsmoleküle zurück, die die Teilchen anstoßen. Mit dieser Erklärung half Einstein, den Grundstein für die Erforschung verschiedener Gruppenprozesse zu legen.

Heute nutzen Wissenschaftler der unterschiedlichsten Disziplinen, von der Biologie bis zur Soziologie, diese Erkenntnisse. Ganz praktisch lässt sich damit unter an-

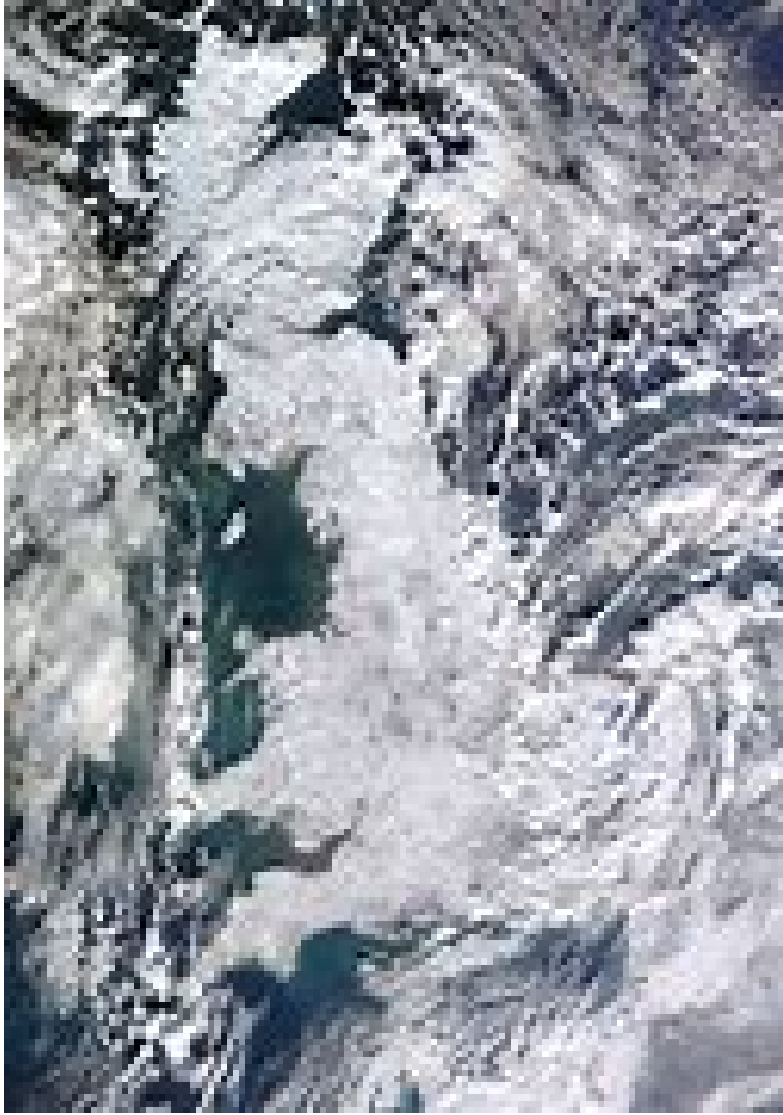


Abbildung 5.3: Die Küste von England

derem die Futtersuche von Ameisen beschreiben: Meist folgen Ameisen den Duftspuren, die ihre Kollegen hinterlassen haben“. Die Wege der Insekten richten sich nach diesen chemischen Markierungen und weisen ihnen den Weg nach Hause oder zu einer Nahrungsquelle. Damit sie neue Futterquellen finden, ist ein zufälliges Abweichen von diesen Geruchsmarken wichtig. Diese Störungen lassen sich durch einen Zufallsfaktor beschreiben, der schon bei Einstein als ungerichteter Einfluss der Wärmebewegung der Moleküle auftaucht. Dieser Zufallsfaktor, den der Franzose Paul Langevin drei Jahre nach Einstein allgemeiner formulierte, dient auch in heutigen Modellen als Variable für eine ungerichtete, scheinbar zufällige Kraft, die verschiedene chaotische Einflüsse repräsentiert.

In [2] werden viele weitere Begriffe definiert und erklärt, die mit der Brownschen Bewegung zusammenhängen. Damit haben wir ein Konzept, dass sowohl für die Physik als auch für die Theorie der Fraktale grundlegende Bedeutung besitzt.



Abbildung 5.4: Robert Brown (1773 - 1858)

5.2 Fraktale in der Pflanzenwelt

Man sieht bei allen pflanzlichen Strukturen, dass es wieder darauf ankommt, eine große Oberfläche auf geringem Raum unterzubringen (Abb. 5.5 und 5.6).



Abbildung 5.5: Pilze an einem Baumstumpf

Jede Scheibe der Pilzkolonie ist für sich ein zweidimensionales Fraktal, alle Scheiben zusammen bilden ein dreidimensionales Fraktal.



Abbildung 5.6: Ein gut entwickelter Kaktus

5.3 Fraktale in der Tierwelt

Warum gibt es eigentlich sowohl Mäuse als auch Elefanten? Warum können Mücken und Libellen, Sardinen und Tunfische, Kolibris und Kondore in jeweils derselben Umwelt leben?

Diese Fragen sind keineswegs so abwegig, wie sie klingen. Sie ergeben sich vielmehr aus einigen nahe liegenden Überlegungen. Zum einen sind die elementaren Bestandteile großer wie kleiner Tiere, die Körperzellen, einander erstaunlich ähnlich. Das gilt insbesondere dann, wenn – wie etwa bei Landsäugetern – der Grundbauplan der Tiere der gleiche ist. Ein Elefant hat nicht größere Leberzellen als eine Maus, sondern nur viel mehr. Gemessen an den Größenunterschieden zwischen den Tieren verschwinden die Unterschiede zwischen ihren Leber- (Lungen-, Nieren-) Zellen. Insbesondere zeigen isolierte Zellen im Reagenzglas in etwa dieselbe Stoffwechselaktivität, einerlei welcher Tierart sie entstammen. Die Menge an Wärme, die eine bestimmte Zelle pro Zeiteinheit produziert, ist für Mäuse und Elefanten ungefähr die gleiche. Ähnliches gilt für den Sauerstoffverbrauch oder – bei einer Nierenzelle – den Flüssigkeitsdurchsatz. [33]

Demnach müsste die stoffwechselbedingte Wärmeerzeugung eines Tieres der Anzahl seiner Zellen oder, was auf dasselbe hinausläuft, seinem Körpervolumen proportional sein. Die Wärmeabgabe dagegen ist – abgesehen von Details wie Fell, Fettschicht, Schweißdrüsen oder Gefieder – im Wesentlichen proportional der Oberfläche. Beide Größen müssen aber miteinander im Gleichgewicht stehen, weil sonst die Körpertemperatur unaufhaltsam in unphysiologische Bereiche ansteigen oder abfallen würde

Wenn ein Tier zehnmal so lang, breit und hoch ist wie ein anderes, hat es das tausendfache Volumen, aber nur die hundertfache Oberfläche. Sein Verhältnis von

Volumen zu Oberfläche ist zehnmal so groß, und damit auch – unter ansonsten gleichen Umständen – das Verhältnis von Wärmeerzeugung zu Wärmeabgabe. Gleichgewicht könnte also nur bei einer bestimmten optimalen Körpergröße herrschen. Ein in diesem Sinne zu großes Tier müsste zusätzlichen Aufwand für die Wärmeabführung treiben, ein zu kleines müsste ständig nachheizen, um nicht auszukühlen.

Was hier am Beispiel des Wärmehaushalts dargestellt wurde, gilt für zahlreiche weitere Stoffwechselaktivitäten wie Atmung, Verdauung und Urinausscheidung: Ihre Intensität müsste eigentlich proportional dem Volumen des zugehörigen Organs sein, aber sie finden alle an Oberflächen statt.

Solche Oberflächen sind beispielsweise Lungenbläschen, Darmzotten, Grenzflächen zwischen Blut und Zellen oder – für den Stoffaustausch zwischen den kompletten Blutkreisläufen von Mutter und Kind – die Plazenta. Da Oberflächen mit zunehmender Körpergröße langsamer anwachsen als Volumina, müsste eine Elefantenlunge weitaus schlechter arbeiten als eine Mäuselunge.

Man weiß aber, dass alle heute lebenden Tiere Ergebnisse einer sehr langen darwinistischen Evolution sind. In einem Jahrmillionen andauernden Ausleseprozess haben sich unter vielen Tierarten stets diejenigen durchgesetzt, die von den verfügbaren Ressourcen den optimalen Gebrauch machen. Ein zusätzliches Heiz- oder Kühlaggregat, welcher Art auch immer, würde so viel Aufwand erfordern, dass der dadurch entstehende Nachteil irgendwelche Vorteile einer abweichenden Körpergröße weit übertroffen hätte. Demnach dürfte es, zumindest bei gleichem Körperbauplan und gleichen Umweltbedingungen, nur sehr geringe Abweichungen in der Körpergröße geben.

Offensichtlich ist etwas falsch an diesen Überlegungen. Große und kleine Tiere existieren nicht nur, sie haben entgegen dieser Theorie auch ihre Wärmeregulation gut im Griff; und das lässt sich nicht durch grundsätzlich unterschiedliche Baupläne erklären. Eine Maus erfriert nicht, solange sie genug zu fressen hat, und heißlaufende Elefanten kommen in der Natur auch nicht vor.

Erst in jüngerer Zeit (1985) wurde eine funktionell begründete Deutung gefunden, und zwar mit Hilfe zweier neuer Konzepte. Man betrachtet den Organismus, was seine Stoffwechselaktivität angeht, einerseits als einen Bioreaktor unter denselben Gesichtspunkten, wie sie für industrielle Anlagen maßgebend sind, andererseits von einem mathematischen Standpunkt als eine fraktale Struktur.

Zunächst zum technischen Standpunkt. Ein industrieller Bioreaktor besteht typischerweise aus einem großen Flüssigkeitsbottich, dem kontinuierlich die umzusetzenden Stoffe (die Reaktanden) zugeführt werden; an anderer Stelle werden die

Reaktionsprodukte abgezapft. Der Reaktor enthält entweder einen Katalysator (in der Regel ein Enzym), der an Partikeln aus einem porösen Gel oder auch eine poröse Membran gebunden ist, oder, in der Lösung suspendiert, lebende Zellen. In beiden Fällen sind also Reaktanden und Katalysator nicht einfach gelöst und damit beliebig beweglich in der Flüssigkeit verteilt; vielmehr findet die chemische Reaktion nur an der Oberfläche und im Inneren des katalytisch aktiven Teils – Gel oder Membran beziehungsweise Zellen – statt. Man bezeichnet das als heterogene im Gegensatz zur homogenen Katalyse. Auf die Dauer stellt sich im Reaktor ein Fließgleichgewicht der an der Reaktion beteiligten Substanzen ein.

Damit der Reaktor effizient, das heißt mit hohem Stoffumsatz pro Volumen, arbeitet, müssen den aktiven Oberflächen ständig frische Reaktanden zugeführt werden. Dem stehen zweierlei Widerstände entgegen: An den aktiven Oberflächen liegt eine äußere Grenzschicht an, die weniger beweglich ist als der Rest der Flüssigkeit; und die Reaktanden müssen durch Diffusion ins Innere der porösen Membran beziehungsweise der Zelle vordringen. Um das zweite Hindernis klein zu halten, verwendet man möglichst dünne Membranen und kleine Partikeln oder Zellen; um die Grenzschicht aufzumischen, hält man die Flüssigkeit mit einem Rührwerk in intensiver, turbulenter Bewegung (CSTR - continuously stirred tank reactor, ständig gerührter Bottichreaktor). Der Vorteil der turbulenten gegenüber der laminaren Strömung ist die gründliche Durchmischung der Flüssigkeit; er wird allerdings mit einem hohen Energieaufwand für das Rühren erkauft.

Man sollte annehmen, dass der Stoffumsatz eines solchen Reaktors seinem Volumen proportional ist: In einem doppelt so großen Bottich müssten pro Zeiteinheit – unter ansonsten gleichen Umständen – doppelt so viele chemische Reaktionen stattfinden. Es sind jedoch im Allgemeinen deutlich weniger als doppelt so viele.

Wie bei der Allometrie in der Biologie beschreibt man diese empirisch gefundene Abhängigkeit durch eine Potenzfunktion V^b mit nicht-ganzzahligem Exponenten b . Der Zahlenwert von b hängt dabei stark von der Zell- oder Partikeldichte im Reaktor ab. Bei niedrigen Zelldichten und hinreichend turbulenter Mischung ist der Gesamtumsatz des Reaktors einfach die Summe der Umsätze der individuellen Zellen in der Suspension. In diesem Fall ist der Umsatz in der Tat dem Volumen proportional ($b=1$): Das System verhält sich isometrisch. Bei hohen Zelldichten aber macht sich bemerkbar, dass die Effizienz des Rührwerks mit der Größe des Reaktors abnimmt, und der spezifische Umsatz sinkt allometrisch mit der Größe des Systems ($b<1$).

Aus der Perspektive des Technikers ist auch ein lebender Organismus ein Bioreaktor mit heterogener Katalyse. Die flüssige, getriebene Phase ist zum Beispiel das Blut, die stationäre Phase, an deren Grenze und in deren Innerem die Reaktionen statt-

finden, das Körpergewebe. Wie im technischen Reaktor wird durch Reaktionen und Transport ein Fließgleichgewicht der chemischen Bestandteile aufrechterhalten.

Ein entscheidender Unterschied liegt jedoch im Mischungsprozess. Im Organismus gibt es kein Rührwerk, sondern nur das Herz, welches das Blut in Bewegung hält. Diese Bewegung ist fast ausnahmslos laminar und erfordert damit wesentlich weniger Energie als die turbulente. Gleichwohl findet eine äußerst effektive Durchmischung zwischen der flüssigen (Blut) und der stationären Phase (Gewebe) statt. Eine in den menschlichen Kreislauf intravenös injizierte Substanz ist bereits nach etwa zwei Minuten im ganzen Körper gleichverteilt. Dieses Ergebnis wäre bei einem gerührten technischen Reaktor gleichen Volumens nur unter sehr hohem Aufwand zu erreichen.

Dieselben Argumente wie bei den technischen Reaktoren gelten auch für die Skalenabhängigkeit des Umsatzes: Da die Reaktanden in einem vergrößerten Organismus längere Transportwege zurücklegen müssen – und wegen der Viskosität des Blutes dessen Fließgeschwindigkeit nicht in gleichem Maße ansteigen kann –, sinkt mit zunehmender Körpergröße die spezifische Umsatzrate. Das wäre bereits der Ansatz einer Erklärung für die Allometrie. Allerdings funktionieren die natürlichen Reaktoren um mehrere Größenordnungen besser als die technischen: bei einem Bioreaktor mit 109 Zellen pro Liter fällt die spezifische Umsatzrate mit wachsendem Volumen entsprechend einem allometrischen Exponenten $b-1$ von ungefähr $-0,3$, bei Organismen dagegen nur mit $-0,25$, und das bei der tausendfach höheren Dichte von 1012 Zellen pro Liter!

Erstaunlich ist also nicht in erster Linie, dass der spezifische Umsatz mit wachsender Körpergröße abfällt, sondern dass dieser Abfall so gering ist. Offensichtlich ist der Mischprozess in lebenden Organismen um ein Vielfaches effizienter als schlichtes Umrühren. Das wiederum ist in der fraktalen Struktur des Gefäßsystems und des Körpergewebes überhaupt begründet.

Die Stoffwechselintensität eines Tieres steigt eben nicht proportional zu seinem Volumen, also zur dritten Potenz der Körperlänge. Sie steigt auch nicht proportional zur Oberfläche, was dem Quadrat der Körperlänge entspräche oder $V \cdot (2/3)$, wobei V für das Volumen steht. Vielmehr liegt die Anstiegsrate dazwischen. Es gilt eine Gesetzmäßigkeit der Form V^b , wobei der Exponent b , ermittelt aus Messungen an einer Vielzahl von Tierarten, ungefähr den Wert $0,74$ hat. In der Körperlänge L statt im Volumen V ausgedrückt, skaliert die Stoffwechselrate wie $L^{3b} = L^{2,22}$. Insbesondere sinkt mit zunehmender Körpergröße die Stoffwechselrate pro Volumeneinheit: Diese so genannte spezifische Stoffwechselrate ist annähernd proportional zu V^{b-1} .
[34]

Das Verhältnis von „Oberfläche“ zu „Volumen“ ist also vom Standpunkt der fraktalen Analyse konstant. Damit ist auch nicht mehr verwunderlich, dass es große und kleine Tiere gibt. Mäuse und Elefanten sind in ihrem Stoffwechsel gleich effizient und können daher ohne weiteres koexistieren. Das eingangs angeführte Paradox der allometrischen Stoffwechselreduktion kommt nur dadurch zu Stande, dass wir Leistungen und Funktionen auf das falsche, weil unstrukturierte Volumen- beziehungsweise Oberflächenmaß bezogen haben.

Ein Fraktal ist, wie jedes mathematische Konzept, in der Natur nicht mit absoluter Vollkommenheit, sondern nur näherungsweise realisiert. So gilt insbesondere die Selbstähnlichkeit nur innerhalb bestimmter Maßstabsbereiche mit unteren und oberen Grenzen. Obere Grenze ist in unserem Beispiel die Größe des gesamten Organs, untere die Größe einer einzelnen Zelle. Bei Annäherung an die Grenzen verschwinden die fraktalen Eigenschaften oder gehen eventuell in ein anderes Skalenverhalten über.

Ein Organismus, als Ganzes betrachtet, ist ein äußerst inhomogenes Fraktal. So gilt die Dimension $D=2,22$, wie wir sie für die Arterien gemessen haben, lediglich für den transportlimitierenden Anteil des Gefäßsystems und davon abhängige Größen wie das Skalieren der Stoffwechselraten für den Gesamtorganismus. In anderen Bereichen, etwa den Kapillaren des Gefäßsystems selbst oder den Alveolen der Lungen, kann die fraktale Dimension wesentlich höhere Werte erreichen, bis zu $D=3$, womit sich diese Gebiete isometrisch verhalten (das heißt mit dem Volumen skalieren).

In der Tierwelt findet man viele schöne Fraktale. Unter Wasser sind sicher die Korallen am schönsten (Abb. 5.7).



Abbildung 5.7: Korallen

Wissenschaftler untersuchten männliche und weibliche Rebhühner (*Alectoris Rufa*), die beide komplizierte schwarz-weiße Gefiedermuster auf der Brust aufweisen.[19]

„Wir haben gezeigt, dass die fraktale Geometrie biologisch bedeutsame Informationen enthüllen kann, die in einem komplexen Gefiedermerkmal kodiert sind: dem schwarz gefleckten Lätzchen des Rebhuhns mit rotem Bein“, schreiben die Forscher in einem Artikel, der am 23. Januar in der Zeitschrift „Proceedings“ der Royal Society veröffentlicht wurde. „Unsere korrelativen Ergebnisse zeigen, dass sowohl bessere Bedingungen als auch eine bessere Immunreaktion aus Lätzchen mit höherer FD vorhergesagt werden können.“

Die Wissenschaftler untersuchten auch die Beziehung zwischen Vogelgesundheit und Fraktalen, indem sie die Ernährung von 33 männlichen und weiblichen Rebhühnern während ihrer Häutungszeit einschränkten, während 35 andere so viel essen konnten, wie sie wollten. Dies hatte zur Folge, dass der erste Vogelsatz etwa 13 Prozent weniger wiegt als die Kontrollgruppe. Die Forscher fotografierten dann die Lätzchen der Rebhühner und stellten fest, dass die Komplexität der Fraktale bei den Vögeln, deren Futter eingeschränkt wurde, signifikant reduziert war. Nach dem Abnehmen wuchsen dieselben Vögel im Gefieder mit einer geringeren fraktalen Dimension als zuvor, während die Vögel, deren Gewicht konstant blieb, keine Veränderung der fraktalen Dimension zeigten. Insgesamt ergab die Studie, dass die Fraktale von Vögeln viel über die Gesundheit von Individuen aussagen.



Abbildung 5.8: Die Tarnung wird oft in die fraktale Struktur einbezogen

6 Fraktale Kunst

Nach den vielen interessanten und schönen Bildern ist es naheliegend, sich etwas mehr mit den Beziehungen zwischen fraktalen Strukturen und verschiedenen Formen der Kunst zu beschäftigen.

6.1 Fraktale Musik

Der intellektuelle Aspekt der Musik, schlägt sich z.B. in imitatorischen Techniken (Kanon, Fuge) und motivischer Arbeit nieder. Dieser mathematische und logische Aspekt der Musik legt heute den Gebrauch von Computern nahe [10]. Ein Hauptgebiet der intellektuellen Beschäftigung des Autors **Arturo Raffaele Grolimund** mit Musik ist die Beziehung von Chaostheorie bzw. Fraktalen zur Musik. In seiner Abhandlung „Fraktale der Musik“ von 1994 untersuchte er fraktale Strukturen, die in der Musik selbst vorkommen. Er zeigt sehr bemerkenswert, dass alle wichtigen Eigenschaften von Fraktalen auch in der Musik auftreten können.

- Das Konzept der Selbstähnlichkeit meint die Tatsache, dass eine Figur sich selbst in kleinerem Maßstab immer wieder enthält. Das kann Ausgangspunkt zur Erzeugung musikalischer Strukturen sein.
- Hier findet sich auch die für Fraktale typische sensitive Abhängigkeit von den Anfangsbedingungen; wie sich der Kanon bzw. das gezeichnete Fraktal entwickelt, hängt wesentlich von der Wahl der ersten Töne ab.
- Das Stück **Colours of Afrika** ist inspiriert von afrikanischer Mbira-Musik. Das Stück enthält viele fraktale Kanons: Bei gleicher Anzahl von Tönen lassen sich doppelt so viele Kanonstimmen unterbringen. Hört man in die Klänge hinein, sind Wellenbewegungen hörbar.
- Das Paradebeispiel eines musikalischen Fraktals ist die Obertonreihe, die physikalische Grundlage des Dur–Dreiklangs und damit des traditionellen musikalischen Harmonieempfindens überhaupt. Auch sie ist selbstähnlich, d.h. sie

enthält sich auf regelmäßige Weise unendlich viele Male selbst. Diese musikalischen Fraktale bzw. ihre Grafiken sind nicht so komplex wie die Mandelbrotmenge oder die Juliamenge. Das kommt wohl daher, dass jeder Tonpunkt ja schon die Obertonreihe als Klangfarbe beinhaltet und somit selber schon ein einfaches Fraktal ist.

- Ebenso lässt sich bei fraktalen polyphonen Strukturen die kontrapunktische Dichte kaum noch erhöhen.

Die alte Forderung nach der Natürlichkeit der Kunst erhält so eine zusätzliche Bedeutung. Die Natur und ihre Prozesse dienen als Vorbild innovativer künstlerischer Gestaltung. Das Prinzip, in einem kleinen Volumen ein Maximum einer bestimmten Substanz unterzubringen, wurde schon oft erwähnt. Das zeigt sich auch hier.

Offensichtlich war Anfang der 90er Jahre fraktale Musik sehr populär. Das zeigt auch ein Artikel [11] einer Gruppe von komponierenden Programmierern. Ein Random-Algorithmus variiert die Obertöne eines obertonhaltigen Klanges, sodass alle harmonischen Proportionen vermieden werden. Diese Random-Variation im Mikrobereich geht in eine ebensolche im Makrobereich über, wird also hörbar und expandiert schließlich soweit, dass der gesamte 128-tonige Umfang der General-Midi-Sound-Moduln ergriffen wird.

- Choral: Harmoniefolgen werden algorithmisch errechnet. Gute Ausbildung im vierstimmigen Satz ist beim Programmierer vorhanden, aber nachweislich ohne Bedeutung. Das Schlagzeug und der Bass sekundieren mit gezielt gesetzten Zufallsstrukturen.
- Einsames Motiv zwischen Dur und Moll: Ein einsames Motiv wird algorithmisch gemorpht. Aufmunternd sekundiert ein Programm, das Dur- und Moll-Dreiklänge erzeugt, wobei Tasteneingaben die Ausgabe in gewissen Grenzen halten. Eine Trommel stabilisiert das Geschehen, das ein rhythmisch aufgelockerter Ganzton-Basslauf beunruhigt.
- Das etwas freundlichere Livespiel: Der verwendete Algorithmus lässt vereinigte Zellen am Leben. Die Musikalisierung des Live-Prozesses ist abendländischen Partitur-Lesegewohnheiten nachempfunden.
- Entwickelnde Variation: Arnold Schönberg hat mit seiner Idee der entwickelnden Variation den Prozess der rekursiven Definition fraktaler Vorgänge kompositorisch vorweggenommen. Das vorliegende Programm greift diese Idee auf und entwickelt Variationen bekannter Motive nach der Wachstumsformel

von Verhulst.

- Fractal Workstation '94: Das Programm ist ein für den „Tag der Offenen Ohren“ des Faches Musik 1994 entwickeltes Programm. Beruhend auf dem Höpferalgorithmus, der 1986 in „Spektrum der Wissenschaft“ als grafisches Wunder vorgestellt wurde, werden verblüffend unterschiedliche Bildentstehungsprozesse instrumentiert.
- Die Hasen auf Wangeroog: Der bekannteste fraktale Algorithmus hat die Formel $X(n+1) = X(n)^2 - C$. Bei verändertem C führt er zu überraschenden Bifurkationen, die sich musikalisch in Motivverzweigungen äußern. Im vorliegenden Orchesterstück spielen acht Ataris dasselbe Kernprogramm in acht unterschiedlichen rhythmischen Mustern (die der Dirigenten-Atari vorgibt). Das Resultat ist eine Abschiedssymphonie auf die Hasen von Wangeroog.

Es gibt noch viele weitere Aktivitäten, die man im Internet finden kann:

- Der Komponist Dmitry Kormann aus Brasilien transkribierte visuelle Muster in Sounds mit Hilfe von Synthesizern und Programmen.
- Ein weiterer wichtiger Komponist war der Spanier Francisco Guerrero Marín.
- Phil Thompson und sein **Organisiertes Chaos**. Er entwickelte das Lebkuchenprogramm, mit denen fraktale Kompositionen erzeugt werden.

Das Schlagzeug ist möglicherweise das Instrument, das für Fraktale besonders geeignet ist (Abb. 6.1) Selbst der beste Schlagzeuger produziert beim Spielen winzige Schwankungen in Rhythmus und Lautstärke. Doch diese sind keineswegs zufällig. Stattdessen bilden sie Fraktale: die Abweichungen sind auf gleich mehreren Ebenen selbstähnlich. Solche Muster absichtlich zu erzeugen, ist nahezu unmöglich, aber unabsichtlich scheinen sie unsere Musik sogar entscheidend zu prägen.

Eine sehr interessante Arbeit kann man unter [4] finden. Dort wird der Lorenz - Attraktor von Abb. 6.2 in Musik überführt. Der Lorenz-Attraktor ist der seltsame Attraktor eines Systems von drei gekoppelten, nichtlinearen gewöhnlichen Differentialgleichungen:

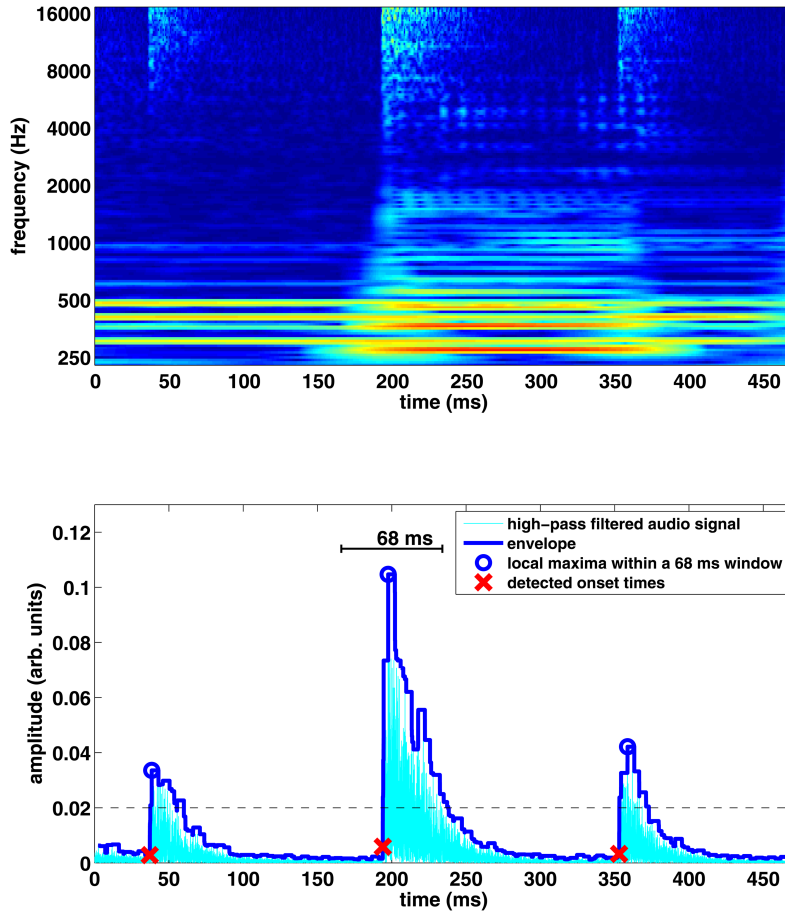


Abbildung 6.1: Die Begleitung eines Volksliedes durch eine Trommel

$$\frac{dx}{dt} = \sigma \cdot (x + y) \quad (6.1)$$

$$\frac{dy}{dt} = R \cdot x - y - x \cdot y \quad (6.2)$$

$$\frac{dz}{dt} = -B \cdot z + x \cdot y. \quad (6.3)$$

Dabei sind B , R und σ Konstanten. Am Anfang befindet man sich in einem bestimmten Punkt. Die Gleichungen 6.1, 6.2 und 6.3 geben für jeden Punkt (x, y, z) eine Geschwindigkeit vor, und nach einer bestimmten Zeit (eine Zehntelsekunde) befindet man sich in einem neuen Punkt, in dem man sich in gleicher Weise verhält. So kommt insgesamt eine dreidimensionale Kurve $x(t), y(t), z(t)$ zustande, deren Eigenschaften noch durch die Konstanten B , σ und R beeinflusst werden können.

Die sich ergebenden Linien kommen sich beliebig nahe, ohne sich zu berühren. Der Lorenz-Attraktor hat eine fraktale Dimension von $D \approx 2,06$. Die für (x, y, z)



Abbildung 6.2: Das Lorenz - Fraktal

berechneten Werte des Lorenz-Kontraktors werden auf eine elegante Art und Weise in Töne übersetzt, die nach einer gewissen Harmonisierung und einer Beschränkung auf Dur eine entsprechende fraktale Tonfolge ergeben. Es wäre sicherlich interessant, auch andere Fraktale in Musik zu übersetzen.

Man könnte diese Art von Musik auch als chaotische Musik bezeichnen. Die numerische Lösung des Systems zeigt bei bestimmten Parameterwerten deterministisch chaotisches Verhalten, die Trajektorien folgen einem seltsamen Attraktor. Damit spielt der Lorenzattraktor für die mathematische Chaostheorie eine Rolle, denn die Gleichungen stellen wohl eines der einfachsten Systeme mit chaotischem Verhalten dar.

Bei youtube kann man sich sehr viele Videos mit fraktaler Musik ansehen und anhören. Viele Stücke sind sehr schön, manche aber auch extrem atonal, je nachdem welches Fraktal verwendet wird. Die Bilder des sich entwickelnden Fraktals sind aber immer außerordentlich eindrucksvoll.

Eine Gitarrenseite schwingt und beschreibt eine halbe Sinus-Welle (Abb. 6.3). Erzeugen wir in ihrer Mitte einen künstlichen Knoten indem wir die Saite an dieser Stelle sanft am Schwingen hindern, verdoppelt sie die Frequenz und klingt eine Oktave höher. Wir können die Saite auch dreimal so schnell schwingen lassen (noch einmal eine Quinte höher) oder viermal, und so weiter.

Interessant sind die folgenden Fragen:

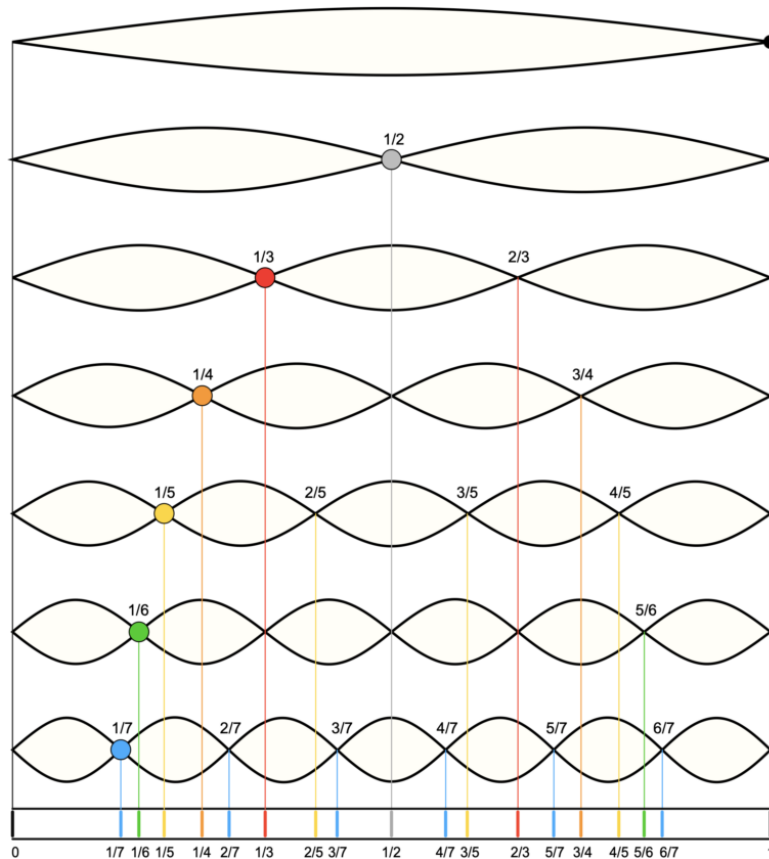


Abbildung 6.3: Schwingungen von Musikinstrumenten

- Welche Obertöne sind vorhanden?
- Wie laut sind diese Obertöne im Verhältnis zueinander?
- Wie ändern sich die Lautstärke und Frequenz der einzelnen Obertöne, während der Ton erklingt?
- Welche Nebengeräusche (Anschlageräusche, Blasgeräusche usw.) kommen hinzu?

Folgende Instrumente haben einen besonders charakteristischen Teiltonaufbau:

- Streichinstrumente besitzen ein sehr reichhaltiges Teiltonspektrum.
- Klarinetten betonen die Lautstärke der ungeraden Teiltöne.
- Beim Fagott ist der Grundton sehr viel schwächer als die ersten Obertöne.
- Glocken betonen oftmals die Terzen sehr stark, und die Obertonzusammensetzung ist komplex.

- Stimmgabeln erzeugen fast nur den Grundton.

Jeder Ton ist Klang, und jeder Klang ist ein Fraktal. Und selbst inmitten eines ganzen Orchesters erkennt man die Oboe an ihrer spezifischen Obertonreihe, ohne deren Aufbau je gelernt zu haben. Auch ihre Melodie ist ein In- und Nacheinander von Fraktalen.

Für alle, die sich nach einer abgespeckten Version von Bach sehnen - einer Version, bei der drei Viertel der Noten des Meisters weggeschnitten wurden - ist Hilfe auf dem Weg. Ein Wissenschaftler und sein Sohn, ein Musiker, haben ein System entwickelt, mit dem sich die Musik auf ihre eigentliche Essenz reduzieren lässt.

In einer in den „Proceedings of the National Academy of Sciences“ veröffentlichten Arbeit schlagen Dr. Kenneth J. Hsu, Geologieprofessor an der Eidgenössischen Technischen Hochschule in Zürich, und sein Sohn Andrew vor, dass mathematische Auszüge aus der Musik von J. S. Bach als Matrizen dienen könnten, auf deren Grundlage neue Kompositionen Bach-ähnlicher Musik erstellt werden könnten, die in ihrer Qualität mit denen des Komponisten selbst vergleichbar sind.

Das System basiert auf der fraktalen Geometrie, dem Studium von Strukturen, bei denen die Form jedes kleinen Teils die Form der gesamten Struktur widerspiegelt. Die Hsus sind der Meinung, dass zumindest einige Musikstücke eine fraktale Form aufweisen, da die grundlegenden musikalischen Muster bestehen bleiben, auch wenn Noten hinzugefügt oder entfernt werden. J. S. Bachs Invention Nr. 1 in C-Dur beispielsweise weist eindeutig eine fraktale Struktur auf, so die Autoren. Man findet dieses kleine Stück ohne Probleme bei youtube.

In der Einleitung ihrer Arbeit erinnern die Autoren daran, dass der österreichische Kaiser Josef II., als er zum ersten Mal die Oper „Die Entführung aus dem Serail“ hörte, zu Mozart sagte, die Musik sei himmlisch, enthalte aber zu viele Noten. (Mozart soll entrüstet geantwortet haben, dass keine einzige Note eingespart werden könne.) Dr. Hsu räumt ein, dass „Music Digests“, die für den österreichischen Kaiser und andere Spießbürger geeignet sind und aus „Halb-Mozarts, Drittel-Chopins oder Viertel-Bachs“ bestehen, vielleicht nie populär werden. Aber die mathematische Reduktion von Kompositionen der Meister auf eine Skelettform könnte Komponisten helfen, wichtige neue Werke zu schaffen, glaubt er.[21]

6.2 Fraktale Bilder

Die fraktale Malerei ist eine relativ weit verbreitete Unterform der digitalen Kunst und beschäftigt sich mit der Erschaffung von digitalen Bildern, die im Wesentlichen aus einem oder mehreren Fraktalen bestehen (Abb. 6.4). Als eigenständige Kunstform ist die Fraktalkunst selbst unter den Künstlern umstritten, da einige argumentieren, man bilde lediglich Teile einer mathematisch-geometrischen Form als Kunstwerk ab, was nicht die eigenständige Schöpfungshöhe erreicht. Es gibt diverse Programme, die zum Erschaffen von Fraktalkunst verwendet werden, zu ihnen gehören die Freeware-Programme *Apophysis* bzw. *Apo7x*, der *Fractal Explorer* sowie *Ultra-Fractal*.



Abbildung 6.4: Fraktales Kunstwerk von Sebastian Baumer

Ornamente und keramische Erzeugnisse können ebenfalls aus fraktalen Strukturen abgeleitet werden (Abb. 6.4 , Abb. 6.5 und Abb. 6.6).

6.3 Fraktale in der Literatur

Vor einiger Zeit haben Physiker fraktale Strukturen auch in den großen Werken der Weltliteratur nachweisen können. Die Forscher um Stanislaw Drózd von der



Abbildung 6.5: Ein fraktaler Teppich



Abbildung 6.6: Fraktale Keramik

Technischen Universität Krakau analysierten die Anzahl der Worte in Sätzen – und das in 113 Büchern.

Sie fütterten ihre Computer mit Werken von so unterschiedlichen Autoren wie James Joyce, Thomas Mann, Honoré de Balzac, William Shakespeare, Virginia Woolf, Umberto Eco, Fjodor Dostojewski, Henryk Sienkiewicz, J. R. R. Tolkien oder Julio Cortázar und suchten jeweils nach selbstähnlichen Längenverhältnissen bei den Sätzen. „Alle getesteten Werke zeigen Selbstähnlichkeit in Bezug auf die Abfolgen und Häufigkeiten ihrer Satzlängen“, berichten die Wissenschaftler in der Fachzeitschrift „Information Sciences“: „Diese Texte sind alle fraktal.“ Die meisten Werke erwiesen sich als **linear fraktal**. Die mit Abstand höchste mathematische Komplexität fanden die Forscher in „Finnegans Wake“ von James Joyce. Bei diesem schwer

verständlichen Werk waren die Ergebnisse der Analysen nicht mehr unterscheidbar von idealen, rein mathematisch konstruierten Multifraktalen. Finnegans Wake gilt als eines der bemerkenswertesten, aber auch der am schwersten verständlichen Werke der Literatur des 20. Jahrhunderts, wozu unter anderem seine ungewöhnliche Sprache beiträgt: Joyce prägt eine eigene Sprache, indem er englische Wörter neu zusammenfügt, umbaut, trennt, oder auch mit Wörtern aus Dutzenden anderen Sprachen mischt (Portmanteaux). Das Ergebnis entzieht sich einem linearen Verständnis und eröffnet Möglichkeiten zu vielfacher Interpretation.

Dem Leser enthüllen sich bei mehrmaligem Lesen immer neue Bedeutungen. Roland McHugh hat mit seinem Standardwerk „Annotations to Finnegans Wake“ knappe Anmerkungen zu Finnegans Wake herausgegeben, wobei seitengleich zu vielen der im Wake verwendeten Wörter Erklärungen in Form etwa geographischer Hinweise oder Hinweise auf Sprachen, denen das jeweilige Wort oder Varianten davon entlehnt sein könnten, angegeben werden.

Doch die polnischen Physiker fanden eine bemerkenswerte Ausnahme: Auch das Alte Testament bestach durch besonders komplexe Multifraktalstrukturen. Es wäre sicherlich eine interessante Frage, ob das Niveau eines Buches durch seine fraktale Dimension ausgedrückt werden kann.

Ein multifraktales System ist eine Verallgemeinerung eines fraktalen Systems, bei dem ein einzelner Exponent (die fraktale Dimension) nicht ausreicht, um seine Dynamik zu beschreiben; stattdessen wird ein kontinuierliches Spektrum von Exponenten (das sogenannte Singularitätsspektrum) benötigt. Das Verhalten um jeden Punkt wird durch ein lokales Potenzgesetz beschrieben:

$$s(\vec{x} + \vec{a}) - \vec{f}(x) \sim a^{h(\vec{x})}. \quad (6.4)$$

Der Exponent wird Singularitätsexponent genannt, da er den lokalen Grad der Singularität oder Regularität um den Punkt beschreibt.

7 Fraktale in Wirtschaft und Technik

Vielen Prozessen in Wirtschaft und Technik ist der fraktale Charakter direkt auf den Leib geschrieben. Produktionszahlen für bestimmte Erzeugnisse, der Bedarf an unterschiedlichen Rohstoffen, Aktien- und Wechselkurse sind durch ein stetiges auf und ab gekennzeichnet.

7.1 Die Fraktale Fabrik

Die Idee der Fraktalen Fabrik, wie sie von Hans-Jürgen Warnecke im Jahr 1992 in die Diskussion gebracht wurde, hat in der Industrie 4.0 und der Smarten Fabrik ihre Entsprechung gefunden. Hauptkennzeichen der Fraktalen Fabrik ist ihre betont dezentrale Organisation, um die Komplexität in den Unternehmen und deren Umwelt beherrschbarer zu machen. Bereits in dem Vorwort zu seinem Buch „Die Fraktale Fabrik. Revolution der Unternehmenskultur“ grenzte Warnecke seinen Ansatz von dem „Computer Integrated Manufacturing (CIM)“ ab, dessen Defizite zu dem Zeitpunkt immer deutlicher zu Tage traten:

Das zugrunde liegende deterministische Weltbild mit bekannten oder bei entsprechendem Forschungsaufwand erkennbaren Zusammenhängen zwischen Ursache und Wirkung ist nicht ausreichend, da es nur für abgegrenzte Teilbereiche der Realität gilt. Demgegenüber empfahl Warnecke, die Realität zu akzeptieren, d.h. die Nichtlinearität der Abläufe in Wirtschaft und Gesellschaft, deren Einfluss sich kein Unternehmen und keine Fabrik entziehen kann, anzuerkennen. Anders lassen sich die Herausforderungen der Zukunft nicht bewältigen. Die bisherigen Formen der Leistungserstellung haben einen sehr hohen Reifegrad erreicht. Sie sind somit in einem Zustand, wo auch mit noch so hohem Aufwand kein größerer Nutzen zu erreichen ist. Als Gestaltungsprinzip, um die sich abzeichnende neue Dynamik unter einen gemeinsamen Nenner bringen zu können, bot sich für Warnecke der zum damaligen Zeitpunkt relativ neue Begriff „Fraktal“ an:

„Fraktale kommunizieren direkt mit entsprechenden Fraktalen der Lieferanten bzw. der Abnehmer. Fraktale können weltweit verteilt sein. Durch Selbstorganisation

wählen sie jeweils die Methoden zur Planung und Steuerung aus und wenden die Automaten und Rechner an, die zum Erfüllen ihrer Aufgabe zweckmäßig sind. Als Gemeinsamkeit aller industriellen Revolutionen kann die Tendenz von der *Zentralisierung* (Dampfmaschine bzw. EDV-Zentrum) zur *Dezentralisierung* (Elektromotor bzw. Arbeitsplatzrechner, Workstation) identifiziert werden. Des weiteren sind jeweils eine zunehmende weltweite Verfügbarkeit und unaufhaltsame Verbreitung feststellbar.

Als passende Organisationsform für die fraktale Fabrik bot sich der modulare Aufbau an. Bei den Informationssystemen sympathisierte man mit der aus Japan stammenden Idee des bionischen Produktionssystems, das u.a. mit neuronalen Netzen und Fuzzy Logic arbeitete – heute bekannt als Bionic Smart Factory 4.0.

Eine fraktale Fabrik definiert man als offenes System, das aus selbständig agierenden und in ihrer Zielausrichtung selbstähnlichen Einheiten besteht und durch dynamische Organisationsstrukturen einen vitalen Organismus bildet. Das Bild der „Fraktalen Fabrik“ lehnt sich eng an die Mathematik der Fraktale zum Beschreiben natürlicher Strukturen an. Aus der Definition gehen drei wesentliche Eigenschaften von Fraktalen hervor:

- Selbstorganisation,
- Selbstähnlichkeit
- Dynamik.

Das Fraktal stellt sich also als eine selbständig agierende Unternehmenseinheit dar, deren Ziele und Leistung eindeutig beschreibbar sind. Die Selbstähnlichkeit und Selbstorganisation spiegelt sich im Zielsystem wider, das sich aus den Zielen der Fraktale ergibt, widerspruchsfrei ist und der Erreichung der Unternehmensziele dient. Fraktale betreiben operativ, taktisch und strategisch Selbstorganisation. Die ständige Anpassung und Optimierung der dezentralen Organisationsstrukturen gewährleistet eine hohe Systematik. Hierfür sind die Fraktale über ein leistungsfähiges Informations- und Kommunikationssystem miteinander vernetzt. Die Informationsversorgung folgt dem „Holprinzip“, die Fraktale legen (bedarfsorientiert) Art und Umfang ihres Zugriffes auf die Daten fest. Die Grenzen zwischen Unternehmensbereichen und auch zwischen Unternehmen sind unscharf, sie ermöglichen den Informationsaustausch und damit kooperative, partnerschaftliche Beziehungen mit Kunden und Lieferanten. Zwischen Fraktalen bestehen ablauffunktionale, prozessbedingte Verbindungen, die im Interesse des Gesamtsystems aufgebaut werden. Prinzipiell wird versucht, Abläufe ganzheitlich zu bewältigen und keine unnötigen Schnittstellen aufzubauen. Damit trägt die Struktur den Aspekten Geschwindigkeit

und situationsbedingte Anpassungsfähigkeit Rechnung.

Die Strukturierung muss derart erfolgen, dass sowohl die Fraktale als auch die gesamte Fabrik aufgrund der Selbstähnlichkeit derselben Systematik unterworfen werden. Die Erfahrungen aus einer Reihe von Projekten zeigen, dass sich ein Unternehmen praktikabel anhand einer Auflösung von sechs Ebenen beschreiben lässt. Der Vorteil dieses Konzeptes besteht darin, daß die Ebenen durchgängig, aber auch jeweils getrennt voneinander betrachtet und bearbeitet werden können.

- Die kulturelle Ebene beschreibt die Kulturelemente, d.h. die Wertvorstellungen im Unternehmen. Es werden Leitbilder, gemeinsame Wertanschauungen und Grundsätze zum Umgang miteinander und mit der Außenwelt geprägt. Darüber hinaus wird der eigenen Organisation ein Platz zugewiesen und ein Zweck vorgegeben.
- Auf der strategischen Ebene wird die Art und Weise festgelegt, in der die im Unternehmen vorhandenen Ressourcen eingesetzt werden, um ein vorgegebenes Ziel zu erreichen. Es gilt ein Unternehmenszielsystem zu erarbeiten.
- Die sozio-psychologische Ebene umfasst die Gesamtheit aller psychologischen, sozialen und informellen Faktoren, die das Beziehungsgefüge aller Mitarbeiter des Unternehmens bestimmen. Als zentrale Größen dieser Ebenen können Aufbauorganisationen und Kommunikation identifiziert werden.
- Die wirtschaftlich-finanzielle Ebene der Fraktalen befasst sich mit dem Modus der Verrechnung von Leistungen. Im Vordergrund stehen hier betriebswirtschaftliche Kenngrößen hinsichtlich Wirtschaftlichkeit und Leistungsrelevanz.
- Die informationelle Ebene hat primär die Gestaltung der formalen Informationsflüsse und damit vor allem die Ablauforganisation zum Gegenstand.
- Zudem gilt es, die Prozess- und Materialflussebenen zu unterscheiden, auf der die Anordnungen und Flussbeziehungen im Unternehmen hinsichtlich des Zielsystems gestaltet werden.

Das dargestellte Ebenenkonzept unterstützt nicht nur die Darstellung des gesamten Unternehmens, sondern bildet auch die einzelnen Fraktale ab. Es dient somit der Komplexitätsreduzierung von Unternehmen durch Strukturbildung, bei Erhalt der Ganzheitlichkeit der Aufgabenerfüllung.

Man sieht in Abb. 7.1, dass sich die gleiche Struktur auf allen Ebenen wiederholt. Offizielle Kommunikation findet nur über die Spitzen der Dreiecke statt. So-

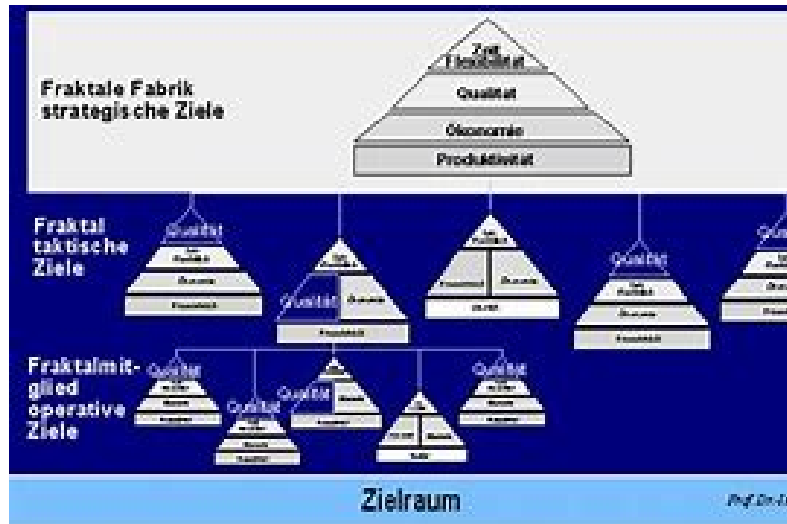


Abbildung 7.1: Struktur einer fraktalen Fabrik

mit kommt eine harmonische Struktur zustande, durch die störende Einflüsse sehr schnell kompensiert werden können.

7.2 Fraktale in der Finanzwelt

Alle zeitlich geordneten Daten in der Finanzwelt zeigen ganz ausgeprägt fraktale Strukturen.

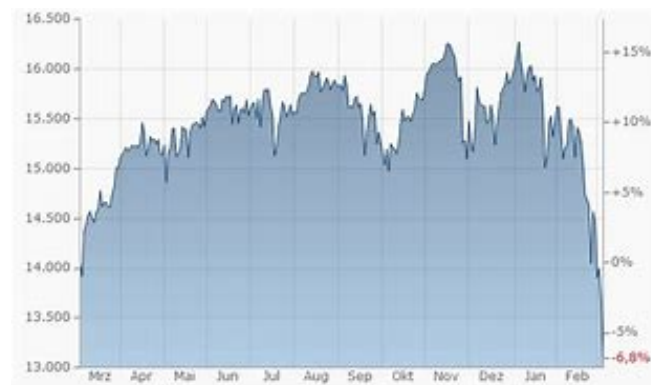


Abbildung 7.2: Der DAX - Kurs im Verlaufe eines Jahres

Die Kurven sehen immer gleich aus, egal, ob man sie über einen Tag, eine Woche oder einen Monat hinaus betrachtet.

Dafür gibt es eine einfache Erklärung. Man verfügt über einen gewissen Bestand an Wertpapieren, der an der Börse verkauft werden könnte, wodurch man einen

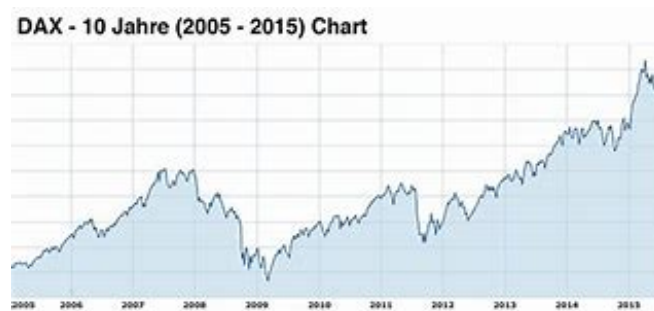


Abbildung 7.3: Der DAX - Kurs im Verlaufe von 10 Jahren



Abbildung 7.4: Der DAX - Kurs seit 1988

bestimmten Erlös erzielen würde. Sinkt der Kurs, wird dieser Erlös geringer, man würde also nicht verkaufen. Steigt der Kurs, wird der Erlös größer. Nun hängt es vom Ermessen des Anlegers ab, mit welchem Gewinnwert er sich zufrieden gibt. An einem bestimmten Punkt des Anstiegs entschließt er (und möglicherweise auch andere) zum Verkauf; er hat einen gewissen Gewinn realisiert. Folgen mehrere Anleger dieser Vorgehensweise, dann sinkt der Kurs ab, und auf diese Weise kommt das Auf und Ab in den Kurven zustande. Hier gehen viele Anleger mehr oder weniger spekulativ vor; sie kaufen eine Aktie, von der sie einen solchen Anstieg erwarten und verkaufen sie wieder, falls der Anstieg zustande kommt und der Gewinn ihren Ansprüchen genügt.

Interessant ist es, wenn man die Kurven von New York, Tokio und Frankfurt miteinander vergleicht. Sie sind nämlich fast identisch, nur zeitlich verschoben. Der Zeitunterschied zwischen Tokio und Frankfurt beträgt sieben Stunden, der Zeitunterschied zwischen Frankfurt und New York sechs Stunden und der zwischen New York und Tokio sogar elf Stunden. Diese Zeitunterschiede werden ausgenutzt und führen zum sogenannten Sekundenhandel. Man kauft beispielsweise in Tokio die Wertpapiere mit einem bestimmten Preis, der Preis steigt, und die Frankfurter Börse beginnt mit diesem höheren Preis; diesen Anstieg kann man ausnutzen und in Frankfurt sofort einen Gewinn realisieren.

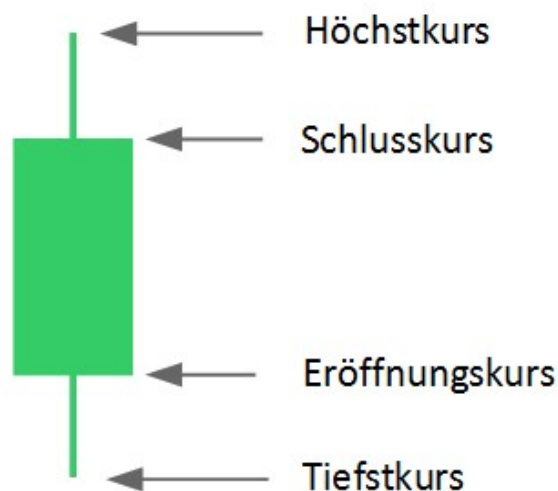


Abbildung 7.5: Das Kerzendiagramm für einen Tag

Der Sekundenhandel ist eine besonders schnelle Form des Daytrading, bei dem Wertpapiere innerhalb kürzester Zeit gekauft und wieder verkauft werden. Während Daytrader-Aktien meist innerhalb eines Tages handeln, werden die Aktien beim Sekundenhandel oft nur wenige Sekunden bis zu einigen Minuten gehalten. Möglich wurden die schnellen Transaktionen durch die Computerbörsen mit ihrer zeitnahen Ausführung, aber erst seit dem Realtime-Trading per Internet können auch normale Privatanleger Daytrading beziehungsweise Sekundenhandel betreiben. Eine besonders wichtige Regel beim Sekundenhandel verbietet es, ein Wertpapier über Nacht zu halten, um während dieser Zeit keine Verluste zu erleiden. Denn um erfolgreich zu sein, ist es im Sekundenhandel enorm wichtig, sofort auf marktbeeinflussende Ereignisse reagieren zu können.

Eine häufig verwendete Darstellungsform sind Kerzendiagramme. Bei der Analyse von Börsenkursen werden sie genutzt, um den Kursverlauf einer Aktie oder eines anderen börsengehandelten Wertes abzubilden.

Jede einzelne Kerze zeigt an, wie sich der Kurs innerhalb eines bestimmten Zeitraums entwickelt hat. Bei einem Tageschart steht jede Kerze für die Kursentwicklung innerhalb eines bestimmten Tages. Wird hingegen ein Wochenchart betrachtet, so zeigt jede Kerze die Kursbewegung innerhalb einer Woche an. Die einzelnen Kerzen zeigen dabei jeweils den ersten Kurs (Eröffnungskurs), den letzten Kurs (Schlusskurs), den höchsten Kurs und den tiefsten Kurs des betrachteten Zeitraums an.

Jede einzelne Kerze besteht aus einem kastenförmigen Körper in der Mitte der

Kerze und zwei Linien, die oberhalb und unterhalb des Kastens herausragen. Die obere und untere Begrenzung des Körpers zeigen den jeweiligen Eröffnungskurs und Schlusskurs an. Liegt der Schlusskurs oberhalb des Eröffnungskurses, wird der Kasten grün eingezeichnet. In diesem Fall ist der Kurs im Handelsverlauf gestiegen. Ist der Kurs hingegen gefallen, liegt der Schlusskurs unter dem Eröffnungskurs und der Kasten wird rot eingezeichnet. Gelegentlich kommt es vor, dass der Kurs an dem selben Punkt schließt, an dem er auch eröffnet hat. In diesem Fall wird an Stelle des Körpers eine gerade Linie eingezeichnet. Diese Kerzenform wird als Doji bezeichnet. Die beiden Linien oberhalb und unterhalb des Körpers werden als Schatten bezeichnet. Die Spitze des oberen Schattens zeigt den höchsten Kurs des betrachteten Zeitraums an, während der untere Schatten den tiefsten Punkt markiert.

Die Gestalt der Kerzen erzählt viel darüber, wie sich der Kurs im betrachteten Zeitraum entwickelt hat. Aus der Länge der Schatten und der Größe der Körper lässt sich schließen, wie sich der Kurs in dem betrachteten Zeitraum verhalten hat. Einige Beispiele dafür finden Sie in den abgebildeten Charts (Abb. 7.6 und 7.7).



Abbildung 7.6: Die Darstellung einer Aufwärtsbewegung

- 1 Ein langer Körper mit kleinen Schatten zeigt eine starke Aufwärtsbewegung (grüner Körper) oder Abwärtsbewegung (roter Körper) an.
- 2 Hat eine Kerze hingegen nur einen sehr kleinen Körper und gleichzeitig lange untere und obere Schatten, so deutet dies auf einen stark schwankenden Handel ohne klaren Trend hin.
- 3 Ein sehr langer Schatten in eine Richtung deutet auf einen Trendumschwung

hin. Ein langer unterer Schatten zeigt an, dass der Kurs zu einem bestimmten Zeitpunkt stark gefallen ist. Dann allerdings setzte eine starke Gegenbewegung ein und der Kurs konnte zum Ende des Handels den Großteil seiner Verluste wieder wettmachen (Abb. 7.6).

- 4 Genau der umgekehrte Fall liegt bei einem langen oberen Schatten vor. Hier ist der Kurs stark gestiegen, konnte sein Hoch aber nicht halten und musste am Ende die meisten seiner zwischenzeitlichen Gewinne wieder abgeben.
- 5 Wenn der höchste Kurs und der tiefste Kurs einer Kerze sehr nahe beieinander liegen, deutet dies auf einen ereignislosen Handel hin.
- 6 Als Gaps werden Lücken zwischen zwei Kerzen bezeichnet. Diese Kurslücken entstehen, wenn die nachfolgende Kerze deutlich oberhalb (Aufwärtsgap) oder unterhalb (Abwärtsgap) der vorherigen Kerze eröffnet. Ein Gap bleibt auf dem Chart nur dann sichtbar, wenn sich der Kurs nicht wieder in Richtung der Vorkerze zurückbewegt und dadurch das Gap schließt.



Abbildung 7.7: Die Darstellung einer Abwärtsbewegung

Ein Aufwärtsgap deutet darauf hin, dass es zur Börseneröffnung deutlich mehr Käufer als Verkäufer gab. Umgekehrt zeigt ein Abwärtsgap an, dass es vor Eröffnung des Handels deutlich mehr Verkauforders als Kauforders gab Abb. 7.7. Aufwärtstrends, Abwärtstrends und Seitwärtsbewegungen In einem Kerzendiagramm lassen sich Aufwärtstrends, Abwärtstrends und Seitwärtsbewegungen relativ leicht erkennen.

In einer Aufwärtsbewegung schließen die Kerzen über dem Schlusskurs ihrer Vorkerzen. Innerhalb der Aufwärtsbewegung gibt es deutlich mehr grüne Kerzen als rote Kerzen. In einer Abwärtsbewegung überwiegen hingegen die roten Kerzen. Die meisten Kerzen innerhalb der Bewegung schließen unterhalb des Schlusskurses ihrer Vorkerzen. Wie der Name schon verrät, bewegt sich der Kurs in einer Seitwärtsbewegung seitlich. Die Kerzen bewegen sich innerhalb eines eng umgrenzten Bereichs auf und ab. In der Seitwärtsbewegung kommt es bei den Kerzen häufig zu Farbwechseln.

Als Nächstes betrachten wir, wie die einzelnen Kerzen im Kerzendiagramm miteinander in Verbindung stehen.



Abbildung 7.8: Eine Stauchungszone

Ein häufig vorkommendes Muster im Kerzendiagramm sind Stauchungszone, in denen mehrere aufeinanderfolgende Kerzen in einem engen Kasten auf und ab pendeln (Punkt 1). Die Körper der einzelnen Kerzen liegen dabei alle mehr oder weniger auf einer Linie. In vielen Fällen liegen zusätzlich alle Kerzenkörper innerhalb der beiden Schatten der größten Kerze. Innerhalb dieses Kastens schwankt der Kurs also unschlüssig zwischen zwei engen Rahmen hin und her, ohne einen klaren Trend ausbilden zu können. Gelingt es dem Kurs schließlich doch aus diesem Kasten auszubrechen, kommt es häufig zu einer heftigen, steilen Kursbewegung in Ausbruchsrichtung (Punkt 2). Allerdings lässt sich vor dem Ausbruch aus dem Kasten nicht voraussagen, in welche Richtung der Ausbruch erfolgen wird. Daher sollte vor einem Einstieg abgewartet werden, bis die erste Kerze außerhalb des Kastens schließt (Abb. 7.8).

Signifikante Hochs oder Tiefs bilden häufig Widerstandszonen für die nachfolgenden Kursbewegungen. An diesen Punkten ist also die Wahrscheinlichkeit höher, dass der Kurs erneut am selben Punkt abprallt und wieder in die Gegenrichtung dreht. Noch stärker sind diese Widerstandszonen, wenn der Kurs bereits mehrfach an der selben Stelle abgeprallt ist.



Abbildung 7.9: Widerstandszonen

In der Abb. 7.9 können Sie solch eine Widerstandszone erkennen. Der Kurs ist hier zweimal, bei Punkt 1 und bei Punkt 2, an der gleichen Linie abgeprallt. Dieses Hoch ist nun aus zwei Gründen interessant: zum einen ist hier die Wahrscheinlichkeit besonders hoch, dass der Kurs an dieser Stelle erneut nach unten drehen könnte. Gleichzeitig würde aber ein Durchbruch durch diese Widerstandslinie als ein besonders positives Zeichen gewertet werden. Viele Händler und Investoren sehen den Bruch einer solchen Widerstandslinie als Indikator für weiter steigende Kurse an und steigen daher kurz nach einem Durchbruch in den Wert ein. Daher kommt es nach einem Durchbruch dieser Linie zu einer erhöhten Nachfrage, was den Kurs weiter nach oben treibt. Vor einem Einstieg empfiehlt es sich allerdings zu warten, ob die Kerze zumindest oberhalb der Widerstandslinie schließt. Manchmal kommt es gerade bei solchen Widerstandslinien zu Bullenfallen, bei denen Investoren in Erwartung steigender Kurse einsteigen, nur um kurz danach mit einer plötzlichen Gegenbewegung konfrontiert zu werden. Sobald sich der Kurs für längere Zeit über die Widerstandslinie bewegt hat, wird die Widerstandslinie zu einer Unterstützungslinie. Entlang der ehemaligen Widerstandslinie bildet sich also eine Unterstützungszone, die es dem Kurs schwerer macht, wieder unter die ehemalige Widerstandslinie zu fallen.

Candlestick-Formationen sind bestimmte Kombinationen von aufeinanderfolgenden Kerzen, die dem Investor Hinweise auf die künftige Kursentwicklung geben.

Während Stauchungszonen, Unterstützungslinien und Widerstände auch in anderen Charts, wie beispielsweise Liniencharts oder Balkencharts, zu erkennen sind, lassen sich Candlestick-Formationen nur im Kerzendiagramm entdecken.

Ein Beispiel für eine Candlestick-Formation sehen Sie in Abb. 7.10. Die drei letzten Kerzen im Kerzendiagramm bilden eine sogenannte Morning-Star-Formation. Die erste Kerze der Formation ist eine lange rote Kerze. Als nächstes folgt eine kleine Kerze, die mit einem Gap eröffnet. Die letzte Kerze ist eine lange grüne Kerze, die ebenfalls mit einem Gap eröffnet hat. Die Morning-Star-Kerzenformation sagt steigende Kurse voraus. In den folgenden Tagen ist also eher mit ansteigenden Kursen zu rechnen.

Generell lassen sich Candlestick-Formationen in zwei Gruppen unterteilen:

- Umkehrsignale deuten an, dass der Kurs nach dem Erscheinen der Candlestick-Formation in die Gegenrichtung drehen wird.
- Fortsetzungsformationen sagen hingegen voraus, dass der bestehende Trend weiterhin Bestand hat.



Abbildung 7.10: Eine Morning-Star-Formation

Die Evening-Star-Formation besteht aus drei aufeinanderfolgenden Kerzen.

- 1 Der Evening-Star folgt immer auf eine Aufwärtsbewegung.
- 2 Die erste Kerze der Evening-Star-Formation ist eine lange grüne Aufwärtskerze. Ist der betrachtete Chart in schwarz und weiß abgedruckt, wäre die erste Kerze eine Kerze mit weißem Körper.

- 3 Die zweite Kerze eröffnet mit einer Kurslücke. Die Kerze eröffnet also oberhalb des Körpers ihrer Vorkerze. Es ist hierbei nicht nötig, dass die Kerze auch über den Schatten der Vorkerze eröffnet. Ebenso wie der Eröffnungskurs liegt auch der Schlusskurs der Kerze über dem Körper der ersten Kerze, sodass die Kurslücke im Chart sichtbar bleibt.
- 4 Die zweite Kerze ist immer eine relativ kleine Kerze. Bei der Evening-Star-Formation ist die Farbe dieser Kerze unwichtig. Handelt es sich bei der zweiten Kerze um ein Doji (also um eine Kerze, bei der der Eröffnungskurs und der Schlusskurs identisch sind) wird die Formation als Evening-Doji-Star bezeichnet. Die zweite Kerze eröffnet also mit einem Kurssprung, kann dann aber dieses Anfangsmomentum nicht halten und schließt am Ende mehr oder weniger am selben Punkt, an dem sie auch eröffnet hat.
- 5 Als letzte Kerze folgt eine rote Abwärtskerze. In der klassischen Evening-Star-Formation muss diese Kerze mit einer Kurslücke nach unten beginnen. Viele Autoren, darunter auch Steve Nison, der die Candlestick-Charts im Westen populär gemacht hat, messen dieser Kurslücke aber nur eine geringe Bedeutung bei und empfehlen, die Formation auch dann zu handeln, wenn keine Lücke vorliegt. Wichtig ist aber, dass es sich bei dieser Kerze um eine lange ausgeprägte Kerze handelt, die deutlich in den Kerzenkörper der ersten Kerze hereinragt.

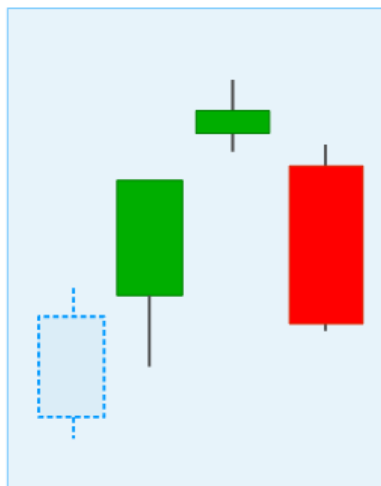


Abbildung 7.11: Eine Evening-Star- Formation

Bei Wechselkursen zwischen unterschiedlichen Währungen finden wir ebenfalls fraktale Strukturen vor.

- Schwach schwankende Aktien schlagen den Markt.



Abbildung 7.12: Der Wechselkurs für Euro - Dollar

Eine der Kernvoraussagen dieses Modells ist, dass Aktien entsprechend ihrer Sensitivität bezüglich einer Markttrendite Beta klassifiziert werden können. Je höher das jeweilige Beta, d.h. die Sensitivität gegenüber Markttrendite ist, desto höher ist auch das Risiko der individuellen Aktie. Arbeiten aus den 70er-Jahren zeigen allerdings, dass dies bei bestimmten Aktien genau nicht der Fall ist.

- „Low Volatility“ als Faktor.

In einem systematischen Faktor-Prozess wird die Volatilität (z.B. über die letzten drei Jahre) als Maß der Schwankung für die Aktien eines Marktes berechnet. Dieses Maß wird insbesondere auch als Risikomaß verstanden. Hiermit können die Aktien hinsichtlich ihres vermeintlichen Risikos in einem regelmäßigen Sortierungsprozess, z.B. halbjährlich, bewertet werden. Wählt man nun die obersten 30 Prozent der Aktien für ein Portfolio aus, also diejenigen, die am wenigsten schwanken, und hält diese bis zum nächsten Umschichtungszeitpunkt, entsteht ein „Low Volatility“-Faktor.

- Investition in unterbewertete und kleine Aktien.

Eine Dekade weiter fokussiert sich die empirische Kapitalmarktforschung auf sogenannte „Value“-Ansätze, also strukturierte Faktor-Investitionen in unterbewertete Aktien. Hierbei kommen klassische Kennzahlen zum Einsatz, wie z.B. das Preis-Buchwert-Verhältnis, das Verhältnis aus Preis zu Gewinnen oder aber auch die Dividendenrendite. Diesen Untersuchungen liegt die Erkenntnis zugrunde, dass Aktien, die als günstig angesehen werden (z.B. über das Preis-Buchwert-Verhältnis gemessen), sich besser entwickeln als Aktien, die als „teuer“ gelten. Weitere Studien zeigten zudem, dass die strukturierte

Investition in kleinere Unternehmen Überrenditen erlauben.

Beide Faktoren gelten hierbei als Risiko-Faktoren, da Aktien, die ein günstiges Preis-Buchwert-Verhältnis aufweisen, oftmals Firmen sind, die unter ökonomischem Stress stehen und mit einem Abschlag gehandelt werden. Analog verhält es sich mit kleinen Firmen, die z.B. spezielle Bilanzrisiken aufweisen können. Insgesamt folgt man damit der Aussage des ursprünglichen Modells: Je mehr Risiken der Investor eingeht, desto höher sei seine zu erwartende Rendite.

- Stark gestiegene Aktien: Momentum.

Anfang der 90er-Jahre wird schließlich durch Jegadeesh und Titmann eine der schwersten Anomalien aus Sicht der Kapitalmarkttheorie veröffentlicht. Das Autoren-Duo zeigt, dass Aktien, die in der Vergangenheit gestiegen sind (z.B. gemessen über die Performance der letzten sechs oder zwölf Monate) und einem systematischen Ranking unterzogen werden, sich besser entwickeln als Aktien, die in letzter Zeit gefallen sind. Aktien, die in Schwung kommen, bleiben somit in Schwung und haben Momentum.

- Sind dies temporäre Anomalien oder ist es strukturell?

Aus Sicht der klassischen Theorie können Überrenditen zum Markt entweder nur durch höhere Risiken entstehend oder andernfalls nur temporär existieren, da die Märkte sich in einen Gleichgewichtsprozess bewegten. Schlimmer noch: Insbesondere die „Low Volatility“ - und die Momentum - Anomalie nutzen rein den Kursverlauf aus der Historie aus und sind somit ein starkes Beispiel dafür, dass sogar die schwache Form der Theorie der Effizienten Märkte, die genau das ausschließt, möglicherweise nicht präzise genug die Marktphänomene beschreibt. Man kann nun einwenden, dass die Beispiele oben nur einen kurzen Zeitraum repräsentieren. Allerdings sind diese Faktoren teilweise bis weit in die Viktorianische Zeit (bei Momentum bis in das Jahr 1801) nachgewiesen.

- Ist es also Zufall oder ist es strukturell? Werden diese Anomalien verschwinden oder ist es die Ursache, dass Investoren doch nicht ganz rational handeln und gewissen „Biases“ unterliegen, wie die Verhaltensökonomie dies unterstellt? Und wie ist es mit den inzwischen weit über 400 gefundenen anderen Anomalien, dem sogenannten „Faktor-Zoo“ und weiterer Diskussionen, in denen auch gefragt wird, warum Momentum und „Low Volatility“ eigentlich nicht in die Kapitalmarktmodelle aufgenommen werden?

In einer seiner letzten Arbeiten skizziert B. Mandelbrot seine Vision trendbasierter, fraktaler Märkte. Dabei handelt es sich um eine Skizze, die mit Hilfe signaltheoretischer Methoden eine alternative Sichtweise auf Faktoren und Markttrenditen erlaubt.

7.3 Die Programmierung von Fraktalen - ihr Nutzen für die Ausbildung

Hier möchte ich die „Fractal Foundation“ als ein Beispiel anführen, dass man sicher an vielen anderen Stellen nachahmen sollte. Ihre Internet-Präsenz beginnt mit der Erklärung: „Wir nutzen die Schönheit der Fraktale, um das Interesse an Wissenschaft, Mathematik und Kunst zu wecken. Unsere Vision ist eine Welt des Staunens und der Neugier, eine Kultur des wissenschaftlichen Forschens, eine Wertschätzung für die Zusammenhänge natürlicher Systeme und ein Verständnis für deren wesentliche Nichtlinearität. Wir sehen eine Welt, in der jeder über mathematische Kenntnisse verfügt und versteht, wie Mathematik ein mächtiges Werkzeug ist, um seine Visionen in die Realität umzusetzen. Wir sind darauf aus, die Kultur zu verändern. Wir haben genug Sporthelden und Filmstars. Was wir brauchen, sind Wissenschaftshelden und Mathe-Stars. Mit groß angelegter Kunst im öffentlichen Raum werben wir für die Schönheit der Mathematik und der Naturwissenschaften und stellen die Kreativität und Intelligenz unserer Schüler unter Beweis. MariaC-The Fractal Foundation wurde 2003 gegründet und ging aus dem in Albuquerque ansässigen "Chaos Club" hervor. Seitdem haben wir über 69 000 Kindern und 55 000 Erwachsenen Fraktale beigebracht. Mit Sitz in Albuquerque haben die meisten unserer Aktivitäten bisher in New Mexico stattgefunden, aber sie weiten sich immer weiter aus.“

Zuerst empfehlen wir Ihnen, den erstaunlichen KOSTENLOSEN Fraktal-Explorer XaoS herunterzuladen

<http://fractalfoundation.org/resources/fractal-software/> ,

mit dem Sie in mathematische Fraktale hineinzoomen können.

Als Nächstes können Sie die Seite Fractivities

<https://fractalfoundation.org/resources/fractivities/>

besuchen, auf der viele verschiedene Projekte für zu Hause, die Schule oder die Natur vorgestellt werden. Wir würden uns freuen, wenn Sie am Fractal Trianglethon teilnehmen würden! Wir brauchen Freiwillige, LehrerInnen und SchülerInnen, die dabei helfen, das größte Sierpinski-Dreieck der Welt zu bauen.

Danach möchten Sie vielleicht den Online-Fraktal-Kurs

<http://fractalfoundation.org/resources/lessons/>

erkunden, der sich ausführlich mit Fraktalen in der Natur und in der Mathematik befasst. Er richtet sich zwar an Schüler der Oberstufe, aber ein Großteil des Materials ist auch für jüngere und fortgeschrittene Schüler geeignet.

Es wird auch eine große Anzahl von Programmen angeboten, mit denen die Schüler leicht einen Einstieg in die Welt der Fraktale schaffen.

Wenn man bereits über Kenntnisse in der Programmierung verfügt, dann kann man auch fertige Programme in mehreren Programmiersprachen finden.

Auf der Webseite

<https://technik.blogbasis.net/ein-fraktal-programmieren-06-08-2013>

wird beispielsweise ein funktionierendes Java-Programm angegeben, das bereits 2013 geschrieben wurde.

8 Chaostheorie

- Chaos und Ordnung

Nach der klassischen griechischen Mythologie bildete das Chaos den ungeordneten Urzustand der Welt. Demnach ging erst aus dem Chaos die spätere geordnete Welt, der Kosmos, hervor. Dieser Mythos von einem zunächst ungeordneten Urstoff findet sich auch in zahlreichen anderen Mythen über die Entstehung der Welt. Mit der christlichen Lehre von der Schöpfung der Weltordnung aus dem Nichts verlor Chaos jedoch in der Spätantike seine ursprüngliche Bedeutung. Unter Chaos versteht man daher heute allgemein Zustände und Vorgänge der Unvorhersagbarkeit und Unberechenbarkeit. Darüber hinaus hat Chaos aber bis in die Gegenwart den Ruf von etwas Unbeschreiblichem, Unheimlichem oder sogar Gefährlichem. Im wissenschaftlichen Bereich ist man auf den meisten Forschungsgebieten beunruhigt, wenn man auf chaotische Erscheinungen (wie Turbulenzen, Vibrationen) stößt, und man versucht häufig, diese Beobachtungen als unbedeutende Störungen abzutun und auszuklammern. Auch im gesellschaftlichen Bereich fürchten sich die meisten vor Chaos. Nicht umsonst werden Extremisten, die eine bestehende politische Ordnung durch Gewaltaktionen zerstören wollen, häufig als Chaoten bezeichnet. Es scheint ein menschliches Grundbedürfnis zu sein, das Unbeschreibliche und Unheimliche der Welt in eine verlässliche Ordnung zu bringen. Dabei wurde in früheren Zeiten eher am Aberglaube festgehalten als auf naturwissenschaftliche Forschung vertraut. In Europa war es bis ins 16. Jahrhundert herrschende (Kirchen-)Lehre, dass die Erde eine Scheibe sei und den Mittelpunkt des Weltalls bilde. Grundlage hierfür waren die frühen Arbeiten des griechischen Geographen und Astronomen Claudius Ptolemäus (um 100 - 160), der ein umfangreiches Kartenwerk von der damals bekannten Welt zusammengetragen hatte. Ptolemäus übernahm darüber hinaus die noch ältere Auffassung des griechischen Philosophen und Naturforschers Aristoteles (384 - 322 v.Chr.), wonach sich die Sonne und alle anderen Himmelskörper um die Erde bewegen (geozentrisches oder ptolemäisches Weltsystem).

Das auf dieser Grundlage entstandene geschlossene Weltbild war der große Schatz der alten Wissenschaften von der Antike bis ins Späte Mittelalter.

Dieses Weltverständnis konnte sich auf eine jahrhundertealte, traditionelle Ordnung berufen, die zudem noch von der herrschenden Kirchenlehre und Bibelauslegung gestützt wurde.

- Determinismus und Reduktionismus

Doch diese alte Ordnung ließ sich auf Dauer nicht mehr mit den tatsächlichen Beobachtungen in der Natur in Einklang bringen. Die im 15. Jahrhundert in Europa entstehenden neuzeitlichen Naturwissenschaften versuchten daher, den Erscheinungen der Natur eine neue, durch Beobachtung überprüfte Ordnung zu geben. Ein erster wichtiger Schritt hierzu waren die Arbeiten des deutsch-polnischen Astronomen Nikolaus Kopernikus (1473 - 1543). Ein weiterer wichtiger Wegbereiter des kopernikanischen Weltsystems war der italienische Physiker und Mathematiker Galileo Galilei (1564 - 1642), der nicht nur die Pendel- und Fallgesetze fand, sondern auch ein verbessertes Fernrohr baute und damit die Phasen der Venus sowie vier Monde des Jupiter entdeckte.

Als entscheidend erwiesen sich jedoch die Forschungen des englischen Physikers und Mathematikers Isaac Newton (1643 - 1727), der bis heute als der Bahnbrecher der neuzeitlichen Naturwissenschaften gilt. Newton entdeckte die Gravitationsgesetze, das Rückstoßprinzip und das Trägheitsgesetz; er fand die Gesetze des Lichtspektrums und entwarf eine überzeugende Theorie vom Licht und den Farben; und er entwickelte, zur gleichen Zeit wie der deutsche Philosoph und Mathematiker Gottfried Wilhelm Leibniz (1646 - 1716), die Differential- und Integralrechnung. In die gleiche Richtung gingen die Arbeiten des französischen Astronomen und Mathematikers Pierre Simon de Laplace (1749 - 1827), der unter anderem die Bewegungen der Himmelskörper genau darstellte. Er und seine Zeitgenossen sahen die Welt als ein Uhrwerk, dessen Vorgänge sich einzeln erforschen, auf bestimmte Gesetzmäßigkeiten reduzieren und nachfolgend verallgemeinern lassen (Reduktionismus). Doch bereits Ende des 19. Jahrhunderts stieß der französische Mathematiker Jules Henri Poincaré (1854 - 1912) auf ein physikalisches Problem, das den Newtonschen Determinismus und den Laplaceschen Reduktionismus wieder in Frage stellte. Poincaré, der auf dem Gebiet der Algebra wichtige Arbeit leistete (Theorie der automorphen Funktionen, Homologietheorie), untersuchte nämlich auch das sogenannte Dreikörperproblem. Mit Hilfe der von Newton entdeckten Gravitationsgesetze lassen sich zwar die Bewegungen von zwei Himmelskörpern eindeutig berechnen (beispielsweise die ellipsenförmige Bewegung des Mondes um die Erde). Sobald jedoch ein dritter Gravitationskörper hinzukommt (zum Beispiel die Sonne mit ihrer Anziehung auf Erde und Mond), lassen sich die zur Berechnung notwendigen Gleichungen nicht

mehr eindeutig lösen. Die Umlaufbahnen können zwar annäherungsweise vorausberechnet werden, aber dennoch tritt bald eine Abweichung ein, so dass langfristige Vorhersagen nicht möglich sind. Daraus folgt auch, dass sich nicht sagen lässt, ob unser Sonnensystem tatsächlich stabil ist - wie ein Uhrwerk verhält es sich jedenfalls nicht. Diese wissenschaftliche Einsicht von Poincaré wurde jedoch zunächst sogar in Fachkreisen als Kuriosität abgetan.[39]

- Seit den 1960er Jahren entstand jedoch eine neue Forschungsrichtung, die Chaos nicht nur als rätselhaften Sonderfall abtat, sondern sich für dessen Eigenschaften interessierte. Forscher verschiedener Fachgebiete entdeckten dabei, dass sich chaotische Systeme im Rahmen einer dynamischen Ordnungsbildung selbst organisieren und verblüffende Ordnungsmuster bilden können. Diese Ausnahmen bestätigen zwar bloß die allgemeine Regellosigkeit des Chaos, aber gleichzeitig lässt sich Chaos nicht länger mit Zufälligkeit gleichsetzen. Die bei der Erforschung chaotischer Zustände und Vorgänge erzielten Ergebnisse fasst man unter dem Begriff „Chaostheorie“ zusammen, und sie beeinflusst heute zahlreiche natur- und geisteswissenschaftliche Bereiche. Viele Verfechter der neuen Forschungsrichtung sind der Auffassung, dass die Chaosforschung (neben der Quantenmechanik und der Relativitätstheorie) die dritte bedeutsame naturwissenschaftliche Errungenschaft des 20. Jahrhunderts ist.

8.1 Nichtlineare Systeme

Aus naturwissenschaftlicher Sicht gehört die Chaostheorie zum Forschungsbereich der nichtlinearen Dynamik. Obwohl im Chaos keine Linearität gemäß Ursache und Wirkung besteht (Kausalbeziehung) und sich chaotische Systeme unvorhersagbar und unberechenbar verhalten, folgen sie selbstverständlich den Naturgesetzen und sind daher auch nicht zufällig. Deshalb spricht man in der Chaosforschung auch von einem gesetzmäßigen (deterministischen) Chaos.

- Bifurkation,
- Phasenraum,
- Attraktor,
- sensitive Abhängigkeit,

- Nichtlinearität,
- exponentielles Fehlerwachstum,
- multiplikative Selbstähnlichkeit,
- Selbstorganisation,
- Bäcker-Transformation,
- Mischung, starke und schwache Kausalität.

Die auftretenden Differential- oder Differenzen-Gleichungen enthalten nichtlineare Funktionen. Diese nichtlinearen Gleichungen zeigen unter bestimmten Umständen interessante Merkmale und Lösungen, beispielsweise Flächen im Phasenraum als Attraktoren, Selbstähnlichkeit und fraktale Strukturen. Wichtige Anwendungen der Nichtlinearen Dynamik finden sich beispielsweise in der Mechanik und der Astrophysik.

$$\ddot{\phi} + 2\frac{\dot{L}}{L}\dot{\phi} + \frac{g}{L}\sin\phi = 0 \quad (8.1)$$

$$L(t) = L_0 - \Delta L_0 \cos(\Omega t). \quad (8.2)$$

Diese beiden Gleichungen beschreiben beispielsweise die Bewegungen einer Schaukel. Hier findet man die Selbstähnlichkeit auf einfache Art und Weise. Die Kurven sind mehr oder weniger gleich, nur die Höhe und die Geschwindigkeit der Bewegungen können sich unterscheiden. Da die Sinus-Funktion auf die nullte Ableitung angewandt wird, handelt es sich um ein nichtlineares System. Im konkreten Fall begrenzt die Sinus-Funktion die Instabilität der auftretenden parametererregten Schwingung, da das System bei größeren Amplituden in stabile Bereiche verstimmt wird. Der nichtlineare Anteil ist die Ursache dafür, dass die Eigenfrequenz von der Schwingungsamplitude abhängt.

Der Phasenraum beschreibt die Menge aller möglichen Zustände eines dynamischen Systems. Ein Zustand wird durch einen Punkt im Phasenraum eindeutig abgebildet. In der Mechanik besteht er aus verallgemeinerten Koordinaten (Konfigurationsraum) und zugehörigen verallgemeinerten Geschwindigkeiten.

Bei n Freiheitsgraden ist der Phasenraum $2n$ -dimensional. Beispielsweise hat ein Gasteilchen im dreidimensionalen Raum $n=3$ Freiheitsgrade, mit den zugehörigen

Impulsen sind das 6 Phasenraumkoordinaten. Ein System (Gas) von N Teilchen hat einen $6N$ -dimensionalen Phasenraum. Es werden aber auch Phasenräume in anderen Anwendungen außerhalb der Mechanik untersucht.

Die zeitliche Entwicklung eines Punktes im Phasenraum wird durch Differentialgleichungen beschrieben und durch Trajektorien (Bahnkurven, Orbit) im Phasenraum dargestellt. Diese sind durch Differentialgleichungen erster Ordnung in der Zeit beschrieben und durch einen Anfangspunkt eindeutig festgelegt (ist die Differentialgleichung zeitunabhängig, sind dies autonome Differentialgleichungen). Dementsprechend kreuzen sich zwei Trajektorien im Phasenraum auch nicht, da an einem Kreuzungspunkt der weitere Verlauf nicht eindeutig ist. Geschlossene Kurven beschreiben oszillierende (periodische) Systeme.

Für Systeme mit bis zu drei Variablen kann der Phasenraum graphisch dargestellt werden. Insbesondere für zwei Variable kann man so die Bewegung (Trajektorien, Phasenraumfluss als Vektorfeld) in einem Phasenraumporträt oder Phasenporträt anschaulich darstellen und qualitativ analysieren.

Der historische Ursprung der Verwendung von Phasenräumen wird häufig auf Joseph Liouville zurückgeführt – wegen des Satzes von Liouville (1838), dass bei konservativen Systemen (mit Energieerhaltung) das Phasenraumvolumen benachbarter Trajektorien zeitlich konstant ist. Liouville hatte aber kein mechanisches System im Auge, sondern bewies den Satz für allgemeine gewöhnliche Differentialgleichungen erster Ordnung, die Verbindung zur Mechanik schlug erst Carl Gustav Jacobi vor. Das Phasenraumkonzept entstand erst, nachdem im weiteren Verlauf des 19. Jahrhunderts die Mathematiker zur Betrachtung höherdimensionaler Räume übergingen. Die erste Verwendung des Phasenraums im heutigen Sinn war bei Ludwig Boltzmann 1872 im Rahmen seiner Untersuchungen der statistischen Mechanik, was 1879 von James Clerk Maxwell übernommen wurde. Das Konzept fand dann Verwendung in den Vorlesungen von Boltzmann und Josiah Willard Gibbs zur statistischen Mechanik, im Artikel zur statistischen Mechanik in der Enzyklopädie der mathematischen Wissenschaften von 1911 von Paul Ehrenfest und Tatjana Ehrenfest (die die Bezeichnung Γ für den Phasenraum einführten) und in der qualitativen Theorie der Differentialgleichungen durch Henri Poincaré.

Ein dynamisches System, dessen Trajektorien den gesamten Phasenraum ausfüllen, also jedem Punkt im Phasenraum beliebig nahe kommen, nennt man ergodisch, siehe auch Ergodenhypothese. Bei konservativen mechanischen Systemen (abgeschlossenen Systemen) ist nach dem Satz von Liouville das Phasenraumvolumen benachbarter Trajektorien zeitlich konstant, bei dissipativen Systemen nimmt es ab (offene Systeme).

In der Hamiltonschen Mechanik ist der Phasenraum ein Beispiel für eine symplektische Geometrie, und die Hamiltonsche Mechanik ist die Geometrie des Phasenraums. Da die Impulse als Ableitungen der Hamiltonfunktion nach den generalisierten Koordinaten definiert sind, ist der Phasenraum dort ein Kotangentialbündel über dem Konfigurationsraum.

In der Quantenmechanik drückt die Heisenbergsche Unschärferelation eine Quantisierung des Phasenraums aus. Die Heisenbergsche Unschärferelation oder Unbestimmtheitsrelation (seltener auch Unschärfeprinzip) ist die Aussage der Quantenphysik, dass zwei komplementäre Eigenschaften eines Teilchens nicht gleichzeitig beliebig genau bestimmbar sind. Das bekannteste Beispiel für ein Paar solcher Eigenschaften sind Ort und Impuls.

Die Unschärferelation ist nicht die Folge technisch behebbarer Unzulänglichkeiten eines entsprechenden Messinstrumentes, sondern prinzipieller Natur. Sie wurde 1927 von Werner Heisenberg im Rahmen der Quantenmechanik formuliert. Die heisenbergsche Unschärferelation kann als Ausdruck des Wellencharakters der Materie betrachtet werden. Sie gilt als Grundlage der Kopenhagener Deutung der Quantenmechanik.

Das Phasenraumporträt gibt eine Möglichkeit, die zeitlichen Entwicklungen dynamischer Systeme graphisch zu analysieren. Dazu werden nur die dynamischen Gleichungen des Systems benötigt, eine explizite Darstellung der Zeitentwicklung, etwa durch analytisches Lösen einer Differentialgleichung, ist nicht nötig.

Als Beispiel folgen einige Elemente der Phasenraumanalyse in einem zweidimensionalen System, das durch die Differentialgleichungen

$$x' = \frac{dx}{dt} \quad , \quad y' = \frac{dy}{dt} \quad (8.3)$$

$$x' = f(x, y) \quad , \quad y' = g(x, y) \quad (8.4)$$

beschrieben ist:

- Einzeichnen des Vektorfelds der Dynamik: Für ein Raster von Punkten wird die Richtung der Bewegung im Phasenraum durch Pfeile dargestellt. Folgt man nun ausgehend von einem bestimmten Startpunkt dem Pfeil, kommt man zu einem neuen Punkt, wo man dieses Vorgehen wiederholen kann. So kann man anhand des Vektorfelds zusätzlich typische Trajektorien in das Phasenraumporträt einzeichnen, die das qualitative Verhalten der zeitlichen Entwicklung einzuschätzen helfen. Beim van-der-Pol-Oszillator zum Beispiel laufen alle Trajektorien auf einen Grenzzyklus zu, was sich anhand von Beispieltrajektorien innerhalb und außerhalb des Zyklus illustrieren lässt. Für

einfache dynamische Systeme kann man Vektorfeld und Beispieltrajektorien oft mit der Hand einzeichnen, bei komplexeren Systemen kann dies durch Computerprogramme geschehen.

- Einzeichnen der Nullklinen: Eine Nullkline bezeichnet eine Kurve im Phasenraum, entlang der sich eine der dynamischen Variablen nicht ändert. Im Fall des obigen zweidimensionalen Systems ist die x-Nullkline durch die Bedingung $x'=f(x,y)=0$ und die y-Nullkline durch $y'=g(x,y)=0$ definiert. Diese Gleichungen lassen sich häufig auch dann nach einer der Variablen auflösen, wenn die Gesamtdynamik nicht analytisch integriert werden kann.
- Bestimmen von Fixpunkten und ihrer Stabilität: Als Fixpunkte werden Zustände bezeichnet, die sich mit der Zeit nicht ändern. Solche Fixpunkte entsprechen den Kreuzungspunkten der Nullklinen im Phasenraum. Im obigen zweidimensionalen System erklärt sich das dadurch, dass an so einem Kreuzungspunkt die Bedingung $g(x, y) = f(x, y) = 0$ erfüllt ist. Durch eine lineare Stabilitätsanalyse kann auch bestimmt werden, ob Trajektorien in der Nähe dieser Punkte angezogen oder abgestoßen werden.
- Finden von Separatrizen: Als Separatrix wird eine Kurve bzw. Fläche bezeichnet, die Phasenraumgebiete mit unterschiedlichem Verhalten voneinander trennt. Gibt es beispielsweise zwei Fixpunkte, die Trajektorien anziehen, gibt es unter Umständen eine Separatrix, die die beiden Einzugsbereiche voneinander trennt. Mit den Orten und der Stabilität aller Fixpunkte bzw. mit dem Vektorfeld der Dynamik können in geeigneten Fällen die Separatrizen ohne weitere Berechnungen gefunden werden.

Die sensitive Abhängigkeit von den Anfangswerten ist eine zentrale Charakteristik chaotischer dynamischer Systeme. Darunter verstanden wird die Eigenschaft solcher Systeme, bei einer nur infinitesimal kleinen Änderung der Anfangsbedingungen ein vollkommen unterschiedliches Systemverhalten im Zeitverlauf zu erzeugen. In diesem Sinn spricht man in der Mathematik von deterministischem Chaos: Die Entwicklung eines chaotischen dynamischen Systems ist als Folge der Unvermeidbarkeit von Messfehlern bei der Bestimmung des Anfangszustandes unvorhersagbar, nicht aufgrund eines stochastischen Verhaltens.

Die Selbstorganisation ist eine Form der Systementwicklung, bei der formgebende oder gestaltende Einflüsse von den Elementen des Systems selbst ausgehen. In Prozessen der Selbstorganisation werden strukturelle Ordnungen, bzw. Musterbildungen erreicht, ohne dass diese nachweislich durch äußere (fremdorganisierte), steuernde Einflüsse entstehen oder linear spezifischen Ursachen zugeordnet werden

können. Selbstorganisation ist eine Eigenschaft komplexer, dynamischer Systeme, die in der Synergetik – der Theorie vom Zusammenwirken der Elemente – untersucht werden. Bei diesem oft spontanen Entstehen von Ordnungsmustern aus der Systemdynamik heraus, spricht man von Emergenz, bzw. von emergenten Phänomenen.

Im politischen oder organisationstheoretischen Gebrauch bezeichnet Selbstorganisation die Gestaltung der Lebensverhältnisse nach flexiblen, selbstbestimmten Vereinbarungen und ähnelt dem Autonomiebegriff. Der politische oder organisationale Gebrauch des Wortes Selbstorganisation wird oft fälschlicherweise mit systemtheoretischen und naturwissenschaftlichen Begründungen legitimiert, steht aber mit diesen Erklärungsmodellen in keinem direkten Ableitungszusammenhang.

Bei Selbstorganisation kann man zwischen autogener (aus eigenen Kräften heraus) und autonomer (selbstbestimmter) Selbstorganisation unterscheiden:

- Selbstreferenz: Selbstorganisierende Systeme sind selbstreferentiell und weisen eine operationale Geschlossenheit auf. Das heißt, „jedes Verhalten des Systems wirkt auf sich selbst zurück und wird zum Ausgangspunkt für weiteres Verhalten“, es wirkt also zirkulär. Operational geschlossene Systeme handeln nicht aufgrund externer Umwelteinflüsse, sondern gemäß der in ihnen entstandenen Formen der Informationserzeugung und -verarbeitung, quasi „aus sich selbst heraus“. Die Ergebnisse interner Verarbeitungsprozesse verändern die Ausgangsbedingungen für nachfolgende Prozesse. Es liegt eine Selbstreferenz und damit eine informationelle Geschlossenheit vor.
- Pfadabhängigkeit: Ein Entwicklungspfad der eingeschlagen wurde, kann nicht so einfach verlassen werden.
- Indeterminiertheit: Welchen Verlauf die Entwicklung nehmen wird, ist letztendlich unvorhersehbar. Die Indeterminiertheit hängt von Zufällen ab, eine kleine Änderung der Ausgangsbedingungen kann zu komplett verschiedenen Pfaden führen.
- Autonomie: Selbstorganisierende Systeme sind autonom, wenn die Beziehungen und Interaktionen, die das System als Einheit definieren, nur durch das System selbst bestimmt werden. Autonomie bezieht sich nur auf bestimmte Kriterien, da eine materielle und energetische Austauschbeziehung mit der Umwelt weiterhin besteht. Bei der vertikalen Autonomie sind Entscheidungsfreiheiten untergeordneter Einheiten scharf abgetrennt. Bei horizontaler Autonomie sind Bereiche auf einer Ebene voneinander unabhängig.

- Zentralisation und Dezentralisation: Bei der Delegation von Entscheidungsbefugnissen auf unterer Ebene spricht man von Dezentralisation, sind die Entscheidungsbefugnisse auf oberer Ebene delegiert hingegen von Zentralisation.
- Redundanz: In selbstorganisierenden Systemen erfolgt keine prinzipielle Trennung zwischen organisierenden, gestaltenden oder lenkenden Teilen. Alle Teile des Systems stellen potentielle Gestalter dar. Mehrere Bereiche können das gleiche tun, was für eine Art Überfluss sorgt (= Redundanz). Redundanz kann Autonomie erhöhen, da es keine strikte Arbeitsteilung gibt.

Um von Selbstorganisation sprechen zu können, müssen folgende (voneinander abhängige) Kriterien erfüllt sein:

- die Evolution eines Systems in eine räumlich/zeitlich organisierte Struktur ohne äußeres Zutun,
- die autonome Bewegung in immer kleinere Regionen des Phasenraumes (sogenannte Attraktoren),
- die Entwicklung von Korrelationen oder raumzeitlichen Mustern zwischen vorher unabhängigen Variablen, deren Entwicklung nur unter dem Einfluss lokaler Regeln steht.

Das Verhalten eines selbstorganisierenden Systems zeigt oft sehr gute Eigenschaften bezüglich der Skalierbarkeit und der Robustheit gegenüber Störeinflüssen oder Parameteränderungen, weshalb sich selbstorganisierende Systeme gut als Paradigma für zukünftige komplexe technische Systeme eignen. Allerdings gibt es keinen einfachen Algorithmus, um die notwendigen lokalen Regeln für ein erwünschtes globales Verhalten zu erzeugen. Bisherige Ansätze bauen zum Beispiel auf manuellem Versuch und Irrtum auf und erwarten ein grundsätzliches Systemverständnis durch den Ingenieur. Als andere Alternative werden oft existierende Systeme in der Natur kopiert, was jedoch das Vorhandensein eines geeigneten Beispiels voraussetzt.

Aktuelle Forschungsarbeiten zielt auf die Anwendung von evolutionären Algorithmen zum Entwurf eines selbstorganisierenden Systems. Ein weiterer Anschlusspunkt ist die Verbindung zur Kybernetik. In den 1940er Jahren entstanden die Wurzeln der Kybernetik, als man Gemeinsamkeiten zwischen dem Gehirn und Computern untersuchte und Schnittstellen verschiedener Einzeldisziplinen erkannte, die menschliches Verhalten, Nachrichtenübertragung, Regelungstechnik, Entscheidungs- und Spieltheorie und statistische Mechanik betrachteten. Gegen Ende des Winters 1943/44

organisierten Norbert Wiener und John von Neumann in Princeton ein gemeinsames Treffen mit Ingenieuren, Neurowissenschaftlern und Mathematikern zu diesem Themenkreis. Ein weiterer Katalysator dieser Entwicklung waren von 1946 bis 1948 die Macy-Konferenzen mit dem Thema „Circular causal, and feedback mechanisms in biological and social systems“ und von 1949 bis 1953 mit dem programmatischen Titel „Cybernetics“.

In gedruckter Form wurde der Begriff von Norbert Wiener erstmals 1948 in *Cybernetics or Control and Communication in the Animal and the Machine* verwendet. Im gleichen Jahr veröffentlichte er in der Zeitschrift *Scientific American* einen grundlegenden Übersichtsartikel zur Kybernetik.

Ab 1948 brachte John von Neumann in seinen Vorlesungen weitere Ergänzungen in die Kybernetik ein. Das Ergebnis dieser Gedankenexperimente war 1953 die Theorie der selbstreproduzierenden Automaten bzw. der Selbstreplikation. Diese Konzepte übertragen Eigenschaften der genetischen Reproduktion auf soziale Meme und lebende Zellen und, seit den 1970ern, auf Computerviren. Norbert Wiener ergänzte 1961 sein Kybernetik-Grundlagenbuch mit zwei weiteren Kapiteln: Über lernende und sich selbst reproduzierende Maschinen sowie Gehirnwellen und selbstorganisierende Systeme.

Der Philosoph und Logiker Georg Klaus etablierte 1953 am Lehrstuhl für Logik und Erkenntnistheorie das Lehrfach Kybernetik an der Humboldt-Universität zu Berlin. Später engagierte er sich für die Gründung einer eigenen Kybernetik-Kommission an der Akademie der Wissenschaften der DDR. Georg Klaus (1912 - 1974) war ein deutscher marxistischer Philosoph sowie Schachspieler und Schachfunktionär. Zum philosophischen Anliegen von Georg Klaus gehörte die Verbindung seiner Philosophie mit den modernen Wissenschaften. Er hatte erkannt, dass auf diesem Gebiet erhebliche Rückstände in der philosophischen Rezeption bestanden. Große Schwierigkeiten hatte die marxistische Philosophie in der Mitte des 20. Jahrhunderts mit einem materialistischen Verständnis von Mathematik und Logik, mit neueren Ergebnissen der Physik (zum Beispiel zu Raum und Zeit) sowie mit Disziplinen wie Semiotik und Kybernetik. Daraus erklärt sich seine intensive Beschäftigung mit der modernen Logik, mit Kybernetik, Semiotik sowie mit einer Allgemeinen Methodologie der Wissenschaften. Seine diesbezüglichen Arbeiten schlossen stets ein, unwissenschaftliche und dogmatische philosophische Interpretationen wissenschaftlicher Ergebnisse zurückzuweisen.

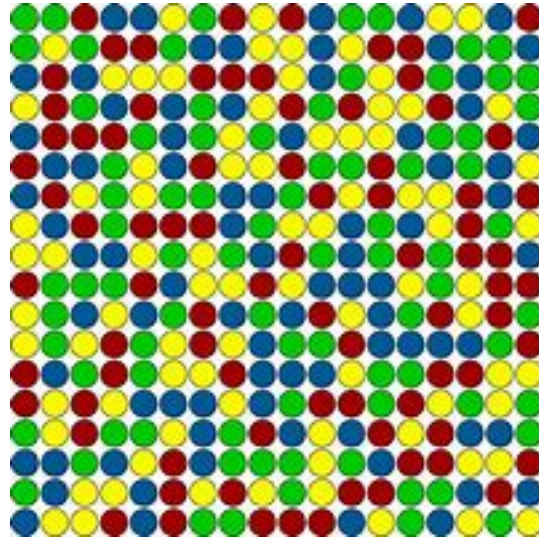


Abbildung 8.1: Ein Vierfarbengitter

8.2 Die Erzeugung von Chaos

Die nachstehende Abbildung wurde dadurch erzeugt, dass man das Ordnen der Farben nicht erlaubt. Es bestand die Forderung, in diesem Gitter 18 x 18 zwei Zeilen und zwei Spalten auszuwählen und dadurch die Eckpunkte von Rechtecken zu bestimmen. Die Färbung der Elemente mit vier Farben sollte so geschehen, dass in keinem Rechteck die Eckpunkte die gleiche Farbe besitzen. Diese Aufgabe hat einen elektronischen Hintergrund, der hier aber keine Rolle spielt.

Es treten für die Farben der Knoten ganz typische Strukturen auf; wir betrachten als Beispiel die Farbe blau.

Man sieht, dass aus irgend einem unerfindlichen Grund die blauen Punkte relativ geradlinig orientiert sind, ebenso die roten und die grünen Punkte. Die gelben Punkte treten aber mehr als schräge Formen auf.

Die fraktalen Strukturen kann man eventuell noch zu einer Vereinfachung der Berechnungen verwenden. Das Ausgangsproblem wurde mit einer logischen Gleichung gelöst, deren Lösungsraum 4^{324} Elemente enthielt.

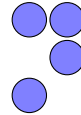
Es ist ein aktuelles Forschungsproblem, ob man nicht in folgender Weise vorgehen kann:

1. Man löst das Problem für ein Gitter 6 x 6 und verwendet diese Lösung entlang der Hauptdiagonalen, insgesamt dreimal. Die Komplexität

ein einzelner Punkt



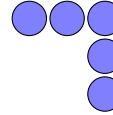
vier zusammenhängende Punkte



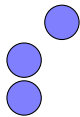
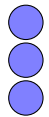
zwei zusammenhängende Punkte



fünf zusammenhängende Punkte



drei zusammenhängende Punkte



zehn zusammenhängende Punkte

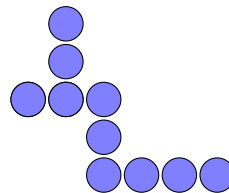


Abbildung 8.2: Kombinationen von blauen Punkten

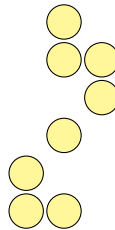


Abbildung 8.3: Kombination von gelben Punkten

2. Man löst das Problem für das nächste Feld 6×6 und verwendet die Lösung zweimal oberhalb der Hauptdiagonalen. Dabei erhöht sich die Zahl der Variablen nicht, nur die Zahl der Nebenbedingungen. Das ist eher günstig, weil es die Anzahl der Lösungen schneller einschränkt.
3. Das Gleiche führt man unterhalb der Hauptdiagonalen durch.
4. Für die Eckfelder verwendet man das gleiche Vorgehen für das Feld 6×6 , man kann aber die Lösung sowohl links unten als auch rechts oben verwenden.

Wenn man zurückgeht, dann man sieht man, dass man diese Darstellung beispielsweise als Perkulationsstrukturen oder als Staubbäden auf einer Fläche auffassen kann. Würde man die vier verschiedenen Farben pulverförmig über einer quadrati-

schen Fläche ausstreuen, erhielte man mit Sicherheit ein ähnliches Bild.

8.3 Der Schmetterlingseffekt

Eine andere Art von undurchsichtigen Verhältnissen kann man unter der Überschrift „Schmetterlingseffekt“ zusammenfassen. Der Schmetterlingseffekt (butterfly effect) ist ein Phänomen der Nichtlinearen Dynamik. Er tritt in nichtlinearen dynamischen, deterministischen Systemen auf und äußert sich dadurch, dass nicht vorhersehbar ist, wie sich beliebig kleine Änderungen der Anfangsbedingungen des Systems langfristig auf die Entwicklung des Systems auswirken.

Die namensgebende Veranschaulichung dieses Effekts am Beispiel des Wetters stammt von Edward N. Lorenz „Kann der Flügelschlag eines Schmetterlings in Brasilien einen Tornado in Texas auslösen?“ Die Analogie erinnert zwar an den Schneeball-effekt, bei dem kleine Effekte sich über eine Kettenreaktion bis zur Katastrophe selbst verstärken. Beim Schmetterlingseffekt geht es jedoch um die Unvorhersehbarkeit der langfristigen Auswirkungen.

Vorarbeiten zu der Theorie leistete Lorenz mit einer Arbeit aus dem Jahre 1963, in der er eine Berechnung zur Wettervorhersage mit dem Computer unternahm. Er untersuchte im Zusammenhang mit langfristigen Wetterprognosen an einem vereinfachten Konvektionsmodell das Verhalten von Flüssigkeiten bzw. Gasen bei deren Erhitzung: hier bilden sich zunächst Rollen (heißes Gas steigt auf einer Seite auf, verliert Wärme und sinkt auf der anderen Seite wieder ab), die bei weiterer Wärmezufuhr instabil werden.

Dieses Verhalten charakterisierte er anhand der drei verbundenen Differentialgleichungen. Das numerische Ergebnis projizierte er in den Phasenraum und erhielt jenen seltsamen Attraktor, der später als Lorenz-Attraktor bekannt wurde: eine unendlich lange Trajektorie im dreidimensionalen Raum, die sich nicht selbst schneidet und aus passendem Blickwinkel die Form zweier Schmetterlingsflügel hat.

Lorenz stieß auf das chaotische Verhalten seines Modells eher zufällig. Um Rechenzeit zu sparen, hatte er bei der numerischen Lösung der Gleichungen auf Zwischenergebnisse bereits durchgeführter Berechnungen zurückgegriffen, hierbei jedoch nur drei Dezimalstellen berücksichtigt, obwohl der Computer mit einer Genauigkeit von sechs Dezimalstellen rechnete. Das Resultat waren zunehmende Abweichungen im Zeitverlauf zwischen den alten und neuen Berechnungen, was Lorenz zu seinen Aussagen über die Sensitivität gegenüber den Anfangsbedingungen bewog. Von nahezu

demselben Ausgangspunkt divergierten die Wetterkurven, bis sie schließlich keine Gemeinsamkeit zeigten.

Bei seiner ersten Berechnung gab er einen Startwert für eine Iteration auf sechs Dezimalstellen genau an (0,506127), bei der zweiten Berechnung auf drei (0,506), und obwohl diese Werte nur um etwa $1/10.000$ voneinander abwichen, wich im weiteren Verlauf diese Berechnung mit der Zeit von der ersten stark ab.

Der Schmetterlingseffekt tritt bei Systemen auf, die deterministisches chaotisches Verhalten zeigen. Diese Systeme besitzen die Eigenschaft, dass sich beliebig kleine Unterschiede in den Anfangsbedingungen (Climaten) im Laufe der Zeit zu starken Unterschieden im System führen; sie sind also sensitiv abhängig von den Anfangswerten. Dieses Phänomen kann mittels der sogenannten Ljapunow-Exponenten quantifiziert werden.

Einige Beispiele:

- Meteorologie: Da die Anfangsbedingungen experimentell immer nur mit endlicher Genauigkeit bestimmt werden können, ist eine Konsequenz dieses Effekts für solche Systeme, dass es unmöglich ist, ihr Verhalten für längere Zeit vorherzusagen. Zum Beispiel kann das Wetter für einen Tag relativ genau prognostiziert werden, während eine Vorhersage für einen Monat kaum möglich ist. Selbst wenn die ganze Erdoberfläche mit Sensoren bedeckt wäre, diese nur geringfügig voneinander entfernt lägen, bis in die höchsten Lagen der Erdatmosphäre reichten und exakte Daten lieferten, wäre auch ein unbegrenzt leistungsfähiger Computer nicht in der Lage, langfristig exakte Prognosen der Wetterentwicklung zu machen. Da das Computermodell die Räume zwischen den Sensoren nicht erfasst, kommt es zu geringfügigen Divergenzen zwischen Modell und Realität, die sich dann positiv verstärken und zu großen Unterschieden führen.

Beispielsweise lassen sich aus den Daten von 1000 Wetterstationen einigermaßen zuverlässige Prognosen über einen Zeitraum von vier Tagen machen. Für entsprechende Vorhersagen über elf Tage bräuchte man bereits 100 Millionen gleichmäßig über die Erde verteilte Messstationen. Absurd wird das Vorhaben, wenn sich die Vorhersage über einen Monat erstrecken soll; dann wäre eine Wetterstation auf je fünf Quadratmillimeter Erdoberfläche nötig.

Allerdings ist das Lorenz-Modell eigentlich viel chaotischer als der tatsächliche Wetterverlauf. Die Gleichungen sind viel instabiler als die grundlegenden physikalischen Gleichungen. Der Mathematiker Wladimir Igorewitsch Arnold



Abbildung 8.4: Der Schmetterlingseffekt

gibt als eine prinzipielle obere Schranke für die Wettervorhersage zwei Wochen an.

- Planetenbahnen: Wenn mehr als zwei Himmelskörper gravitativ aneinander gebunden sind, können minimale Änderungen der Ausgangssituation im Laufe der Zeit zu großen nichtvorhersagbaren Änderungen der Bahnen und Positionen führen. Dieses Verhalten ist Thema des Dreikörperproblems.

Der Lorenz-Attraktor wurde bereits in Abb. 6.2 dargestellt, die entsprechenden Gleichungen sind in 6.3 angegeben.

Das Dreikörperproblem zeigt die gleiche Empfindlichkeit gegenüber den Anfangsbedingungen. Es besteht darin, eine Lösung (Vorhersage) für den Bahnverlauf dreier Körper unter dem Einfluss ihrer gegenseitigen Anziehung (Newtonsches Gravitationsgesetz) zu finden. Um quantitative Resultate zu erlangen, muss es im allgemeinen Fall bislang numerisch gelöst werden.

Das Dreikörperproblem galt seit den Entdeckungen von Johannes Kepler und Nikolaus Kopernikus als eines der schwierigsten mathematischen Probleme, mit dem sich im Laufe der Jahrhunderte viele bekannte Mathematiker wie Alexis-Claude Clairaut, Leonhard Euler, Joseph-Louis Lagrange, Thorvald Nicolai Thiele, George William Hill und Henri Poincaré beschäftigten. Im allgemeinen Fall erfolgt die Bewegung chaotisch und kann nur numerisch berechnet werden.

Den Spezialfall, dass einer der drei Körper eine verschwindend kleine Masse hat und seine Wirkung auf die beiden anderen vernachlässigt werden kann, bezeichnet man als eingeschränktes Dreikörperproblem. Er spielt in der Astronomie eine wichtige Rolle (z. B. bei Forschungssatelliten wie bei der Planetary Grand Tour), die auf

das Problem der Lagrange-Punkte führt.

Das Zweikörperproblem ist durch die Keplerschen Gesetze analytisch lösbar. Dagegen sind die Integrale im Fall von mehr als zwei Himmelskörpern keine algebraischen Integrale mehr und nicht mehr mit elementaren Funktionen lösbar. Karl Frithiof Sundman konnte Anfang des 20. Jahrhunderts als Erster eine analytische Lösung des Dreikörperproblems in Form einer konvergenten Potenzreihe angeben, unter der Annahme, dass der Gesamtdrehimpuls des Systems nicht verschwindet und es deshalb nicht zu einem Dreierstoß kommt, bei dem der Abstand aller drei Körper Null beträgt. Für praktische Berechnungen ist Sundmans Lösung allerdings nicht brauchbar, da bei der Summe mindestens 10 hoch 8.000.000 Terme berücksichtigt werden müssten, um eine hinreichende Genauigkeit zu erzielen.

Die Stabilität eines Dreikörpersystems wird durch das Kolmogorow-Arnold-Moser-Theorem beschrieben: falls ein ungestörtes System nicht entartet ist, dann werden für genügend kleine autonome hamiltonsche Störungen die meisten nicht resonanten Tori lediglich leicht deformiert, so dass auch im Phasenraum des gestörten Systems invariante Tori existieren, die von den Phasenbahnen dicht und quasiperiodisch umsponnen werden, wobei die Frequenzen rational unabhängig sind. Diese invarianten Tori bilden die Mehrheit in dem Sinne, dass das Maß des Komplements ihrer Vereinigung klein ist, wenn die Störung schwach ist.

Näherungs- oder exakte Lösungen sind in manchen Fällen möglich: Wenn die Masse eines der Himmelskörper klein ist, dann löst man das Dreikörperproblem iterativ, heutzutage mit Computern, oder berechnet Bahnstörungen, die der kleinste (leichteste) Körper durch die größeren (schwereren) erleidet. Exakt lösbar ist der schon erwähnte Sonderfall des Gleichgewichts der Anziehungskraft zweier großer (schwerer) Körper auf einen verschwindend kleinen (leichten) Körper (unter Berücksichtigung der im sich drehenden Bezugssystem auftretenden Scheinkräfte).



Abbildung 8.5: Das Dreikörperproblem

Ein ebenfalls sehr chaotisches Verhalten kann man mit dem Stichwort „Herdentrieb“ beschreiben. Es wurde schon im Kapitel über „Fraktale in der menschlichen Gesellschaft“ angedeutet. Es soll aber jetzt wesentlich breiter dargestellt werden.

Als Herde bezeichnet man in der Zoologie überwiegend eine Ansammlung großer, in der Regel gleichartiger, oft ausschließlich pflanzenfressender Amnioten, vor allem großer Säugetiere und großer Laufvögel. Die Bezeichnung ist unabhängig davon, ob es sich um Wildtiere oder um Haustiere handelt. Insbesondere in Herden zusammenlebende, sowohl wilde als auch domestizierte Huftiere werden als Herdentiere bezeichnet.



Abbildung 8.6: Eine Gnu-Herde in Tansania

Bei einer Herde handelt es sich um einen mehr oder weniger einheitlich koordinierten Sozialverband von weniger als zehn bis einigen tausend Individuen. Je nach Größe kann eine Herde ein anonymer Sozialverband sein, in dem die meisten Individuen einander nicht kennen, oder ein individualisierter Sozialverband, in dem die Tiere miteinander vertraut sind. Unter bestimmten Umständen vereinigen sich vor allem bei Wiederkäuern kleinere Gruppen, bei denen die Gruppenmitglieder engere Bindungen zueinander haben, zu großen anonymen Herden. Solche großen Herden können dann auch aus Tieren verschiedener Arten zusammengesetzt sein, beispielsweise aus Gnus, Zebras und Straußen.

Kleinere Herden können entweder locker und ohne ein führendes Tier organisiert sein, wie bei männlichen Hirschen außerhalb der Paarungszeit, oder hierarchisch mit einem Leit- oder Alphetier, wie bei Pferden. Das Herdenverhalten ist von vielen Faktoren abhängig, sei es die Verfügbarkeit der Nahrung, sei es artspezifisches Fortpflanzungsverhalten. Durch eine große Herde mit vielen wachsenden Tieren sinkt die Wahrscheinlichkeit für das einzelne Tier, von einem Raubtier erbeutet zu werden. Pinguine stehen beim Überwintern in großer Zahl dicht zusammen, das reduziert den Verlust an Körperwärme. Das Herdenverhalten gilt als evolutionäre Anpassung.

Wir kehren noch einmal zum Thema von Kapitel 4 zurück und betrachten das Verhalten von Menschenmassen unter dem Gesichtspunkt der Chaostheorie. Massenpsychologie ist ein Teilgebiet der Sozialpsychologie und beschäftigt sich mit dem Verhalten von Menschen in Menschenansammlungen. Ausgang für die Theoriebildung der Massenpsychologie ist die zum allgemeinen Erfahrungsschatz gehörende Tatsache, dass große Menschenmassen ein oft überraschend und irrational erscheinendes Verhalten zeigen, zum Beispiel die Auslösung einer Panik aufgrund eines eher unbedeutenden Anlasses.

Wichtige Entscheidungen in einer Gruppe werden nicht von einzelnen Individuen getroffen, sondern von der Masse durch eine Abstimmung herbeigeführt, um durch die Zusammenarbeit ein Ziel zu erreichen. In der Geschichte sind große Menschenmassen imstande gewesen, dramatische und plötzliche soziale Veränderungen außerhalb der etablierten Rechtsprozesse einzuleiten. Kollektive Zusammenarbeit wird von einigen verdammt, von anderen unterstützt. Sozialwissenschaftler haben einige unterschiedliche Theorien aufgestellt, um massenpsychologische Phänomene zu erklären und zu erläutern, inwiefern sich das Gruppenverhalten vom Verhalten der Einzelpersonen innerhalb der Gruppe signifikant unterscheidet.

- Ansteckungstheorie: eine frühe Theorie zum kollektiven Verhalten hat der französische Soziologe Gustave Le Bon mit seinem Hauptwerk *Psychologie der Massen* (1895), formuliert. Nach Le Bons Ansteckungstheorie (Contagion Theory) üben soziale Gruppen eine hypnotische Wirkung auf ihre Mitglieder aus. Geschützt in der Anonymität der Menge, geben Menschen ihre persönliche Verantwortung auf und ergeben sich den ansteckenden Gefühlen der Masse. Die Menschenmenge entwickelt so ein Eigenleben, wühlt die Gefühle auf und verleitet die Personen tendenziell zu irrationalen Handeln. Wie Clark McPhail ausführt, offenbaren systematische Untersuchungen allerdings, dass „die verrückte Masse“ kein Eigenleben getrennt von den Gedanken und Intentionen ihrer Mitglieder führt. Norris Johnson, der eine Panik während eines Who-Konzerts 1979 erforschte, kam zu dem Schluss, dass die Masse aus vielen Kleingruppen bestand, deren Mitglieder vorwiegend versuchten, einander zu helfen.

Le Bons Arbeiten bilden den Ausgangspunkt von Sigmund Freuds Studie „Massenpsychologie und Ich-Analyse“, allerdings arbeitete Le Bon über „instabile, sich zusammenrottende Massen, während Freud eher höher organisierte, stabile Massen wie Kirche und Heer analysierte.“

Wilhelm Reich formulierte aus seiner eigenen Weiterentwicklung der Psychoanalyse 1933 sein Werk „Die Massenpsychologie des Faschismus“; Elias Ca-

netti belegt diese These in seinem eher literarisch fundiertem Werk „Masse und Macht“.

- Annäherungstheorie: die Annäherungstheorie (Convergence Theory) postuliert, dass das Massenverhalten nicht von der Masse selbst ausgeht, sondern von einzelnen Individuen in die Gruppe hineingetragen wird. Die Gruppenbildung selbst läuft auf die Annäherung von Individuen mit ähnlicher Gesinnung hinaus. Mit anderen Worten ausgedrückt: Die Ansteckungstheorie besagt, dass Gruppen Menschen zu bestimmtem Handeln veranlassen; die Annäherungstheorie dagegen sagt das Gegenteil: Menschen, die in einer bestimmten Weise handeln wollen, schließen sich zusammen.

Ein Beispiel für die Annäherungstheorie ist ein Phänomen, welches sich manchmal beobachten lässt, wenn in einer zuvor homogenen Gegend vermehrt Immigranten auftauchen und Mitglieder der bereits existierenden Gemeinschaft sich (offenbar spontan) verbünden, um die Zuzügler zu bedrohen. Anhänger der Konvergenztheorie glauben, dass in solchen Fällen nicht die Masse den Rassenhass oder Gewalt erzeugt, sondern, dass die Feindseligkeit längere Zeit in vielen Bewohnern gebrodelt hat. Die Masse entsteht aus der Annäherung derjenigen Menschen, die gegen die neuen Nachbarn sind. Die Konvergenztheorie besagt, dass das Verhalten der Masse selbst nicht irrational ist, vielmehr die Personen Ansichten und Werte ausdrücken, die in der Gruppe existieren, so dass die Reaktion des Pöbels nur das rationale Produkt von weit gestreuten populären Gefühlen ist.

- Führungstheorie: nicht zwangsläufig, aber oft werden Massen bei Le Bon (und Gabriel Tarde) von selbstgewählten – oder zumindest kollektiv anerkannten – Führern angeleitet und dabei mitunter zu Taten verführt, die sie außerhalb der Masse, als Individuen wahrscheinlich nicht begehen würden. Ein Sonderfall solcher Massenführerschaft besteht dann, wenn der Führer es versteht, die geballte Gemeinschaftssolidarität der Masse und damit auch ihr kraftvolles Selbstwertgefühl auf sich zu beziehen. Der Führer „verkörpert“ dann die Masse, ihre Ziele und Werte, ihr Denken und ihre Emotionen; Er stellt sich als ihr „höchster Diener“ dar und wird erst durch diese scheinbare Unterwerfung ihr Herr. Dies kann soweit führen, dass die außerordentlich ausgeprägte gegenseitige Sympathie, die die Angehörigen der Masse füreinander aufbrachten, nun mehr und mehr auf den Führer konzentriert wird: die Masse beginnt ihren Führer zu lieben.

Hier entsteht, was in der Sozialpsychologie unter dem Begriff der charismatischen Massenführerschaft zu verstehen ist. Die Führerverherrlichung mündet ein in die Zuschreibung besonderer „genialer“, „wunderbarer“, ja fast „göttli-

cher Eigenschaften“, die den Führer zur Leitung der Masse besonders befähigen und das „blinde Vertrauen“, das man in ihn setzt, rechtfertigen. Phantasievolle Legenden, Anekdoten und Gerüchte der Masse, mithin aber auch gezielte Massenpropaganda des Führers und seiner Gefolgschaft bestätigen und festigen solche wertgebenden Zuschreibungen. Die Masse beginnt an ihren Führer zu glauben wie an eine Heilsgestalt. Blinder Massengehorsam, bis „in den Tod“ wird dem Führer zuweilen freiwillig entgegengebracht und nach seiner Etablierung von ihm wie selbstverständlich eingefordert. Charismatische Massenführerschaft dieser extremen Form, die im Falle religiöser, militärischer und politischer Führerschaft am häufigsten auftritt, wird begünstigt durch den „Messianismus der Massen“ (Michael Günther), einer Ausformung der besonders ausgeprägten Religiosität der Massen, von der schon bei Le Bon die Rede ist.

Dieser Messianismus schafft ein Machtvakuum, eine „charismatische Lücke“: die Masse entwickelt insbesondere beim Vorherrschen starker kollektiver Emotionen (wie Todesangst oder höchster Verwirrung) und geringem horizontalen Organisationsgrad den Wunsch nach Klarheit und Führung. Ihr Überlebenswille konzentriert sich auf die Hoffnung einer begnadeten Leitung, dies umso stärker, je verzweifelter ihre Lage erscheint. Ist ihr Selbstwertgefühl so beschädigt, ist die Masse überzeugt, sich nicht mehr aus eigener Kraft aus hoffnungsloser Situation zu befreien, ist sie auch bereit zur Unterordnung, zur Anerkennung einer höheren Wertigkeit, als der eigenen. Der Hoffnungsträger, der geschickt die charismatische Lücke nutzt, wird wie ein Messias begrüßt, der vom Schicksal geschickt wurde. Rettung aus der Not scheint in greifbare Nähe gerückt, Freude und Erleichterung breiten sich aus, Dankbarkeit wird dem entgegengebracht, der die "letzte Hoffnung" verkörpert. Der Führer nutzt die Suggestibilität der höchst emotionalisierten Masse gekonnt aus und verstärkt seinen überweltlichen Nimbus, indem er die kollektiven Hoffnungen der Masse aufgreift und sich als Messias präsentiert, der von höheren Mächten geschickt wurde, um eine bestimmte Sendung zu erfüllen.

Das komplexe Zusammenspiel von Massenhoffnungen und Führungschancen des Charismatisierten funktioniert allerdings nur solange, wie sich die Führungsfigur auch bewährt: schlägt alles fehl, enttäuscht der Charismatisierte allzu offensichtlich die überbordenden Hoffnungen der Masse, wird ihm die Legitimität, die er als Massenführer gewonnen hatte, auch rasch wieder entzogen, die Masse folgt und gehorcht ihm nicht mehr und entzieht ihm ihre Zuneigung. Die Instabilität charismatischer Massenmacht liegt begründet in der Möglichkeit der raschen Entzauberung des Charismas bei mangelnder Bewährung. Auch daher sind charismatisierte Massenführer oftmals bemüht, ihre Macht mit rationaler- und traditionaler Herrschaftspsycholo-

gie zu vereinen, die mehr Stabilität garantiert und auch über Niederlagen und Schicksalsschläge hinweghelfen kann – die charismatische Massenmacht für sich genommen kaum überstehen könnte. Paradigmen solcher Verschränkungen charismatischer Massenföhrerschaft mit rationaler und traditionaler Herrschaftspsychologie lassen sich quer durch die Geschichte finden: Alexander der Große gibt ein Beispiel, ebenso Gaius Julius Cäsar und Napoleon Bonaparte. Das 20. Jahrhundert brachte besonders viele charismatisierte Massenföhrer an die Macht, so Benito Mussolini, Wladimir Iljitsch Lenin, Josef Stalin, Adolf Hitler, Mao Tse-Tung und zahlreiche weniger bekannte Föhrergestalten.

- Anwendungsgebiete: Neben der Politik ist der Finanzmarkt ein wichtiges Anwendungsgebiet, in welchem sich die massenpsychologische Forschung etablieren kann. Die Zusammenföhrung von Wissen über Anlegerverhalten mit den Erkenntnissen der Massenpsychologie offenbart neue Modelle und Herangehensweisen für realistischere Erklärungskonzepte der Finanzmarktdynamik. Denn der Konjunktur von Boom und Depression ist ein wiederkehrendes Element in der Finanzmarktgeschichte und traditionellen ökonomischen Theorien und Finanzmarktmodelle (z. B. Markteffizienzhypothese) versagen aber bei der Erklärung und Vorhersage solcher Trends und den ihnen zugrundeliegenden Verhaltensweisen der Marktteilnehmer. Denn sie berücksichtigen nicht den gesamten Menschen, sondern nur eine akademische Abstraktion jener Aspekte des menschlichen Verhaltens, die sie für ökonomisch relevant halten. Und sie vergessen auch die Gesellschaft, mit der die Märkte untrennbar verbunden sind.

An diesem Punkt setzt die Massenpsychologie, die unter anderem auf den Konzepten der gegenseitigen sozialen und psychologischen Ansteckung sowie der menschlichen Neigung zur Orientierung und Nachahmung anderer im sozialen Umfeld basiert, an. Die Erforschung kollektiver Dynamiken liefert noch einen weiteren Beitrag zum besseren Verständnis der Prozesse an den Finanzmärkten, indem sie auf den Zusammenhang zwischen kurzfristigen Entwicklungen und langfristigen Veränderungsprozessen hinweist. Basierend auf dem Prinzip langer und kurzer Zyklen wird in der Massenpsychologie zwischen bewusst erkannten, kurzlebigen Auswirkungen und den ihnen zugrunde liegenden langsamen, subtilen und oftmals nicht erkannten Entwicklungen unterschieden. Durch diese zentrale Erkenntnis, die im Wesentlichen schon auf Gustave LeBon zurückgeht, leistet die massenpsychologische Forschung einen Beitrag zur Beschreibung des Zusammenwirkens eines seit den 1960er Jahren angehäuften Schuldenbergs und der periodischen Entstehung von Boom-Krisen-Zyklen während der vergangenen Jahrzehnte.

- Die Weisheit der Vielen: die Weisheit der Vielen – weshalb Gruppen klüger sind als Einzelne („The Wisdom of Crowds. Why the Many Are Smarter than the Few and How Collective Wisdom Shapes Business, Economies, Societies and Nations“) ist der Titel eines Buchs von James Surowiecki, das 2004 erschienen ist. Er argumentiert darin, dass die Kumulation von Informationen in Gruppen zu gemeinsamen Gruppenentscheidungen führen, die oft besser sind als Lösungsansätze einzelner Teilnehmer (sogenannte kollektive Intelligenz).

Das Buch präsentiert zahlreiche Fallstudien und Anekdoten, um seine Argumentation zu illustrieren. Dabei werden viele Fachgebiete berührt, hauptsächlich aber die Ökonomie und die Psychologie.

Die einleitende Geschichte erzählt von Francis Galtons Überraschung, dass Besucher der westenglischen Nutztiermesse 1906 im Rahmen eines Gewinnspiels das Schlachtgewicht eines Rindes äußerst genau schätzten, wenn man als Schätzwert der Gruppe den Median aller 787 Schätzungen annahm. Der Mittelwert der Einzelschätzungen stimmte sogar exakt und war damit besser als die jedes einzelnen Teilnehmers, darunter manche Experten wie z. B. Metzger.

Das Buch bezieht sich auf unterschiedliche Gruppen unabhängig entscheidender Personen, nicht auf Phänomene der Massenpsychologie. Er zieht Parallelen zu statistischen Auswahlverfahren, wonach eine unterschiedliche Gruppe individuell entscheidender Menschen eher die Gesamtheit aller möglichen Ausgänge eines Ereignisses repräsentieren kann und damit in der Lage ist, bessere Voraussagen für die Zukunft zu treffen. Surowiecki unterteilt Entscheidungen in drei Hauptgruppen auf, die er als Problemfelder klassifiziert:

- Kognition: dieses Problemfeld umfasst Entscheidungen, bei denen es eine konkrete Lösung gibt, die durch den Einsatz der kognitiven Fähigkeiten erkannt werden kann. Surowiecki argumentiert, dass dies einer Gruppe viel genauer, schneller und unabhängiger von politischen Kräften gelingen könne als Experten oder Expertengruppen.
- Koordination: Koordination von Verhalten enthält die Optimierung der Nutzung eines Restaurants oder unfallfrei zu fahren. Das Buch enthält viele Beispiele aus der experimentellen Ökonomie, dieser Abschnitt beruht aber mehr auf natürlich vorkommenden Phänomenen, wie Fußgänger, die die Gehweg-Benutzung optimieren oder die Auslastung populärer Restaurants. Er untersucht, wie geteilte Überzeugungen/Normen

innerhalb einer Kultur erstaunlich genaue Voraussagen über die Reaktionen anderer Mitglieder dieser Kultur erlauben.

- Kooperation: wie Gruppen von Menschen ein Vertrauensnetzwerk aufbauen können, ohne dafür eine zentrale Kontrolle über ihr Verhalten oder eine direkte Durchsetzung der Regeln zu benötigen. Dieser Abschnitt spricht sich besonders für einen freien Markt aus.

Nicht alle Gruppen sind weise. Beispiele für solche Überlegungen sind zum Beispiel aufgebrachte Menschenmengen oder Investoren an der Börse nach einem Boom oder Crash. Untersuchungen sind dahingehend nötig, um mehr Beispiele für fehlerhafte Gruppenintelligenz aufzudecken und zu vermeiden. Dennoch ist es möglich, Schlüsselkriterien zu definieren, die eine weise Gruppe von einer irrationalen Gruppe unterscheiden.

- Meinungsvielfalt: Jeder Mensch besitzt unterschiedliche Informationen über einen Sachverhalt, so dass es immer zu individuellen Interpretationen eines Sachverhaltes kommen kann.
- Unabhängigkeit: Die Meinung des Einzelnen ist nicht festgelegt durch die Ansicht der Gruppe.
- Dezentralisierung: Hier steht die Spezialisierung im Mittelpunkt des Fokus, um das Wissen des Einzelnen anzuwenden.
- Aggregation: Es sind Mechanismen vorhanden, um aus Einzelmeinungen eine Gruppenmeinung zu bilden.

Surowiecki untersuchte Situationen, in denen die Gruppe einen sehr schlechten Ruf aufbaute und argumentierte, dass in diesen Situationen das Wissen oder die Zusammenarbeit fehlerhaft sei. Dies geschah seiner Ansicht nach dadurch, dass die Gruppenmitglieder zu sehr auf die Ansichten anderer Menschen hörten und ihnen nacheiferten, statt sich selbst ein Bild über die Situation zu machen und zu differenzieren. Er nennt verschiedene Details von Experimenten, wonach die Gruppengewohnheiten durch einen ausgewählten Sprecher bekannt werden. Er behauptet obendrein, dass der Hauptgrund für die intellektuelle Konformität einer Gruppe hauptsächlich darin besteht, systematische Fehlentscheidungen zu treffen.

Wenn die entscheidende Instanz nicht in der Lage ist, die Gruppe zu akzeptieren, so führe das laut Surowieckis Aussagen dazu, dass das Personenrecht und

das Recht zur Selbstinformation verloren gehen. Die Zusammenarbeit in der Gruppe kann auf diese Weise nur so gut, beziehungsweise eher schlechter als besser sein, als das klügste Mitglied (Die Möglichkeit besteht dem Anschein nach). Detaillierte Fallbeispiele schließen folgende Fehler ein:

- Zentralismus: Das Unglück der Weltraumfähre Columbia, dessen Verschulden sich auf die bürokratische Hierarchie des NASA-Managements verschob, da es nichts von den Warnungen der Ingenieure gewusst haben will.
- Meinungsunterschiede: Die US-amerikanische Gemeinschaft konnte das Attentat des 11. September 2001 nicht verhindern, da Informationen von einer Unterbehörde vermutlich nicht an eine andere weitergeleitet worden sind. Laut Surowiecki arbeiten Gruppen am besten, wenn sie sich ihre Arbeit selbst aussuchen und sich selbst Informationen, die sie benötigen, besorgen (in diesem Fall IQ-Forscher). Die Isolation des SARS-Virus dient als Beispiel für die Unmöglichkeit der Koordination von Forschung. Er legt die Isolation des Virus als ein Beispiel für den freien Datenfluss zur Koordinierung von Forschung, durch Labore rund um die Welt ohne einen zentralen Kontrollpunkt aus.
- Ambivalenz: Wo Übergänge sichtbar werden und verlangsamt dargestellt werden, kann es zu einer Informationsflut kommen, welche die entscheidenden Individuen, unter der Berücksichtigung der getroffenen Wahl nicht bemerken: Vorausgesetzt dies geschieht, fällt es dem Einzelnen leichter, sein Benehmen auf die Gruppe abzustimmen, da er das Benehmen der Gruppe leicht kopieren kann. Verlust der Unabhängigkeit in der Gruppe

Auch in der menschlichen Gesellschaft trifft man Herdenverhalten an, etwa bei Restaurant- oder Theaterbesuchen, in der Mode oder beim Kauf von Bestsellern. Auf Finanzmärkten neigen Anleger manchmal dazu, sich in ihren Kauf- und Verkaufsentscheidungen wie eine Herde zu verhalten und mehrheitlich in ein Handelsobjekt zu investieren bzw. zu desinvestieren. Herdenverhalten ist eine Ausprägung massenpsychologischer Contagion-Effekte und kann somit eine Ursache für Finanzmarktkrisen oder Wirtschaftskrisen sein. Auch Hamsterkäufe zeigen Herdenverhalten wie vor Naturkatastrophen oder während der Covid-19-Pandemie ab März 2020, als es in deutschen Läden Regallücken bei bestimmten Waren (beispielsweise Mehl, Nudeln, Toilettenpapier) gab.

Einer der bedeutendsten Forscher zum Herdenverhalten ist Abhijit Banerjee, der 1992 das Herdenverhalten (herd behaviour) als gleichgerichtete Entscheidungen definierte, denen unterschiedliche private Informationen zugrunde liegen. Er implizier-

te hierbei eine asymmetrische Information, die darin besteht, dass schlecht informierte Marktteilnehmer aus dem Marktverhalten der besser informierten Marktteilnehmer eine Gewinnchance oder Verlustgefahr erblicken und deren Verhalten imitieren.

Herdenverhalten ist mithin auf den sozialen Effekt der Imitation zurückzuführen. Es liegt vor, wenn sich ein einzelner Entscheidungsträger unter Berücksichtigung des Verhaltens anderer Akteure für dasselbe Verhalten wie diese entscheidet, selbst wenn seine eigene unabhängige Entscheidung anders ausfallen würde.

Neben diesen Informationskaskaden gibt es als Ursachen noch das Reputationsmodell im Rahmen des Schönheitswettbewerbs und Netzwerkeffekte. Bereits John Maynard Keynes untersuchte 1936 das Herdenverhalten im Rahmen des Schönheitswettbewerbs (beauty contest) bei Investitionsentscheidungen. Dem liegt die Annahme zugrunde, dass Manager aus Gründen der Reputation auf Märkten mit unvollkommener Information dem Verhalten anderer Manager folgen. Dieser Schönheitswettbewerb führt dazu, dass Finanzanalysten nicht die erwartete Ertragskraft von Unternehmen prognostizieren, sondern das, was die Mehrheit aller Analysten prognostizieren wird. Etwaige Prognosefehler werden hierdurch verstärkt. Ein Netzwerkeffekt liegt vor, wenn ein bestimmter Ursache-Wirkungszusammenhang zwischen der Handlung eines Marktteilnehmers und den Handlungen der anderen besteht. Der Nutzen für den einzelnen Marktteilnehmer steigt, wenn weitere Marktteilnehmer sich für dasselbe Produkt entscheiden. Beispielsweise steigt in sozialen Netzwerken wie Facebook deren Nutzen, je mehr Benutzer diese Online-Community aufweist.

Dem Herdenverhalten können verschiedene massenpsychologische oder marktpsychologische Ursachen zugrunde liegen. Der Verbraucher kann von der Furcht getrieben sein, angesichts von Regallücken seinen Bedarf nicht decken zu können, wenn er nicht sofort kauft. Auch die Erwartung eines Verbrauchers, dass andere Verbraucher nach ihm auch hamstern werden, drängt ihn zu Hamsterkäufen. Ebenfalls seine Befürchtung, dass es künftig zu Lieferengpässen kommen könnte, zwingt ihn zu nicht bedarfsgerechten Kaufentscheidungen. Zuweilen werden auch Ohnmachtsgefühle der Verbraucher als Ursache gesehen.

Das Verhalten ist irrational, zumal Nahrungs- und Genussmittel oder Toilettenpapier als Massenprodukte jederzeit reproduzierbar sind. In Frankreich und Italien ist von Hamsterkäufen (französisch *achat de ravitaillement*, italienisch *acaparmiento*) unter anderem der Rotwein betroffen, ein nicht jederzeit reproduzierbares Produkt. In der Türkei ist das dem Kölnisch Wasser ähnliche „Kolonya“ durch Hamsterkäufe knapp geworden.

Es ist jedoch schwierig, Herdenverhalten explizit nachzuweisen; ein gemeinsamer Kauf/Verkauf eines bestimmten Wertpapiers durch viele Wirtschaftssubjekte muss nicht notwendigerweise auf Herdenverhalten (und somit auf Informationsasymmetrien) zurückzuführen sein; es kann sich auch um Zufall handeln. Wenn neue Informationen den aktuellen Preis des Papiers falsch erscheinen lassen und wenn diese Informationen zum gleichen Zeitpunkt vielen Wirtschaftssubjekten bekannt werden (vollkommene Information), können daraus viele gleichzeitige, unbeeinflusst vom Verhalten anderer Anleger getroffene Verkaufsentscheidungen resultieren. Eine spezifische Art Herdenverhalten stellt das Noise Trading dar, wenn „Noise Trader“ durch den vorhandenen Noise dazu motiviert werden, in steigende Kurse hinein zu kaufen (Hausse) oder in fallende hinein zu verkaufen (Baisse).

Die Folge von Herdenverhalten sind starke Preisschwankungen des betroffenen Handelsobjekts. Zudem beschleunigen Hamsterkäufe die Warenrotation und verringern die logistische Reichweite.

Als Marktverhalten ist das Herdenverhalten insbesondere bei Noise Tradern bekannt, die oft vom Herdenverhalten geleitet sind und durch Stimmungen oder Gruppen dazu motiviert werden, zu kaufen oder in fallende hinein zu verkaufen. Dies ist der so genannte „Stimmungs-Noise“. Steigende oder fallende Kurse sind ein Indiz dafür, dass sich bereits andere Marktteilnehmer zuvor genauso entschieden haben. Dieser Noise kann sowohl bei Kauf- als auch bei Verkaufsentscheidungen und auch bei Halte-Entscheidungen zugrunde liegen. Herdenverhalten ist somit ein Zeichen für fehlende Effizienz von Märkten.

Spekulation wird für einen Markt erst problematisch, wenn nicht mehr mit Hilfe von Fundamentaldaten spekuliert wird, sondern Herdenverhalten einsetzt. Dann können Spekulationsblasen entstehen, die meist auf Herdenverhalten zurückzuführen sind. Spekulationsblasen können durch die Erwartung der Mehrheit der Marktteilnehmer von künftigen Gewinnchancen begründet werden.

Gewinnmitnahmen können auch auf Herdenverhalten beruhen, wenn eine Vielzahl von Anlegern ein hohes Kursniveau zum Verkauf nutzt und sich weitere Anleger dem anschließen. Auch der Bank Run ist ein typisches Herdenverhalten, denn Anleger beobachten eine vielleicht zufällige Massenabhebung von Bargeld und schließen sich dieser blindlings an im Vertrauen darauf, dass diese einen bestimmten Grund haben muss; die massenhaften Abhebungen kulminieren schließlich im Dominoeffekt. Anleger ziehen ihre Einlage ab, weil sie befürchten, diese Einlage infolge des sequentiellen Auszahlungsprinzips („Wer zuerst kommt, mahlt zuerst.“) andernfalls nicht mehr abziehen zu können, da die Bargeldbestände aufgebraucht sind. Folglich ist es für jeden Einleger rational, der Herde zu folgen. Ein Bank Run ist umso eher zu erwarten, je schlechter die Bankkunden informiert sind und je mehr sie

„überreagieren“. Hamsterkäufe sind geeignet, durch stark zunehmende Nachfrage zur Knappheit bestimmter Güter oder Dienstleistungen und damit zur Marktenge beizutragen.

Herdenverhalten kann zu selbsterfüllenden Prophezeiungen führen: Verhalten sich Marktteilnehmer in einer bestimmten Weise, so kann dies dazu führen, dass sich die einer Anlage zugrundeliegenden Fundamentaldaten durch das Herdenverhalten selbst ändern: Sie entwickeln sich in die Richtung, die die Herde einschlägt – folglich ist es rational, nicht aus der Herde auszuscheren, wodurch sich letzten Endes das erwartete Ergebnis einstellt.

Massenhysterie bezeichnet eine starke emotionale Erregung in großen Menschenmengen, etwa (euphorisch) aus Anlass von Rock- und Popkonzerten, großen Sportereignissen oder (trauernd) nach dem Tod von berühmten Personen. Der Begriff ist von dem der moralischen Panik abzugrenzen, der gezielt der sozialen Kontrolle gilt.

Dieser Gebrauch geht auf den *The Quarterly Christian Spectator* 1830 zurück und wurde unter anderem bei einem Choleraausbruch gebraucht. Marshall McLuhan begann das Phänomen in *Understanding Media* 1964 wissenschaftlich zu beschreiben.

In diesem Sinne wurde und wird beispielsweise die überschießende Begeisterung für die Beatles ebenso dem Bereich der Massenhysterie zugeordnet wie die Trauer um Rudolph Valentino, Josef Stalin oder Eva Perón. Auch die mittelalterliche Tanzwut, der Hexenwahn der Frühen Neuzeit und andere massenhaft auftretende Ängste (etwa die Kommunistenangst im McCarthyismus) werden häufig als Massenhysterie bezeichnet. Der Begriff wird gelegentlich auch als gleichbedeutend mit Massenpanik benutzt. Die Sozialpsychologie beschäftigt sich unter dem Thema Massenpsychologie wissenschaftlich mit dem Verhalten von Menschen in Menschenansammlungen.

Der Begriff Massenpanik bezeichnet ein Unglück mit einer großen Zahl von Beteiligten auf engem Raum, bei dem die räumliche Beengtheit mitursächlich für den Verlauf des Unglücks ist. Er legt die Vorstellung nahe, dass eine Menschenmasse bei Großveranstaltungen oder Schadensereignissen in Panik gerät und es zu unkontrollierten Fluchtbewegungen kommt. Ursache einer Massenpanik können gefährliche äußere Umstände (wie ein Brand oder der Einsturz eines Gebäudes) oder das Verhalten Einzelner innerhalb einer Menschenmasse sein. Die Begriffe Massenunfall, Massenunglück und Massenpanik werden in den Medien häufig synonym verwendet. Eine Massenpanik tritt nur im Verlauf eines sehr geringen Anteils von Massenunglücken auf.

Starkes Gedränge oder Katastrophen mit vielen Beteiligten können eine Massenpa-

nik auslösen, die mit unkontrollierter Angst und massiven Fluchtbewegungen einhergeht. In einer solchen Situation gibt es nur wenige Interventionsmöglichkeiten. Die größten Einflussmöglichkeiten bestehen in der Entstehungsphase bzw. davor. Wichtig sind gezielte, klare, häufige, regelmäßige und strukturierte Aufforderungen und Informationen. Dies kann z. B. durch Lautsprecherdurchsagen oder durch Abläufe geschehen, die Gelassenheit demonstrieren (z. B. Fortsetzung der Veranstaltung wie eines Fußballspiels). Auch Aufmerksamkeit erweckende Interventionen (z. B. ein schriller Pfeifton) oder das Stellen einfacher Aufgaben können eine panische Menge erreichen (z. B.: Achten Sie auf Kinder!).

Es ist sicher klar geworden, dass alle diese beschriebenen Strukturen fraktalen Charakter besitzen. Es handelt sich stets um eine große Menge von Individuen, die auf einen engen Raum zusammengedrängt sind. Diese Mengen können sich sowohl positiv als auch negativ verhalten. Als Individuum ist man immer gut beraten, wenn man (wie bei der Vermeidung von Covid - Erkrankungen) einen vernünftigen Abstand zu anderen Individuen einhält. Der Abstand muss groß genug sein, um die Verbreitung von negativen Faktoren zu verhindern, er muss aber auch klein genug sein, die Verbreitung von positiven Erscheinungen zu ermöglichen.

Kehren wir noch einmal zur Tierwelt zurück. Fische, Vögel, Heuschrecken, Mücken, Ameisen, Bienen - hier findet man beliebig viele weitere fraktale, chaotische, sich selbst organisierende Strukturen vor.



Abbildung 8.7: Ameisen auf einer Melone



Abbildung 8.8: Ein Heuschreckenschwarm

Zuletzt noch ein Chaos, das sicher vielen vertraut ist.



Abbildung 8.9: Ein Autostau

8.4 Chaos in der Natur

Erdbeben sind in zweifacher Hinsicht mit Fraktalen und Chaos verbunden. Zunächst besitzen die Aufzeichnungen von Schwingungen, die durch ein Erdbeben hervorgerufen werden, fraktalen Charakter. Und dass nach einem Erdbeben chaotische Zustände herrschen können, ist aus der Literatur, aus dem Fernsehen und aus Filmen hinreichend bekannt.

Stürme treten in mehreren Abstufungen auf:



Abbildung 8.10: Die Folgen eines schweren Erdbebens in Haiti

- Ab einer Windgeschwindigkeit von rund 75 Stundenkilometer gilt ein Wind als Sturm.
- Als Orkan werden im weiteren Sinn Winde mit der Stärke 12 auf der Beaufortskala bezeichnet, im engeren Sinn werden darunter Nordatlantiktiefs verstanden, in denen solche Winde mit der Stärke 12 auftreten. Früher wurden alle Winde mit Orkanstärke als Orkane bezeichnet.
- Tornados sind die verheerendste Form von Stürmen. Ein Tornado ist ein kleinräumiger Luftwirbel in der Erdatmosphäre mit annähernd senkrechter Drehachse. Er hängt zusammen mit konvektiver Bewölkung (Cumulus und Cumulonimbus) und unterscheidet sich damit von Kleintromben (Staubteufeln). Der Wirbel erstreckt sich durchgehend vom Boden bis zur Wolkenuntergrenze, muss dabei aber nicht durchweg kondensiert sein. Diese Definition geht auf Alfred Wegener (1917) zurück und ist heute noch allgemein anerkannt.

Für die Entstehung eines Tornados müssen zunächst die Voraussetzungen für hochreichende Feuchtekongvektion gegeben sein. Diese sind bedingte Labilität, also eine hinreichend starke vertikale Temperaturabnahme, genügendes Feuchteangebot (latente Wärme) in den unteren 1–2 km der Atmosphäre sowie Hebung der Luftmasse, um die Feuchtekongvektion auszulösen. Hebungsmechanismen können thermischer (Sonneneinstrahlung) oder dynamischer (Fronten) Natur sein. Wesentlicher Energielieferant solcher Stürme und Gewitter allgemein ist die im Wasserdampf der feuchten Luftmasse gespeicherte latente Wärme, welche bei der Kondensation freigesetzt wird. Erst diese zusätzliche Wärmemenge ermöglicht ein hochreichend freies Aufsteigen der Luft (Feuchtekongvektion), da die Atmosphäre gegenüber trockener



Abbildung 8.11: Das Auge eines Orkans

Konvektion, abgesehen von bodennaher Überhitzung, stabil ist. Im letzteren Fall kann es lediglich zur Bildung von Kleintromben kommen. An der Böenfront eines Schauers oder Gewitters können Kleintromben, die sogenannten Böenfrontwirbel oder Gustnados, entstehen. Diese können sich zu Tornados entwickeln, sofern sie Kontakt zu dem feuchtkonvektiven Aufwind bekommen und so verstärkt werden.

Im Anfangsstadium ist ein Tornado zunächst fast unsichtbar. Erst wenn im Inneren des Wirbels durch den Druckabfall und die damit einhergehende adiabatische Abkühlung Wasserdampf kondensiert oder Staub, Trümmer, Wasser und dergleichen aufgewirbelt werden, tritt der Tornado auch optisch in Erscheinung. Eine durchgehende Kondensation von der Wolke bis zum Boden ist nicht in jedem Fall zu beobachten. Eine solche von der Mutterwolke ausgehende Kondensation wird als Trichterwolke (englisch: funnel cloud) bezeichnet. Erreicht der Luftwirbel den Boden nicht, so spricht man von einer Blindtrombe. Für einen Tornado ist der Bodenkontakt des Luftwirbels entscheidend, nicht dessen durchgehende Sichtbarkeit. Sind zum Beispiel unter einer Trichterwolke Windwirkungen nachweisbar, also im Regelfall Schäden

am Boden, so handelt es sich um einen Tornado. Die Gestalt des Luftwirbels ist sehr vielfältig und reicht von dünnen schlauchartigen Formen bis zu einem mehr oder weniger breiten, sich nach oben erweiternden Trichter (siehe nebenstehende Abbildungen und Weblinks). Dabei kann der Durchmesser einige Meter bis hin zu 500 m und sogar bis über 1 km betragen. Nicht selten treten – besonders bei großen Durchmessern – mehrere Wirbel auf, die um ein gemeinsames Zentrum kreisen, was als Multivortex-Tornado bezeichnet wird. Staub, Trümmer und kondensiertes Wasser können mitunter verhindern, dass ein Multivortex-Tornado als solcher erkannt wird, weil die Einzelwirbel nicht sichtbar sind.

Die Kraft eines Tornados kann vielfältige Schäden verursachen. Er kann Häuser und Autos zerstören und stellt eine Gefahr für Tiere und Menschen dar. Auch Steinhäuser sind nicht sicher. Indirekt entstehen viele Schäden durch umherfliegende Trümmer. Hauptursache der Schäden ist der Staudruck des Windes und oberhalb von circa 300 km/h auch zunehmend indirekte Schäden durch umherfliegende Trümmer. Die frühere Annahme, der starke Unterdruck innerhalb eines Tornados, der bis zu 100 hPa betragen kann, lasse Gebäude gleichsam explodieren, ist nicht mehr haltbar. Auf Grund ihrer hohen und auf engem Raum wechselnden Windgeschwindigkeiten stellen Tornados prinzipiell eine Gefahr für den Flugverkehr dar; Unfälle sind aber auf Grund der Kleinräumigkeit dieser Wettererscheinung selten. Zu einem spektakulären Fall kam es am 6. Oktober 1981, als der NLM-Cityhopper-Flug 431 in einen Tornado geriet und nach Abriss der rechten Tragfläche abstürzte. Alle 17 Personen an Bord starben.

Dauer und Geschwindigkeiten Die Dauer eines Tornados beträgt zwischen wenigen Sekunden und mehr als einer Stunde, durchschnittlich liegt sie unter zehn Minuten. Die Vorwärtsbewegung eines Tornados folgt der zugehörigen Mutterwolke und liegt im Schnitt bei 50 km/h, kann aber auch deutlich darunter (praktisch stationär, nicht selten bei Wasserhosen) oder darüber (bis über 100 km/h bei starker Höhenströmung) liegen. Dabei ist die Tornadospur im Wesentlichen linear mit kleineren Abweichungen, welche durch die Orographie und das lokale Windfeld in der Umgebung der Gewitterzelle bedingt sind.

Die interne Rotationsgeschwindigkeit des Windes ist jedoch meist wesentlich höher als die der linearen Bewegung und für die schweren Verwüstungen verantwortlich, die ein Tornado hinterlassen kann. Die höchste je registrierte Windgeschwindigkeit innerhalb eines Tornados wurde während des Oklahoma Tornado Outbreak am 3. Mai 1999 bei Bridge Creek, Oklahoma (USA) mit einem Doppler-Radar bestimmt. Mit 496 ± 33 km/h lag sie im oberen Bereich

der Klasse F5 der Fujita-Skala; die obere Fehlergrenze reicht sogar in den F6-Bereich. Dies ist damit die höchste je gemessene Windgeschwindigkeit auf der Erdoberfläche überhaupt. Oberhalb der Erdoberfläche erreichten nur Jetstreams höhere Windgeschwindigkeiten. In der offiziellen Statistik fällt dieser Tornado aber mit Rücksicht auf den wahrscheinlichsten Wert und die Unsicherheiten unter F5.

In den USA sind etwa 88 Prozent der beobachteten Tornados schwach (F0, F1), 11 Prozent stark (F2, F3) und unter 1 Prozent verheerend (F4, F5). Diese Verteilungsfunktion ist weltweit sehr ähnlich und in dieser Form von mesozyklonalen Tornados dominiert, welche das volle Intensitätsspektrum ausfüllen. Die Intensität von nicht-mesozyklonalen Tornados geht dagegen kaum über F2 hinaus.

Tornados entstehen über Land am häufigsten im Frühsommer, wobei das Maximum mit zunehmenden Breitengraden später auftritt. Über Wasser wird das Maximum im Spätsommer erreicht, weil dann die Wassertemperatur und folglich die Labilität am höchsten ist. Ähnliches gilt für den Tagesgang. Tornados über Land treten am wahrscheinlichsten in den frühen Abendstunden auf, während bei Wasserhosen das Maximum in den Morgenstunden liegt. Ferner zeigt sich bei Wasserhosen ein klimatologischer Unterschied im Jahresgang, je nachdem, ob diese an Land ziehen oder über dem Wasser verbleiben. Die jahreszeitliche Verteilung für den ersten Fall gleicht der für Tornados über Land, während reine Wasserhosen das besagte Spätsommermaximum zeigen.

Tornados werden weltweit überall da beobachtet, wo es Gewitter gibt. Schwerpunkte sind Regionen mit fruchtbaren Ebenen in den Subtropen bis in die gemäßigten Breiten. An erster Stelle steht der Häufigkeit nach der Mittlere Westen der USA, wo die klimatischen Bedingungen für die Bildung von schweren Gewittern und Superzellen aufgrund der weiten Ebenen (Great Plains) östlich eines Hochgebirges (Rocky Mountains) und nördlich eines tropischen Meeres (Golf von Mexiko) sehr günstig sind. Für Wetterlagen mit hohem Unwetterpotential bedingt das Gebirge relativ trockene und kühle Luftmassen im mittleren bis oberen Bereich der Troposphäre bei südwestlichen bis westlichen Winden, während in den tieferen Schichten feuchtwarme Luftmassen aus der Golfregion ungehindert nach Norden transportiert werden können. Dadurch kommt eine labile Schichtung der Atmosphäre bei einem großen Angebot latenter Wärme mit einer Richtungsscherung des Windes zusammen.

Weitere wichtige Regionen sind Mittel-, Süd- und Osteuropa, Argentinien, Südafrika, Bengalen, Japan und Australien. Zahlreiche, wenn auch im Mittel

schwächere, meist nicht-mesozyklonale Tornados treten im Bereich der Front Range (Ostrand der Rocky Mountains), in Florida und über den Britischen Inseln auf.

Jährlich werden in den USA etwa 1200 Tornados registriert, die meisten entstehen in Texas, Oklahoma, Kansas und Nebraska entlang der Tornado Alley mit etwa 500 bis 600 Fällen pro Jahr. Dies ist durch die oben genannten besonderen klimatischen Bedingungen gegeben, welche die Voraussetzungen für die Entstehung speziell von mesozyklonalen Tornados weit häufiger als in anderen Regionen bieten. Darüber hinaus gibt es in den USA mehrere regionale Häufungen, z. B. in Neuengland und in Zentral-Florida.

In Europa liegt die jährliche Zahl der Tornadobeobachtungen bei 330, davon 160 über Wasser, unter Einbeziehung der Dunkelziffer schätzungsweise bei 590 Tornados, davon geschätzt 290 Wasserhosen (2020: 800 gemeldete Ereignisse). Wie in den USA sind auch die meisten europäischen Tornados schwach. Verheerende Tornados sind zwar selten, doch sind bisher acht F4- und zwei F5-Ereignisse aus Deutschland dokumentiert. Letztere wurden bereits von Alfred Wegener 1917 in einer Arbeit zur Tornadoklimatologie Europas beschrieben. Weitere verheerende Fälle sind aus Nordfrankreich, den Benelux-Staaten, Österreich, Oberitalien sowie aus der Schweiz bekannt.

In Deutschland liegt die Zahl der jährlich beobachteten Tornados bei durchschnittlich 30 bis 60 mit einer noch recht hohen Dunkelziffer vor allem schwächerer Ereignisse. Genaue Zahlen sind nur auf der Tornadoliste.de verfügbar. Nach den derzeit vorliegenden Zahlen muss jährlich mit etwa fünf oder mehr F2, mit einem F3 alle zwei bis drei Jahre und alle 20 bis 30 Jahre mit einem F4 gerechnet werden. Ein F5 ist nach derzeitigen Erkenntnissen ein Jahrhundertereignis oder noch seltener. Von der Fläche gesehen hat Deutschland also so viele und auch starke Tornados wie Texas in den USA

Eine Übersicht zur räumlichen und zeitlichen Verteilung von Tornados innerhalb von Deutschland und deren Intensität findet sich in den Weblinks. Generell ist festzustellen, dass das Tornadorisiko im Westen der Norddeutschen Tiefebene am höchsten ist.

In Österreich wurden im Schnitt der vergangenen 30 Jahre jährlich etwa drei Tornados beobachtet. Allerdings ist seit 2002 durch die vermehrte Spotter- und Statistiktätigkeit v. a. ehrenamtlicher Helfer eine mittlere Anzahl von etwa fünf Tornados/Jahr zu beobachten. Unter Einbezug einer möglicherweise recht hohen Dunkelziffer sowie der nach wie vor sehr unterrepräsentierten F0-Fälle könnte die tatsächliche, gemittelte, jährliche Anzahl bei bis zu zehn

Tornados liegen.

Dabei treten jedes Jahr mehrere F0- und F1-Fälle auf. Im Schnitt kann zudem mit einem F2 jährlich, bzw. einmal in zwei Jahren, alle fünf bis zehn Jahre auch mit einem F3 gerechnet werden. Bisher ist auch ein F4-Tornado in Österreich dokumentiert.

Die höchste Tornadodichte ist dabei in der Südoststeiermark zu beobachten (um drei Tornados/ 10.000 km^2 / Jahr), gefolgt von dem Gebiet um den Hausruck in Oberösterreich, dem Wiener Becken, der Region um Linz, dem westlichen Weinviertel, dem Klagenfurter Becken, Bodensee-Region sowie dem Inntal im Bereich von Innsbruck.



Abbildung 8.12: Ein Tornado in Oklahoma

Generell ist das Auftreten von Tornados starken Schwankungen unterworfen, was sich in Häufungen innerhalb recht kurzer Zeitspannen äußert, gefolgt von recht langen Abschnitten relativer Ruhe. Die Ausbrüche sind durch den engen Zusammenhang mit bestimmten Wetterlagen begründet, wo mehrere Faktoren für die Tornadoentstehung zusammenkommen (siehe oben unter Entstehung). Größere Ereignisse dieser Art mit verheerenden Tornados sind vor allem aus den USA bekannt. Für West- und Mitteleuropa sind hier die Jahre 1925, 1927 und 1967 zu nennen mit dem Schwerpunkt Nordfrankreich/Benelux/Nordwestdeutschland. Diese Region kann auch als europäische tornado alley angesehen werden. Der zahlenmäßig bedeutendste Ausbruch in Europa mit insgesamt 105, aber meist schwächeren Tornados (max. F2) traf am 23. November 1981 die Britischen Inseln. Derzeit erlaubt die Datenbasis für Mitteleuropa keine Aussage, ob Tornados auf Grund der globalen Klimaerwärmung häufiger auftreten, da der Anstieg der beobachteten Fälle vor allem auf eine bessere Erfassung in den letzten Jahren zurückzuführen ist.

- Überschwemmungen, Strudel und Turbulenzen

Als Überschwemmung bezeichnet man einen Zustand, bei dem eine normalerweise trockenliegende Bodenfläche vollständig von Wasser bedeckt ist. Flutkatastrophen waren neben besonders starken Erdbeben bislang die für Menschen folgenreichsten Naturkatastrophen. Überschwemmungen sind meistens Naturereignisse.

Sturmfluten schieben Meerwasser über Flussmündungen dort, wo es keine Sperrwerke gibt, tief ins Binnenland hinein sowie über Deiche an der Küste und an küstennahen Flussabschnitten hinweg oder bewirken Deichbrüche; Tsunamis können auch hoch gelegene, nicht durch Deiche geschützte Küstenabschnitte überfluten.

Binnengewässer treten bei Hochwasser über die Ufer, wenn das Wasser nicht zügig genug nach unten fließt (als Oberflächenwasser talwärts bzw. auf nicht versiegelten Böden ins Grundwasser versickernd). Das Hochwasser kann durch ergiebige Niederschläge, aber auch durch den Bruch von Dämmen oder Stau-
mauern infolge zu hohen Wasserdrucks entstanden sein. Gletscher hindern nachfließendes Wasser am Abfließen und bilden auf diese Weise einen Eisstausee. Sich füllende Grundwasserreservoirs führen zu steigenden Wasserständen, wenn das Grundwasser über wasserundurchlässigen Bodenschichten liegt. Es findet dann von unten oder (unterhalb der Erdoberfläche) von den Seiten her seinen Weg in tiefliegende Teile von Bauwerken. Dies betrifft vor allem Gebiete mit einem ständig hohen Grundwasserpegel. Der mangelhafte, teils auch ausbleibende Abfluss großer Wassermassen im Binnenland und auf Inseln kann (zumindest teilweise) von Menschen ausgelöst sein:

Die Versiegelung großer Flächen erschwert das Versickern von Oberflächenwasser, das ungehemmt in Bäche und Flüsse gelangt. Allerdings bewirkt Starkregen selbst eine Versiegelung tiefer gelegener Flächen, da er die Fähigkeit zur Aufnahme von Wasser auch in an sich nicht versiegelten Böden ab einem bestimmten Zeitpunkt unterbindet. Die Begradigung und Einengung von Flüssen führt zu einer Erhöhung der Fließgeschwindigkeit des Wassers, das schnell tiefer gelegene Gebiete erreicht und dort auf langsamer fließendes Wasser an Engstellen aufläuft. Es gibt vor allem in Ballungsräumen zu wenige Flächen, in denen sich das Wasser ausbreiten kann. Dies führt zwangsläufig zu Hochwasser bei starkem Wasserzustrom durch heftige Niederschläge und/oder in großen Mengen nachrückendes Oberflächenwasser.

Zusammengebrochene Konstruktionen oder Treibgut können sich an Brücken, Sperrwerken, Rechen oder Überläufen, bzw. Ablässen verkanten und dadurch einen Wasserstau auslösen (Verklausung); Auch in großer Entfernung zu größeren Flüssen kann es zu (meist lokal begrenzten) Überschwemmungen kom-

men. Zumeist sorgen hier Starkregenfälle dafür, dass die Kanalisation und Entwässerungsgräben die Wassermassen nicht bewältigen können. Zur Zwischenlagerung von Wassermassen geschaffene Einrichtungen wie Regenrückhaltebecken oder geplante Überschwemmungsflächen erweisen sich oft als unterdimensioniert. In den genannten Fällen ist ein Hauptgrund für den Eintritt von Überschwemmungen, dass die Eintrittswahrscheinlichkeit und die Wirkung von Starkregenereignissen unterschätzt werden. Obwohl Wasserfluten eine Naturgewalt darstellen, handelt es sich bei einer Überschwemmung dann nicht um ein Naturereignis, wenn sie durch einen Wasserrohrbruch ausgelöst wurde oder wenn Gebiete (insbesondere im Zusammenhang mit Kampfhandlungen) absichtlich unter Wasser gesetzt werden.

Überschwemmungen können unter Umständen erhebliche Wasserschäden am Eigentum von Menschen hervorrufen sowie die Gesundheit und das Leben von Menschen und Nutztieren gefährden. Besteht eine solche Gefahr, so sprechen Rettungskräfte von einem Wassernotstand. Zu unterscheiden sind temporäre Überschwemmungen, die durch das Abfließen oder Hochpumpen des eingedrungenen Wassers enden, von dauerhaften Überschwemmungen. Letztere drohen insbesondere tiefliegenden Küstengebieten infolge des klimabedingten Anstiegs des Meeresspiegels.

Bei einer Überschwemmung in bergigem Gelände können Schäden auch unterhalb des eigentlichen Wasserspiegels ausgelöst werden, indem Grundstücke in Hanglage unterspült werden. Durch den Materialabtrag kann der Hang oberhalb des Wasserspiegels abrutschen und Gebäude oder deren Inhalt (z. B. ein Auto in einer Garage) nach unten stürzen.

Nicht jede Überschwemmung stellt jedoch für Menschen ein Problem dar. So gäbe es z. B. in Wüstengebieten ohne das regelmäßige Über-die-Ufer-Treten von Flüssen wie dem Nil abseits von Oasen kein fruchtbares Land für den Acker- und Gartenbau, und unter ökologischen Aspekten werden naturbelassene Feuchtgebiete, die regelmäßig überschwemmt werden, als überaus wertvoll bewertet.

Auch in Deutschland gibt es Kulturlandschaften, deren Bewohner ihren Wohlstand der Fruchtbarkeit ihres Landes verdanken, welche wiederum durch Sedimente zu erklären ist, die nach Überschwemmungen auf dem Boden zurückgelassen wurden. Dies trifft beispielsweise auf das Artland zu, welches größtenteils im Hase-Binnendelta liegt.

Strudel sind meist harmlos.



Abbildung 8.13: Eine Überschwemmung in Japan



Abbildung 8.14: Kleine Strudel sind hübsch und harmlos

Die turbulente Strömung ist die Bewegung von Fluiden, bei der Verwirbelungen in einem weiten Bereich von Größenskalen auftreten. Diese Strömungsform ist gekennzeichnet durch ein dreidimensionales Strömungsfeld mit einer zeitlich und räumlich scheinbar zufällig variierenden Komponente. Den räumlichen Aspekt verdeutlicht nebenstehendes Bild, den zeitlichen z. B. das Rauschen des Windes.

Turbulenz führt zu verstärkter Durchmischung und infolge zu effektiv erhöhten Diffusionskoeffizienten. Bei großräumiger Turbulenz ist der Beitrag der molekularen Diffusion vernachlässigbar. Die Vermischung betrifft auch die innere Energie (Wärmetransport) und den Impuls.

Der Druckverlust eines durch ein Rohr strömenden Fluids beruht auf der Diffusion des Impulses zur Rohrwandung und ist bei turbulenter Strömung größer als bei laminarer Strömung. Die Verwirbelung entsteht durch den Geschwindigkeitsunterschied der Strömung in Rohrmitte gegenüber der Strömung nahe der Wandung. Mit steigendem Durchfluss nimmt die Intensität der Turbulenz zu und der Druckverlust erhöht sich annähernd mit der zweiten

Potenz.



Abbildung 8.15: Turbulenzen sind selten gefährlich

Turbulente Strömungen sind durch folgende Eigenschaften gekennzeichnet:

1. ausgeprägte Selbstähnlichkeit bei Mittelwertbildung bezüglich Länge und Zeit, mit großer Ausdehnung zulässiger Längen- und Zeitskalen, ungeordnete und schwer vorhersagbare raumzeitliche Struktur. Ein Wirbelsturm ist mehrere Kilometer groß, während die kleinsten in ihm enthaltenen Wirbel kleiner als einen Millimeter sind.
2. empfindliche Abhängigkeit von Anfangsbedingungen. Die Windgeschwindigkeit nahe der Erdoberfläche schwankt sehr stark und ist schwer vorhersagbar, wenn topografische Unregelmäßigkeiten Turbulenzen hervorrufen. Vor der Errichtung einer Windkraftanlage werden daher in der Regel Windgutachten auf Grundlage lokaler Messungen erstellt.
3. empfindliche Abhängigkeit von Randbedingungen. Die turbulente Strömung zwischen Strahl und Wand im Wandstrahl-Kippglied gibt den Impuls in den langsamen Teil der Grenzschicht weiter und lässt ihn abströmen, was den Strahl an der Wand hält (Coanda-Effekt).
4. empfindliche Abhängigkeit von Randbedingungen. Schnee auf einer Tragfläche dämpft die großräumige Turbulenz in der Ablöseblase und führt zu Strömungsabriss schon bei geringeren Anstellwinkeln. Riblets auf Oberflächen können in turbulenter Strömung den Reibungswiderstand verringern, ebenso die Dimples genannten kleinen Vertiefungen auf der Oberfläche von Golfbällen.

Turbulenz kann folgendermaßen definiert werden:

1. Zufälligkeit des Strömungszustandes: nicht vorhersagbare Richtung und Geschwindigkeit.
2. Diffusivität: starke und schnelle Durchmischung, im Gegensatz zur langsameren molekularen Diffusion.
3. Dissipation: kinetische Energie wird auf allen Skalen fortgesetzt in Wärme umgewandelt und teilt sich aus den Skalen größerer Ausdehnung in hierarchischer Weise in kleinere Elemente auf. Turbulenter Fluss bleibt also nur erhalten, wenn von außen Energie zugeführt wird.
4. Nichtlinearität: der laminare Fluss wird instabil, wenn die Nichtlinearitäten an Einfluss gewinnen. Mit zunehmender Nichtlinearität kann eine Sequenz verschiedener Instabilitäten auftreten, bevor sich volle Turbulenz ausbildet.
Entstehung

Die Lineare Stabilitätstheorie beschäftigt sich mit dem Umschlag laminarer Strömungen in turbulente Strömungen. Sie betrachtet dazu das Anwachsen wellenförmiger Störungen mit kleiner Amplitude, d. h. das Anwachsen der Tollmien-Schlichting-Wellen aufgrund der Kelvin-Helmholtz-Instabilität.

Turbulenz kann auch durch spezielle Formgebung bewirkt werden, etwa in statischen Mischern oder durch die Dimples genannten kleinen Vertiefungen auf der Oberfläche von Golfbällen.

Die Ungefährlichkeit von Turbulenzen verschwindet natürlich, wenn man die Natur verlässt, ganz im Gegenteil: Turbulenzen auf den Finanzmärkten oder in der Politik können oft zu Katastrophen führen.

8.5 Attraktion

Ein Attraktor verbindet das Chaos mit der Ordnung. Er ist ein Punkt im Phasenraum, der auf das System anziehend wirkt. Beispielsweise besitzt ein Pendel genau einen Attraktor, nämlich den Ruhepunkt.

Das Pendel wird nach rechts oder links um einen Winkel α bis zu einem Gipfelpunkt angehoben. Wenn es dort losgelassen wird, bewegt es sich auf einem Kreisbogen nach unten und wird immer schneller. Die höchste Geschwindigkeit erreicht es im Nullpunkt, dann bewegt es sich nach oben bis zu einem Winkel, der wegen der

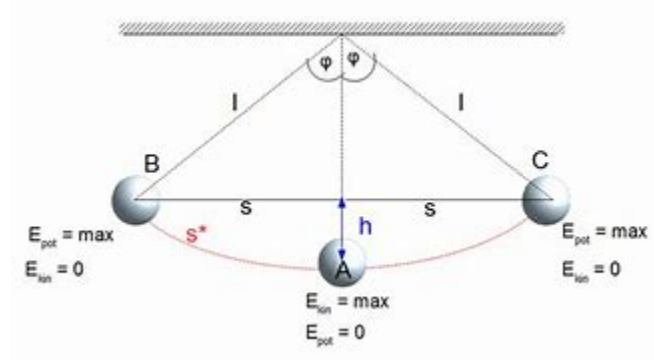


Abbildung 8.16: Darstellung der Pendelbewegung

Reibung etwas kleiner als α ist. So geht es eine Weile hin und her, der Winkel α wird immer kleiner, und schließlich kommt das Pendel bei $\alpha = 0^\circ$ zur Ruhe. Bei einer Pendeluhr kompensiert man den Reibungswiderstand durch eine Feder, die ab und an aufgezogen werden muss.

Bei der höchsten Auslenkung links und rechts ist der Impuls (das Produkt von Masse und Beschleunigung) gleich null. Das trifft auch auf den maximalen Ausschlag rechts zu. Im Ruhezustand ist die Auslenkung gleich null, aber der Impuls ist maximal.

Ein seltsamer Attraktor ist ein Attraktor, also ein Ort im Phasenraum, der den Endzustand eines dynamischen Prozesses darstellt, dessen fraktale Dimension nicht ganzzahlig und dessen Kolmogorov-Entropie echt positiv ist. Es handelt sich damit um ein Fraktal, das nicht in geschlossener Form geometrisch beschrieben werden kann. Gelegentlich wird auch der Begriff chaotischer Attraktor bevorzugt, da die „Seltsamkeit“ dieses Objekts sich mit den Mitteln der Chaostheorie erklären lässt. Der dynamische Prozess zeigt ein aperiodisches Verhalten.

Von einem seltsamen Attraktor spricht man, wenn folgende Bedingungen erfüllt sind:

- Chaotisches Verhalten: Beliebig kleine Änderungen des Anfangszustandes führen zu völlig unterschiedlichen Verläufen.
- Fraktale Struktur: Der Attraktor besitzt eine nicht-ganzzahlige Dimension.
- Keine Aufteilungsmöglichkeit: Jede Bahn, die im Einzugsbereich startet, nähert sich beliebig stark an jeden Punkt des Attraktors an.

An einigen klassischen Beispielen lassen sich die Eigenschaften seltsamer Attrakto-

ren recht gut studieren. Das verwendete dynamische System kann diskret oder kontinuierlich sein. Kontinuierliche Systeme werden meist durch Differentialgleichungen beschrieben, im Phasenraum bilden diese Systeme eine vom Ausgangszustand ausgehende Linie, die Trajektorie. Wegen der Eindeutigkeit der Ableitung können sich Trajektorien in keinem Punkt schneiden. Daraus lässt sich folgern, dass Attraktoren in zwei Dimensionen nur eine einfache Struktur aufweisen können; seltsame Attraktoren und damit chaotische Systeme gibt es in kontinuierlichen dynamischen Systemen erst bei einem Phasenraum mit mindestens drei Dimensionen.

Ein relativ einfaches Beispiel für einen seltsamen Attraktor ist der Hénon-Attraktor (benannt nach Michel Hénon), der als diskretes System im zweidimensionalen Raum durch folgende Gleichungen definiert ist:

$$x_{k+1} = y_k + 1 - a \cdot x_k^2 \quad (8.5)$$

$$y_{k+1} = b \cdot x_k. \quad (8.6)$$

Jeder Abbildungsschritt lässt sich in drei Teilschritte zerlegen: Eine Falt-Streck-Operation durch Addition von $1 - a \cdot x_k^2$ zu y , eine Kontraktion um den Faktor b entlang der x -Achse und eine Spiegelung unter Vertauschung der x - und y -Achse. Bei der Wahl der Parameter $a = 1,4$ und $b = 0,3$ ergibt sich das typische Bild des Hénon-Attraktors, wenn man die Bahn eines Ausgangspunktes verfolgt, der hinreichend nahe am Attraktor liegt. Oberflächlich sieht der Attraktor aus wie eine Linie, die in einige wenige Falten gelegt ist, da er aber invariant zu der beschriebenen Falt-Streck-Stauch-Spiegel-Operation ist, ist jede einzelne Linie wieder unendlich oft geteilt in annähernd parallel laufende Linien. Ein Querschnitt durch den Attraktor hat die Struktur einer Cantor-Menge, der Attraktor hat damit eine fraktale Struktur.

Betrachtet man die Umgebung eines Punktes auf dem Attraktor, d. h. eine Kreisscheibe mit kleinem Durchmesser, so wird diese durch einen Abbildungsschritt in eine langgezogene Ellipse überführt, die entlang der Linien des Attraktors gestreckt ist, durch den Kontraktionsschritt aber eine geringere Fläche besitzt. Durch fortgesetzte Anwendung der Abbildungsvorschrift überdeckt das Abbild der Punktumgebung immer größere Bereiche des Attraktors, während seine Fläche gegen Null geht.

Der De Jong – Attraktor wird durch die folgenden Gleichungen definiert:

$$x_{t+1} = \sin(a \cdot y \cdot t) - \cos(b \cdot x \cdot t) \quad (8.7)$$

$$y_{t+1} = \sin(c \cdot x \cdot t) - \cos(d \cdot y \cdot t) \quad (8.8)$$

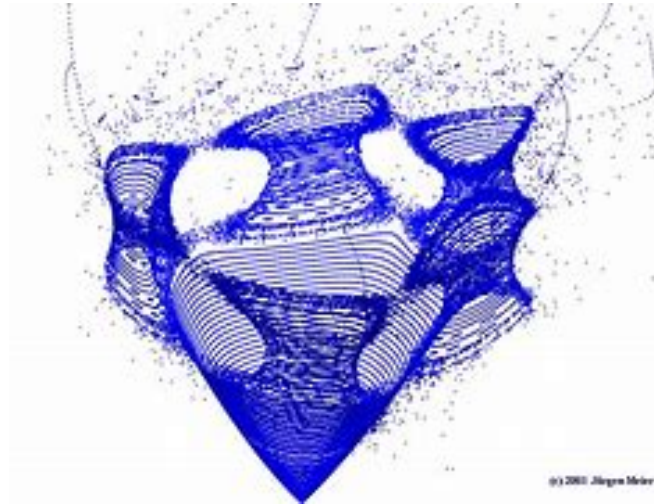


Abbildung 8.17: Ein quadratischer Hénon-Attraktor

wobei a , b , c und d die Anfangsbedingungen sind.

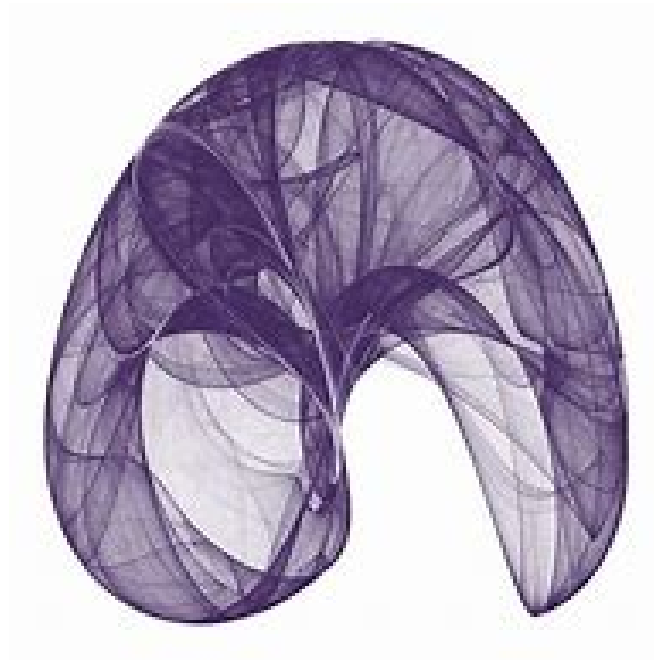


Abbildung 8.18: Ein De Jong - Attraktor

Auch hier können die Elemente, die diesen Gleichungen folgen müssen, ein bestimmtes Gebiet nicht verlassen.

Die Bäcker-Abbildung beschreibt ein einfaches, stark chaotisches System, für das alle wichtigen charakteristischen Größen exakt berechnet werden können. Diese auf E. Hopf zurückgehende iterierte Abbildung ist auf dem Einheitsquadrat definiert durch

$$y_{n+1} = \begin{cases} a \cdot y_n & \text{für } 0 \leq x_n < \frac{1}{2} \\ \frac{1}{2} + a \cdot y_n & \text{für } \frac{1}{2} \leq x_n < 1 \end{cases}$$

mit $0 < a \leq \frac{1}{2}$.

Für $a = \frac{1}{2}$ ist die Abbildung flächenerhaltend, ansonsten dissipativ.

Gewaltige Attraktoren verbergen sich hinter dem Begriff „Schwarzes Loch“. Ein Schwarzes Loch ist ein Objekt, dessen Masse auf ein extrem kleines Volumen konzentriert ist und infolge dieser Kompaktheit in seiner unmittelbaren Umgebung eine so starke Gravitation erzeugt, dass nicht einmal das Licht diesen Bereich verlassen oder durchlaufen kann. Die äußere Grenze dieses Bereiches wird Ereignishorizont genannt. Nichts kann einen Ereignishorizont von innen nach außen überschreiten – keine Information, keine Strahlung und schon gar keine Materie. Dass ein „Weg nach außen“ nicht einmal mehr denkbar ist, beschreibt die allgemeine Relativitätstheorie schlüssig durch eine extreme Krümmung der Raumzeit.

Es gibt unterschiedliche Klassen von Schwarzen Löchern mit ihren jeweiligen Entstehungsmechanismen. Am einfachsten zu verstehen sind stellare Schwarze Löcher, die entstehen, wenn ein Stern einer bestimmten Größe seinen gesamten nuklearen Brennstoff verbraucht hat und kollabiert. Während die äußeren Hüllen dann in einer Supernova abgestoßen werden, fällt der Kern durch seinen Schweredruck zu einem extrem kompakten Körper zusammen. Für ein hypothetisches Schwarzes Loch von der Masse der Sonne hätte der Ereignishorizont einen Durchmesser von nur etwa sechs Kilometern, das entspricht dem 230.000-sten Teil des jetzigen Sonnendurchmessers. Am anderen Ende des Spektrums gibt es supermassereiche Schwarze Löcher von millionen- bis milliardenfacher Sonnenmasse, die im Zentrum von Galaxien stehen und eine wichtige Rolle in deren Entwicklung spielen.

Außerhalb des Ereignishorizonts verhält sich ein Schwarzes Loch wie ein normaler Massenkörper und kann von anderen Himmelskörpern auf stabilen Bahnen umrundet werden. Der Ereignishorizont erscheint von außen visuell als vollkommen schwarzes und undurchsichtiges Objekt, in dessen Nähe der dahinterliegende Raum wie durch eine optische Linse verzerrt abgebildet wird.

Die Bezeichnung Schwarzes Loch wurde im Jahr 1967 durch John Archibald Wheeler geprägt. Zu jener Zeit galt die Existenz der erst theoretisch beschriebenen Schwarzen Löcher zwar als sehr wahrscheinlich, war aber noch nicht durch Beobachtungen bestätigt. Später wurden zahlreiche Beispiele für Auswirkungen Schwarzer Löcher beobachtet, z. B. ab 1992 die Untersuchungen des supermassereichen Schwarzen Lochs Sagittarius A* im Zentrum der Milchstraße im Infrarotbereich. 2016 wurde die Fusion zweier Schwarzer Löcher über die dabei erzeugten Gravitationswellen durch LIGO beobachtet und 2019 gelang eine radioteleskopische Aufnahme eines Bildes des supermassereichen Schwarzen Lochs M87* im Zentrum

der Galaxie M87.



Abbildung 8.19: Ein Schwarzes Loch

Ein typisches Attraktionsverhältnis findet man im Verhältnis zwischen Jäger und Beute. Als Beispiel soll das Verhältnis zwischen Gepard und Antilope dienen. Man stellt zu einem bestimmten Zeitpunkt fest, dass in einem bestimmten Gebiet wenig Geparden und viel Antilopen vorhanden sind. Auf Grund der guten Futterlage vermehren sich die Geparden sehr schnell, ihr Wert erreicht ein Maximum, die Anzahl der Antilopen ein Minimum. Zu diesem Zeitpunkt beginnt die Abnahme der Anzahl der Geparden, da das Futter nicht für alle ausreicht, und die Zahl der Antilopen kann wieder zunehmen, da weniger Feinde vorhanden sind. Die Zahl der Antilopen und der Geparden schwingt zeitlich versetzt zwischen einem Minimum und einem Maximum hin und her. Man kann dieses Verhältnis durch zwei zeitlich versetzte Sinusschwingungen darstellen.

Ein Torus-Attraktor entsteht durch Zusammenführung zweier Teilsysteme mit insgesamt drei Variablen in ein Gesamtsystem.

- Ein Punkt (des einen Teilsystems) fährt im Raum auf einer Ellipsenbahn umher und um diesen Punkt kreist zu jeder Zeit ein weiterer Punkt (des anderen Teilsystems)-
- Durch die Kopplung wird der weitere Punkt der Punkt des Gesamt-Grenzzykklus und liegt auf einem Torus.
- Hierbei herrscht asymptotische Vorhersagbarkeit vor, da sich das System auf der Oberfläche des Torus befindet und nicht willkürlich im Phasenraum.

Die Störung einer geordneten Umwelt durch Turbulenzen führt ebenfalls meistens zu chaotischen Verhältnissen, die die betroffenen Regionen und Menschen vor große

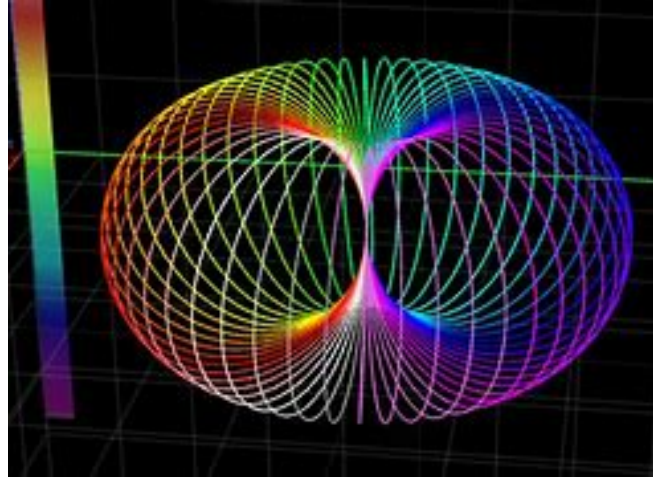


Abbildung 8.20: Darstellung eines Torus-Attraktors

Probleme stellt. Hierzu gehören mindestens

- geologische Turbulenzen (Vulkanausbrüche, Erdbeben,...)
- Turbulenzen im Wasser (das Umspülen von Pfeilern und Felsen, Tsunamis, Ebbe und Flut, Überschwemmungen, ...)
- Turbulenzen in der Atmosphäre (Gewitter, Regenschauer, ...)

Vom chaotischen Verhalten turbulenter Strömungen sind die Menschen seit Jahrhunderten fasziniert. Leonardo da Vinci hielt das Phänomen in einer bekannten Zeichnung fest. Die gezielte Untersuchung kam erst viel später. In Laborexperimenten wird die Entstehung der Turbulenz erst seit gut hundert Jahren untersucht, beispielsweise in geometrisch einfachen Umgebungen wie Rohren.

Eines dieser Experimente – ein besonders berühmtes – ist bis heute an der Universität Manchester zu besichtigen: Osborne Reynolds untersuchte dort im Jahr 1883, wann die Wasserströmung in einer geraden Glasröhre mit kreisförmigem Querschnitt – einem idealisierten Heizungsrohr – turbulent wird. Um die Strömung sichtbar zu machen, färbte er einen Teil des Wassers ein und beobachtete anschließend, ob und wann das gefärbte Wasser verwirbelte. Und in der Tat: Ab einer gewissen Strömungsgeschwindigkeit konnte er Turbulenz beobachten.

Die Erforschung der Turbulenz war fruchtbar für die Mathematik: In der zweiten Hälfte des 20. Jahrhunderts wurde dadurch die Entwicklung der Theorie dynamischer Systeme („Chaostheorie“) angeregt: Die beteiligten Forscher wollten unter anderem chaotische turbulente Strömungen verstehen. Ihre mathematischen Untersuchungen lieferten mehrere bahnbrechende Resultate:

Überrascht stellten die Wissenschaftler fest, dass in manchen simplen Modellsystemen praktisch keine Vorhersage über den zukünftigen Zustand möglich ist – und dies, obwohl sie die Gesetze, nach denen sich diese Systeme zeitlich entwickelten, genau kannten. Die Systeme reagierten einfach zu empfindlich auf Störungen. Man musste eine Art mathematischen Bankrott erklären: Die Gleichungen, die das ins Chaos driftende System beschrieben, ließen sich nicht allgemein lösen. Doch es gab bald einen Hoffnungsschimmer. Nicht die gesamte Mathematik dieser Systeme war unlösbar. Man fand spezielle, einfache Lösungen der Gleichungen, die typische und durchschnittliche Zustände repräsentierten – zum Beispiel typische Strömungsmuster. Diese speziellen Lösungen lassen sich einfacher untersuchen. Dies sind vor allem Fixpunkte, also unveränderliche Lösungen, oder periodische Orbits, bei denen das System nach der Periodendauer zum gleichen Zustand zurückkehrt. Diese Lösungen und ihre Eigenschaften verraten einiges darüber, wie Chaos genau entsteht, und es konnten sogar allgemeine Mechanismen der Entstehung identifiziert werden.

Sehr vertraut und wünschenswert sind Attraktoren des Alltags: jeder Supermarkt, jedes Modegeschäft, jede Tankstelle, jedes Ärztehaus wirkt als Attraktor und zieht die Kunden oder Patienten aus einer gewissen Umgebung an. Die Anstrengungen der Werbung sind darauf gerichtet, solche Attraktoren zu erzeugen.

8.6 Bifurkation

Ein dynamisches System kann durch eine Funktion $F(x)$ beschrieben werden, die die zeitliche Entwicklung des Systemzustands x bestimmt. Diese Funktion sei nun von einem Parameter μ abhängig, was man durch die Schreibweise $F(x, \mu)$ ausdrückt. Wenn nun das System für Parameterwerte unterhalb eines bestimmten kritischen Werts μ_c ein qualitativ anderes Verhalten aufweist als für Werte oberhalb von μ_c , dann spricht man davon, dass das System bei μ_c eine Bifurkation im Parameter μ erfährt. Der Parameterwert μ_c wird dann als Bifurkationspunkt bezeichnet.

Was eine „qualitative Änderung“ ist, kann man formal mit dem Begriff der topologischen Äquivalenz bzw. der topologischen Konjugation beschreiben: Solange für zwei Parameterwerte μ_1 und μ_2 die Systeme $F(x, \mu_1)$ und $F(x, \mu_2)$ zueinander topologisch äquivalent sind, liegt keine qualitative Änderung im obigen Sinne vor. Die Änderung am Bifurkationspunkt besteht in den meisten Fällen entweder in einer Änderung der Anzahl von Attraktoren wie Fixpunkte oder periodischen Orbits, oder eine Änderung der Stabilität dieser Objekte.

Bifurkationen lassen sich in Bifurkationsdiagrammen graphisch darstellen. Bei ei-

nem eindimensionalen System werden dabei die Fixpunkte des Systems gegen den Parameter μ aufgetragen. Für jeden Parameterwert wird so die Anzahl und die Lage dieser Punkte angezeigt. Zusätzlich kann man stabile und instabile Fixpunkte z. B. durch verschiedene Färbung unterscheiden. Bei einem System mit mehreren Variablen kann man ähnliche Diagramme zeichnen, indem man nur einen Unterraum des Phasenraums betrachtet, etwa durch einen Poincaré-Schnitt.

Das bekannteste Bifurkationsdiagramm ist das Feigenbaumdiagramm, das sich aus einer logistischen Gleichung ableitet und eine Periodenverdoppelungs-Bifurkation abbildet. Man erkennt, dass bei kleinen Parameterwerten nur ein stabiler Fixpunkt existiert, der am ersten Bifurkationspunkt in einen Orbit aus zwei alternierenden Häufungspunkten übergeht. Dieser Orbit verdoppelt dann an weiteren Bifurkationspunkten jedes Mal wieder seine Periode (kommt also erst nach 2, 4, 8 etc. Durchläufen wieder an den gleichen Punkt), bis er bei einem Parameterwert von etwa 3,57 in einen chaotischen Zustand übergeht, wo überhaupt keine Periode mehr erkennbar ist. All diese Übergänge lassen sich mithilfe des Bifurkationsdiagrammes gut veranschaulichen.

Ein typisches Beispiel einer Bifurkation ist das Knicken eines Stabes unter Druckbelastung. Man stelle sich einen im Boden eingespannten, senkrecht stehenden, masselosen Stab mit einer Last mit dem Gewicht μ an der Spitze vor. Die Winkelabweichung des Stabes aus der Senkrechten entspricht dabei der Variablen x .

Solange das Gewicht klein genug bleibt, ist $x = 0$ eine stabile Gleichgewichtslage des Systems, d. h. für kleine Abweichungen richtet sich der Stab selbständig wieder in die Senkrechte ($x = 0$) aus. Wird das Gewicht μ kontinuierlich gesteigert, so wird bei einem bestimmten Gewicht (der Knicklast oder auch Verzweigungslast) die senkrechte Gleichgewichtslage instabil. Gleichzeitig entstehen (für ein ebenes System) zwei neue (stabile) Gleichgewichtslagen (indem der Stab nach links oder rechts abknickt.) Der Übergang des Systems von einer (stabilen) zu drei (einer instabilen, zwei stabilen) Gleichgewichtslagen ist die Bifurkation, die in diesem Fall eine Pitchfork-Bifurkation ist.

Man unterscheidet die folgenden Typen:

- die Pitchfork-Bifurkation
- die Saddle-Node-Bifurkation
- die Hopf-Bifurkation
- die Transkritische Bifurkation

- die Flip-Bifurkation

Eine ganz regelmäßige Bifurkation tritt auf, wenn man einen Stammbaum einer Person von unten nach oben betrachtet: jede Person besitzt zwei Eltern, vier Großeltern, acht Urgroßeltern usw. Diese Struktur wird schnell fraktal, wenn man Geschwister, Stiefeltern, Stiefgeschwister usw. betrachtet. Außerdem gehen die Verzweigungen in unterschiedlicher Höhe zu Ende, weil man oft keine weiteren Informationen hat.

Die Pitchfork-Bifurkation, auch Heugabel- oder Stimmgabel-Bifurkation genannt, ist ein bestimmter Typ einer Bifurkation eines nichtlinearen Systems.

Die Normalform der Pitchfork-Bifurkation ist:

$$\frac{dx}{dt} = r \cdot x + \alpha \cdot x^3 = r \cdot x + \alpha \cdot x^3 \quad (8.9)$$

mit $\alpha = \pm 1$, wobei r der für das Auftreten der Bifurkation zu variierende Parameter ist.

Für $\alpha = 1$ erhält man die subkritische Pitchfork-Bifurkation, für $\alpha = -1$ erhält man die superkritische Pitchfork-Bifurkation

Die Pitchfork-Bifurkation hat folgende Gleichgewichtspunkte:

$$x_1^* = 0 \quad (8.10)$$

$$x_2^* = \pm \sqrt{-\frac{r}{\alpha}} \quad (8.11)$$

Die Sattel-Knoten-Bifurkation ist ein weiterer Typ einer Bifurkation eines nichtlinearen dynamischen Systems. Ihre Normalform lautet

$$\frac{dx}{dt} = \mu - x(t)^2. \quad (8.12)$$

Diese Normalform hat für $\mu \geq 0$ Fixpunkte $x_{1,2} = \pm\sqrt{\mu}$. Das bedeutet, es existiert für $\mu < 0$ kein Fixpunkt, für $\mu = 0$ genau ein Fixpunkt und sonst zwei. Der erste Fixpunkt ist stabil (Knoten), der zweite instabil (Sattel). Am Bifurkationspunkt $\mu = 0$ kollidieren Sattel und Knoten. Betrachtet man ein System mit höherer Ordnung in

$$\frac{dx}{dt} = \mu - x(t)^2 + O(x^3), \quad (8.13)$$

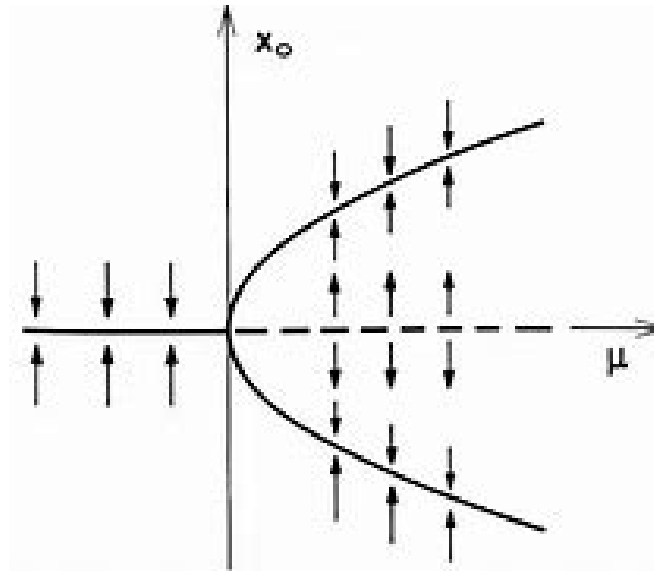


Abbildung 8.21: Die Stimmgabel-Bifurkation

so beeinflussen diese Terme in einer genügend kleinen Umgebung um den Sattel-Knoten-Punkt $\mu = 0$ das Verhalten des Systems nicht.

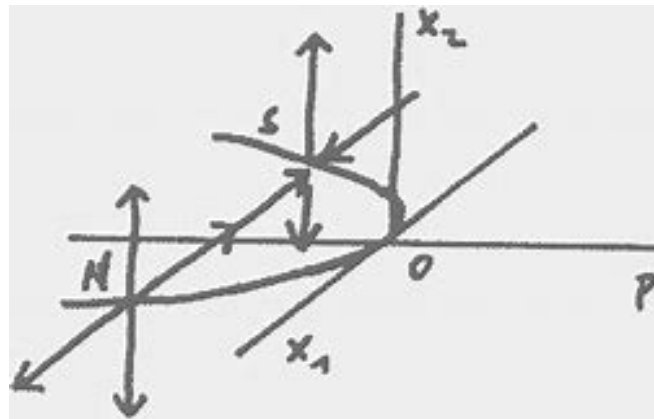


Abbildung 8.22: Die Sattel-Bifurkation

Eine Hopf-Bifurkation oder Hopf-Andronov-Bifurkation ist ein Typ einer lokalen Bifurkation in nichtlinearen Systemen. Sie ist benannt nach dem deutsch-amerikanischen Mathematiker Eberhard Frederick Ferdinand Hopf bzw. nach Alexander Alexandrowitsch Andronow, der sie mit Witt und Chaikin in der Sowjetunion in den 1930er Jahren behandelte. Die Wurzeln der Theorie gehen aber auf Henri Poincaré Ende des 19. Jahrhunderts zurück.

Bei einer Hopf-Bifurkation überquert an einem Gleichgewichtspunkt (Fixpunkt) des Systems ein Paar komplex konjugierter Eigenwerte der aus der Linearisierung des Systems resultierenden Jacobimatrix die imaginäre Achse der komplexen Ebene; am Bifurkationspunkt selbst sind die konjugierten Eigenwerte also rein imaginär. Die Hopf-Bifurkationen können nur in zwei- oder höherdimensionalen Systemen

auftreten, da die Linearisierung des Systems mindestens zwei Eigenwerte besitzen muss.

Die Normalform der Hopf-Bifurkation ist

$$\frac{dz}{dt} = z((\lambda + \mathfrak{I}) + (\alpha + \mathfrak{I}\beta)|z|^2). \quad (8.14)$$

Dabei sind z eine komplexe Größe, t die Zeit, $\mathfrak{I}$ die imaginäre Einheit, λ , α und β sind reelle Parameter, λ ist ein Eigenwert.

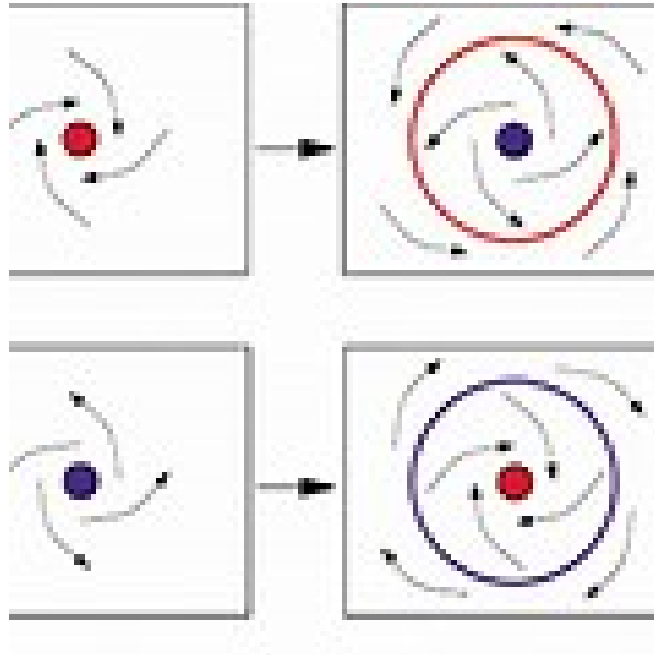


Abbildung 8.23: Die Hopf-Bifurkation

Die transkritische Bifurkation beschreibt einen Vorgang, bei dem die Stabilität („anziehend“ oder „abstoßend“) zweier Ruhelagen eines Systems vertauscht wird. Sie ist damit ein bestimmter Typ einer Bifurkation eines nichtlinearen Systems.

Die Normalform der transkritischen Bifurkation ist:

$$\frac{dz}{dt} = \mu \cdot x - x^2, \quad (8.15)$$

wobei μ der Bifurkationsparameter ist. Die transkritische Bifurkation hat folgende Gleichgewichtspunkte:

$$x_1^* = 0 \quad (8.16)$$

$$x_2^* = \mu. \quad (8.17)$$

Setzt man $x = (x_{1/2}^* + \delta)$ mit $\delta \ll 1$ in die Normalform ein (d. h. man stört den Fixpunkt) und vernachlässigt alle Terme der Ordnung δ^2 , erhält man

$$\frac{d\delta}{dt} = \begin{cases} \mu\delta, & \text{bei } x_1^*, \\ -\mu\delta & \text{bei } x_2^*. \end{cases} \quad (8.18)$$

für die zeitliche Entwicklung der Störung δ .

Für $\mu < 0$ ist also x_1^* ein stabiler Fixpunkt (d. h. die Störung nimmt mit der Zeit ab) und x_2^* ein instabiler (die Störung wächst). Für $\mu > 0$ ist es umgekehrt. Bei dem kritischen Wert des Bifurkationsparameters $\mu = 0$ ist der (in diesem Fall einzige) Fixpunkt $x^* = 0$ indifferent stabil.

Die diskrete logistische Abbildung

$$x_{n+1} = rx_n(1 - x_n) \quad (8.19)$$

folgt ebenfalls einer transkritischen Bifurkation. Sie besitzt die Fixpunkte $x_1^* = 0$ und $x_2^* = 1 - \frac{1}{r}$. Der Ursprung x_1^* ist hier stabil für $r < 1$ und instabil für $r > 1$, während x_2^* für $1 < r < 3$ stabil ist und diese Stabilität für $r > 3$ verliert.

Die logistische Gleichung kann aus der kontinuierlichen Normalform durch den Übergang $\frac{dx}{dt} \rightarrow x_{n+1} - x_n$ und die Transformation $\frac{x_n}{1+\mu} \rightarrow x_n, r = 1 + \mu$ gewonnen werden.

9 Zellulare Automaten

Zelluläre oder auch zellulare Automaten dienen der Modellierung räumlich diskreter dynamischer Systeme, wobei die Entwicklung einzelner Zellen zum Zeitpunkt $t + 1$ primär von den Zellzuständen in einer vorgegebenen Nachbarschaft und vom eigenen Zustand zum Zeitpunkt t abhängt. Ein Zellularautomat ist durch folgende Größen festgelegt:

- ein Raum R (Zellularraum)
- eine endliche Nachbarschaft N
- eine Zustandsmenge Q
- eine lokale Überföhrungsfunktion $\delta: Q^N \rightarrow Q$.

Der Zellularraum besitzt eine gewisse Dimensionalität, er ist in der Regel ein- oder zweidimensional, kann aber durchaus auch höherdimensional sein. Man beschreibt das Aussehen eines Zellularautomaten durch eine globale Konfiguration, welche eine Abbildung aus dem Zellularraum in die Zustandsmenge ist, das heißt, man ordnet jeder Zelle des Automaten einen Zustand zu. Der Übergang einer Zelle von einem Zustand (lokale Konfiguration) in den nächsten wird durch Zustandsübergangsregeln definiert, die deterministisch oder stochastisch sein können. Die Zustandsübergänge erfolgen für alle Zellen nach derselben Überföhrungsfunktion und gleichzeitig. Die Zellzustände können wie die Zeitschritte diskret sein. In der Regel ist die Anzahl der möglichen Zustände klein: Nur wenige Zustandswerte reichen zur Simulation selbst hochkomplexer Systeme aus.

Man unterscheidet zwei verschiedene Nachbarschaften (auch Nachbarschaftsindex genannt):

- Die Moore-Nachbarschaft ist eine Nachbarschaftsbeziehung in einem quadratischen Raster. Alle Flächen, welche mindestens eine Ecke mit der Basisfläche gemeinsam haben, gelten als Nachbarn. Sie ist nach Edward F. Moore benannt und wird auch als 8er-Nachbarschaft bezeichnet. Sie entspricht den

Zugmöglichkeiten des Königs beim Schach.

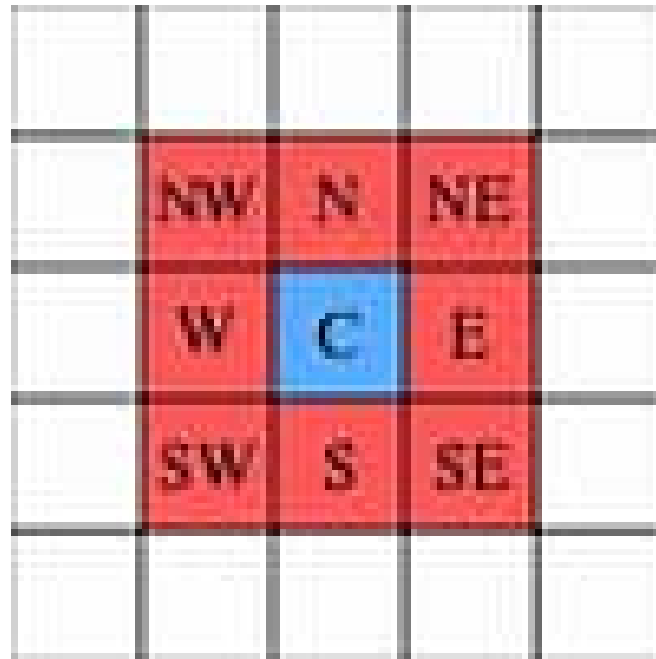


Abbildung 9.1: Die Moore-Nachbarschaft

- Die Von-Neumann-Nachbarschaft berücksichtigt nur vier Zellen aus ihrer Nachbarschaft, die direkt oberhalb und unterhalb der betrachteten Zelle bzw. links und rechts von ihr liegen.

Viele Fragen im Zusammenhang mit Algorithmen erfordern eine genaue mathematische Definition eines Algorithmus. Alan Mathison Turing (1912 - 1954) war ein britischer Logiker, Mathematiker, Kryptoanalytiker und Computerwissenschaftler. Er gilt heute als einer der einflussreichsten Theoretiker der frühen Computerentwicklung und der Computerwissenschaft. Turing schuf einen großen Teil der theoretischen Grundlagen für die moderne Informations- und Computertechnologie. Auch seine Beiträge zur theoretischen Biologie erwiesen sich als bahnbrechend.

Das von ihm entwickelte Berechenbarkeitsmodell der Turing-Maschine bildet eine der Grundlagen der Theoretischen Informatik. Während des Zweiten Weltkriegs war er maßgeblich an der Entschlüsselung deutscher Funkprüche beteiligt, die mit der deutschen Rotorchiffriermaschine Enigma verschlüsselt wurden. Die meisten seiner Arbeiten blieben auch nach Kriegsende unter Verschluss. Die Erkenntnisse, die Turing bei der Kryptoanalyse der Fish-Chiffren gewann, halfen später bei der Entwicklung des ersten digitalen, programmierbaren elektronischen Röhrencomputers, ENIAC.

1953 entwickelte Turing eines der ersten Schachprogramme, dessen Berechnungen er aus Mangel an Hardware selbst durchführte. Der Turing Award, die wichtigste

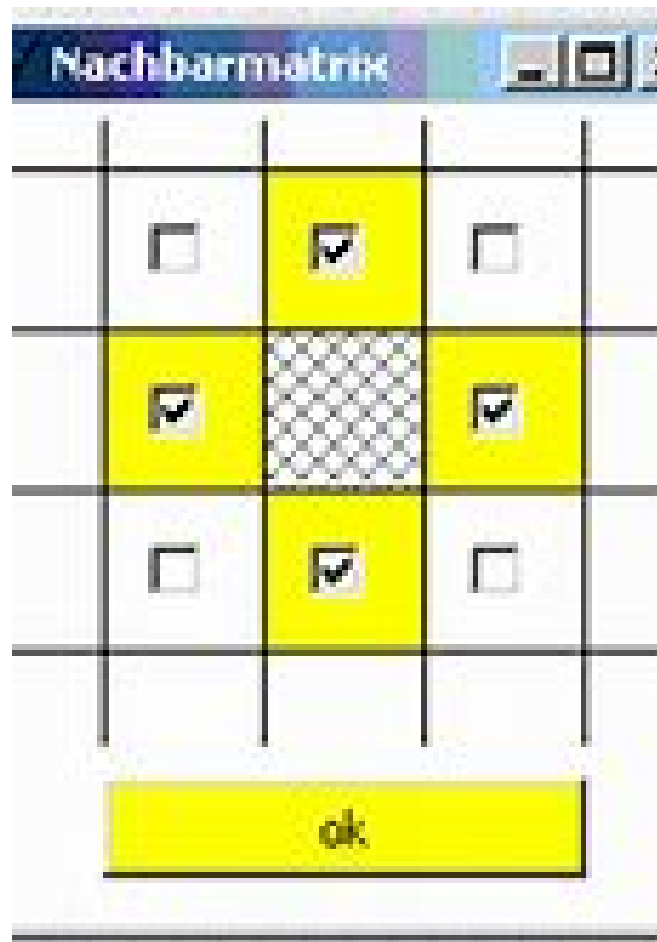


Abbildung 9.2: Die Von Neumann-Nachbarschaft

Auszeichnung in der Informatik, ist nach ihm benannt, ebenso wie der **Turing-Test** zum Nachweis des Vorhandenseins von Künstlicher Intelligenz.

Bei jedem Schritt liest der Schreib-Lese-Kopf das aktuelle Zeichen, überschreibt es mit einem anderen oder dem gleichen Zeichen und bewegt sich dann nach links oder rechts oder bleibt stehen. Welches Zeichen geschrieben und welche Bewegung ausgeführt wird, hängt von dem Zeichen ab, das an der aktuellen Position gefunden wird, sowie von dem Zustand, in dem sich die Turingmaschine gerade befindet. Zu Beginn befindet sich die Turingmaschine in einem bestimmten Startzustand und geht bei jedem Schritt in einen neuen Zustand über. Die Anzahl der Zustände, in denen sich die Turingmaschine befinden kann, ist endlich. Ein Zustand kann mehrmals durchlaufen werden, er sagt nichts über die auf dem Band vorhandenen Zeichen aus.

Eine Turingmaschine bleibt stehen, wenn für den aktuellen Zustand und das gelesene Bandzeichen kein Übergang zu einem neuen Zustand definiert ist. Es hängt also im allgemeinen von der Kombination aus Zustand und Symbol ab, ob die Turingmaschine weiterrechnet oder anhält. Zustände, in denen die Turingmaschine

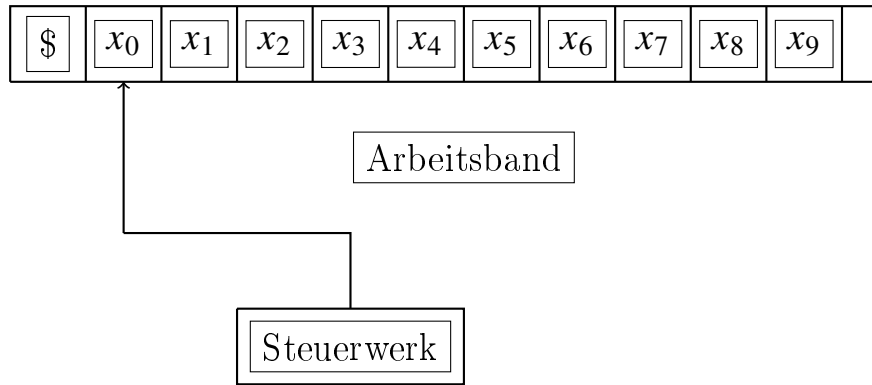


Abbildung 9.3: Eine Einband-Turing-Maschine

unabhängig vom gelesenen Bandsymbol anhält, nennt man Endzustände. Es gibt aber auch Turingmaschinen, die bei bestimmten Eingaben nie anhalten.

Eine Turingmaschine ist ein Tupel

$$M = (Z, E, A, d, q, F)$$

mit

- Z eine endliche, nichtleere Menge von Zuständen des Steuerwerks,
- E das Eingabealphabet,
- A das Arbeitsalphabet mit $E \subseteq A$,
- d die Übergangsrelation, $d \subseteq Z \times A \times A' \times Z$; hierbei ist $A' = A \cup \{L, R\}$,
- $q \in Z$ der Startzustand,
- $F \subseteq Z$ die Menge der Endzustände.

Das Konzept der Einband - Turingmaschine ist nahezu identisch mit dem Konzept des eindimensionalen zellulären Automaten. Ein eindimensionaler zellulärer Automat besteht aus einer Reihe von Zellen, wobei jede Zelle bestimmte Anfangswerte und eine Liste von Regeln hat, die festlegen, wie diese Werte in jedem Takt geändert werden. Man kann jetzt davon ausgehen, dass eine ganze Folge von Zellen mit dem Wert 0 gefüllt sind, es tritt nur einmal der Wert 1 auf. Genauso würde der Anfangszustand der Turing-Maschine aussehen. Die Anfangssituation hat also die folgende Gestalt:

...001000000...

Als Regel kann man jetzt festlegen, dass jede Zelle ihren eigenen Wert und den Wert des linken Nachbarn addiert und den Wert des linken Nachbarn durch den neuen Wert ersetzt. Man beginnt mit der zweiten Zelle und erhält $0+0=0$, danach folgt $1+0=1$, $0+1=1$. Es entsteht

...01100000...

Nach weiteren Schritten erreicht man

...01210000...
 ...01210000...
 ...01331000...
 ...01464100...
 ...

Man hat also mit dieser einfachen Vorschrift einen Algorithmus realisiert, der der Reihe nach die auftretenden Koeffizienten der binomischen Formeln

$(a+b)^0$
 $(a+b)^1$
 $(a+b)^2$
 $(a+b)^3$
 $(a+b)^4$
 ...

berechnet. Der Algorithmus arbeitet sehr schnell, weil er im wesentlichen nur aus parallel ausführbaren Operationen besteht.

9.1 Conway's Game of Life

Das Spiel des Lebens (Conway's Game of Life) ist ein vom Mathematiker John Horton Conway 1970 entworfenes Spiel, basierend auf einem zweidimensionalen zellulären Automaten. Es ist eine einfache und bis heute populäre Umsetzung der Automaten-Theorie von Stanisław Marcin Ulam.

Das Spielfeld ist in Zeilen und Spalten unterteilt und im Idealfall unendlich groß. Jedes Gitterquadrat ist ein zellulärer Automat (Zelle), der einen von zwei Zuständen einnehmen kann, welche oft als lebendig und tot bezeichnet werden. Zunächst

wird eine Anfangsgeneration von lebenden Zellen auf dem Spielfeld platziert. Jede lebende oder tote Zelle hat auf diesem Spielfeld genau acht Nachbarzellen, die berücksichtigt werden (Moore-Nachbarschaft). Die nächste Generation ergibt sich durch die Befolgung einfacher Regeln.

Das Spiel kann manuell auf einem Stück Papier oder mit Computerhilfe simuliert werden. Da ein reales Spielfeld immer einen Rand hat, muss das Verhalten dort festgelegt werden. Man kann sich den Rand zum Beispiel durch tote Zellen belegt denken, so dass manche Gleiter ihre Bewegungsrichtung dort ändern. Eine andere Möglichkeit ist ein Torus-förmiges Spielfeld, bei dem alles, was das Spielfeld nach unten verlässt, oben wieder hereinkommt und umgekehrt, und alles, was das Spielfeld nach links verlässt, rechts wieder eintritt und umgekehrt.

Anstatt auf einer quadratisch gerasterten Ebene kann die Simulation auch auf einer sechseckig gerasterten Ebene erfolgen. Dann beträgt die Zahl der Nachbarzellen nicht acht, sondern sechs. Es gibt auch dreidimensionale Game of Life-Simulationen.

Eine weitere Variationsmöglichkeit ist, mehr als zwei mögliche Zustände der Gitterzellen einzuführen.

Die Folgegeneration wird für alle Zellen gleichzeitig berechnet und ersetzt die aktuelle Generation. Der Zustand einer Zelle (lebendig oder tot) in der Folgegeneration hängt nur vom aktuellen Zustand der Zelle selbst und den aktuellen Zuständen ihrer acht Nachbarzellen ab.

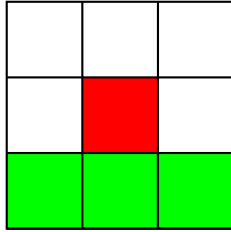
Die Abbildung 9.4 zeigt die von Conway zu Anfang verwendeten Regeln.

Mit diesen vier einfachen Regeln entsteht aus bestimmten Anfangsmustern im Laufe des Spiels eine Vielfalt komplexer Strukturen. Einige bleiben unverändert, andere oszillieren und wieder andere wachsen oder vergehen. Manche Strukturen, sogenannte Gleiter, bewegen sich auf dem Spielfeld fort. Sogar logische Funktionen wie UND und ODER lassen sich durch bestimmte Anfangsmuster simulieren. Damit können dann sogar komplexe Funktionen der Schaltungslogik und digitalen Rechnertechnik nachgebaut werden.

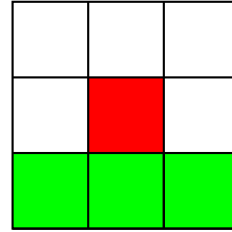
Die Beschäftigung mit Game of Life kann unter verschiedenen Sichtweisen erfolgen, wie etwa:

Das Verhalten als Ganzes: Für einige Leute ist es interessant, was für ein Verhalten bestimmte Regelwelten aufweisen, zum Beispiel ob sie explodieren oder implodieren, ob sie langsam schrumpfen oder ob sie langsam aushärten.

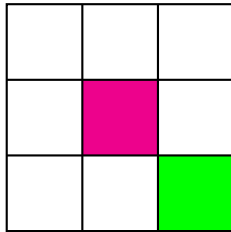
Eine tote Zelle mit genau drei lebenden Nachbarn wird in der Folgegeneration neu geboren.



Eine lebende Zelle mit zwei oder drei lebenden Nachbarn bleibt in der Folgegeneration am Leben.



Lebende Zellen mit weniger als zwei lebenden Nachbarn sterben in der Folgegeneration an Einsamkeit.



Lebende Zellen mit mehr als drei lebenden Nachbarn sterben in der Folgegeneration an Überbevölkerung.

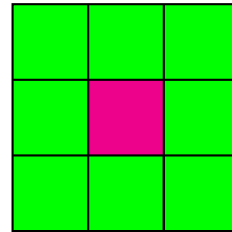


Abbildung 9.4: Regeln von Conway

- Der biologische Aspekt. Game of Life als Mikrokosmos:

Für andere ist Game of Life wie der Blick in ein Mikroskop. Man beobachtet die kleinen Strukturen, die man abzählen und bewerten kann. Hier freut man sich somit besonders, wenn eine neue „Lebensform“ auftaucht. Explodierende, expandierende oder gar aushärtende Regelwelten sind hierbei uninteressant.

- Der ökonomische Aspekt:

Game of Life als Modell des Computerhandels der Finanzmärkte. Gemäß der Algorithmen des Computerhandels kauft man ein Produkt, wenn einige, aber nicht zu viele und nicht zu wenige Nachbarn es ebenfalls bereits besitzen. Wenn zu wenige es haben, verkauft man, bevor es ganz wertlos wird. Wenn zu viele es haben, verkauft man, bevor die Blase platzt.

- Der chemische Aspekt: Energie und Materie.

Wenn man die Häufigkeit und die Komplexität der Game-of-Life-Objekte mit dem Aufwand an Energie und Zwischenschritten vergleicht, die benötigt wer-

den, um eine bestimmte chemische Verbindung zu erhalten, so kann man die unterschiedlichen Life-Objekte auf unterschiedliche energetische Niveaus setzen. Objekte, die bei jedem Ablauf vorkommen, wären dann auf dem Niveau von Wasser, Kohlenstoffdioxid und Natriumchlorid. Objekte, wie Unruh und Fontäne, wären dann beispielsweise auf einem Niveau wie Salzsäure und Natronlauge, und Objekte wie die Segler (LWSS, MWSS und HWSS), die auch zufällig entstehen können, wären schon auf dem Niveau relativ komplexer Verbindungen.

- Der physikalische Aspekt: Kräfte und Anfangswertproblem.

Selbst die einfachsten physikalischen Gesetze können beliebig komplexes Verhalten als Gesamtes zeigen. Rein deterministisch/mechanisch können (beliebig) kleine Abweichungen der Startbedingung zu ganz unterschiedlichen Ergebnissen führen. Somit lässt sich ein Anfangswertproblem formulieren, worauf chaotisches Verhalten folgt. Es folgen Endzustände, Schwingungen, Wachstum, aber auch dauerhaft unregelmäßiges Verhalten.

- Game of Life als Automat:

Es gibt den Typus des Game-of-Life-Interessierten, der hauptsächlich an der Konstruktion von Automaten interessiert ist, also solchen Strukturen, die wie eine Maschine oder Fabrik arbeiten. Es gibt einen Verband aus Strukturen, der entfernt Ähnlichkeit mit einem Rollfeld eines Flughafens hat, auf dem ständig Flugzeuge starten, und dazwischen die Fahrzeuge, die den Betrieb aufrechterhalten, zu ihren Stationen fahren.

- Game of Life als Rechnermodell: Es ist möglich, mithilfe komplexer Startmuster eine Universelle Turing-Maschine und deren Eingabe zu modellieren. Conway's Game of Life ist damit Turing-vollständig. Theoretisch lässt sich jedes algorithmische Problem, das man mit einem Computer lösen kann, auch allein durch Game of Life berechnen. Das wurde bereits dargestellt.
- Game of Life in der Theoretischen Informatik als Entscheidungsproblem:

Man kann zeigen, dass es keinen Algorithmus gibt, der als Eingabe zwei beliebige Game-of-Life-Konfigurationen erhält und in allen Fällen entscheiden kann, ob eine Konfiguration aus der anderen entstehen kann oder nicht. Diese Frage ist damit unentscheidbar.

Auf dem Spielfeld zeigt sich mit jedem Generationsschritt eine Vielfalt komplexer Strukturen. Einige typische Objekte lassen sich aufgrund eventuell vorhandener be-

sonderer Eigenschaften in Klassen einteilen: sie verschwinden, bleiben unverändert, verändern sich periodisch (oszillieren), bewegen sich auf dem Spielfeld fort, wachsen unaufhörlich usw.

- Statische Objekte bilden eine Klasse von Objekten, die sich im Spielverlauf ohne äußere Einflüsse nicht mehr verändern, also „stabile Zellsysteme“ darstellen (Abb. 9.5).

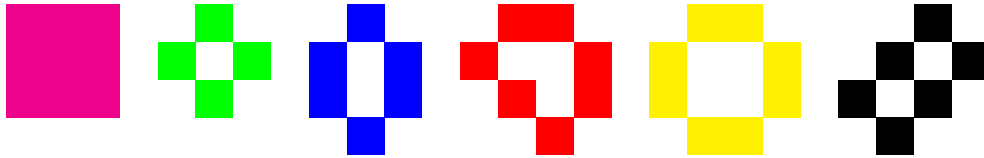


Abbildung 9.5: Stabile Zellsysteme

- Oszillierende Objekte (Abb.9.6): Hierbei handelt es sich um Objekte, die sich nach einem bestimmten Schema periodisch verändern, d. h. nach einer endlichen, festen Anzahl von Generationen wieder den Ausgangszustand erreichen. Die einfachste zyklische Konfiguration ist eine horizontale oder vertikale Reihe von drei lebenden Zellen. Beim horizontalen Fall wird direkt ober- und unterhalb der Zelle in der Mitte eine lebende Zelle geboren, während die äußeren beiden Zellen sterben; so erhält man eine vertikale Dreierreihe. Eine Reihe von zehn horizontal oder vertikal aneinander hängenden Zellen entwickelt sich sogar zu einem Objekt, das einen Zyklus von fünfzehn Generationen hat, dem Pulsator.

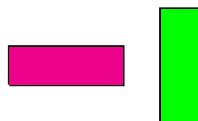


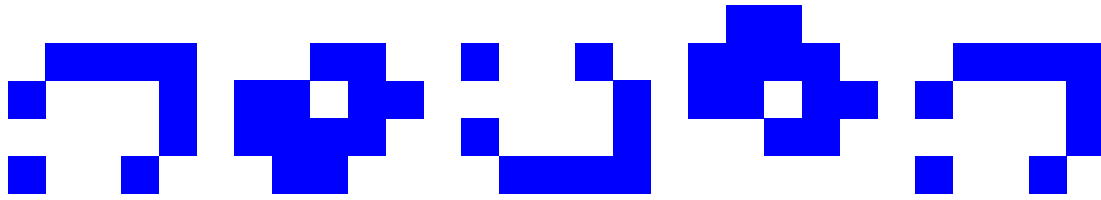
Abbildung 9.6: Oszillierende Objekte

Man findet im Internet eine große Zahl von Programmen, die Startsituationen unter unterschiedlicher Zielstellung über viele Generationen hinweg verfolgt. Als Beispiel seien etwa [40] und [41] angeführt.

Als nächstes soll ein oszillierendes Muster dargestellt werden, das aus 12 Elementen besteht und eine Periode 15 besitzt.



Die folgende Konstellation erreicht bereits nach dem vierten Schritt den Anfangszustand.



Jede Figur lässt sich aus fünf Teilrechtecken zusammensetzen, die Zahl der Elemente bewegt sich immer von 9 nach zwölf und zurück. Bei jedem Schritt werden drei Elemente neu erzeugt, im nächsten Schritt sterben drei Elemente.

Die Ameise ist eine Turingmaschine mit einem zweidimensionalen Speicher und wurde 1986 von Christopher Langton entwickelt. Sie ist ein Beispiel dafür, dass ein deterministisches (das heißt nicht zufallsbedingtes) System mit einfachen Regeln und einfachem Anfangszustand sowohl für den Menschen visuell überraschend ungeordnet erscheinende als auch regelmäßig erscheinende Zustände annehmen kann.

Die Ameise bewegt sich in einem unendlichen Quadratgitter aus Feldern, die entweder schwarz oder weiß sein können. In der Ausgangssituation sind alle Felder weiß und die Ameise sitzt auf einem der Felder und schaut in eine bestimmte Richtung (in dieser Darstellung nach unten). Der Übergang zum nächsten Zustand erfolgt nach folgenden Regeln:

- Auf einem weißen Feld drehe 90 Grad nach rechts; auf einem schwarzen Feld drehe 90 Grad nach links.
- Wechsle die Farbe des Feldes (weiß nach schwarz oder schwarz nach weiß).
- Gehe zum nächsten Feld in der aktuellen Blickrichtung.

In den ersten ca. 500 Schritten treten wiederholt symmetrische Muster auf. Danach bildet die Ameise während rund 10.000 Schritten ein komplexes, chaotisches Muster. Schließlich geht sie dazu über, eine regelmäßige Struktur („Ameisenstraße“ zu bauen: Sie gerät alle 104 Schritte in denselben (lokalen) Zustand; jeweils diagonal um 2 Felder verschoben, und baut die Straße bis ins Unendliche weiter.

Greg Turk und Jim Propp untersuchten eine Verallgemeinerung der klassischen Langton-Ameise. Ein Feld durchläuft dabei einen Zyklus von zwei oder mehr Feldzuständen (Farben): Bevor die Ameise zum nächsten Feld geht, ändert sie den

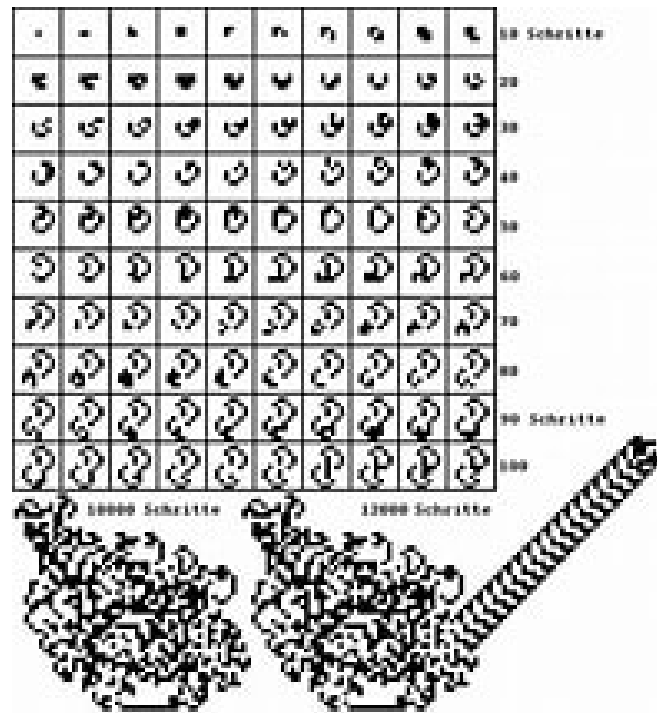


Abbildung 9.7: Das Bild einer Ameisenstraße

Zustand des aktuellen Felds zum nächsten im Zyklus. Jedem Zustand ist eine Schwenkrichtung zugeordnet, entweder nach Links oder nach Rechts um 90 Grad. Die ursprüngliche Langton-Ameise wird durch die Regel 'RL' beschrieben.

Einige Regeln erzeugen symmetrische Abbildungen, andere scheinbar vollkommen chaotische, wobei teilweise unbekannt ist, ob diese nach hinreichend vielen Schritten eine Ameisenstraße erzeugen.

Eine weitere Form der dynamische Ordnungsbildung im Chaos gelten in der Chaosforschung solitäre Wellen oder Solitonen. Bereits der schottische Ingenieur und Schiffsbauer John Scott Russell (1808 - 1882) hatte in den dreißiger Jahren des vorigen Jahrhunderts eine eigenartige Welle beobachtet, die Naturwissenschaftlern bis dahin undenkbar erschien. Er bemerkte nämlich bei einem Ausritt an einem Schiffskanal eine Welle, die sich in dem Kanal über längere Zeit mit gleichbleibender Form und Geschwindigkeit fortpflanzte. Russell wusste, dass sich Wellen normalerweise aufgrund von vielen kleinen Störungen rasch in chaotischen Turbulenzen auflösen.

Ein „Soliton“ ist dagegen eine Welle, die in einem chaotischen System über längere Zeit stabil bleibt. Die ungewöhnliche Stabilität von Solitonen entsteht durch nichtlineare Wechselwirkungen, bei denen die verschiedenen Schwingungen in ihnen rückgekoppelt werden. Die Schwingungen in Solitonen weisen daher auch eine große Selbstähnlichkeit auf.

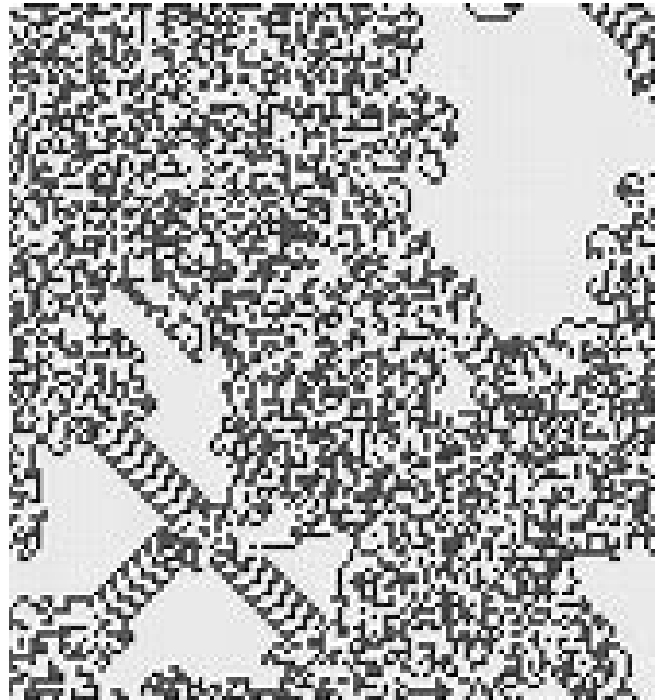


Abbildung 9.8: Langtons Ameisenstraße

Die niederländischen Mathematiker Diederik Johannes Korteweg (1848 - 1941) und G. de Vries (1858 - ?) entwickelten bereits Ende des vorigen Jahrhunderts die nichtlineare KdV-Gleichung, (so genannt nach den Anfangsbuchstaben ihrer Nachnamen), mit der man auch Solitonen berechnen kann. Diese können nämlich nur in einem eng begrenzten Bereich nichtlinearer Rückkopplung entstehen: Denn wenn die Welle zu stark ist, bricht sie bald in sich zusammen, und wenn sie zu schwach ist, verebbt sie rasch. Solitonen lassen sich auch an bestimmten Flussmündungen beobachten, wo die Gezeiten regelmäßig Flutwellen den Fluss hinauf drücken. Nichtlineare Solitonen treten aber nicht nur in engen Schiffskanälen oder Flussmündungen auf, sondern auch in den Weiten der Ozeane. Dies ist dann der Fall, wenn unterseeische Beben oder Vulkane seismische Wellen auslösen, und die bekannteste Form hierfür ist vermutlich der im Pazifischen Ozean auftretende Tsunami.

Dieses japanische Wort bedeutet „große Woge im Hafen“, und es beschreibt anschaulich, welche Zerstörungen ein bis zu 40 Meter hoher Tsunami anrichten kann, wenn er auf eine Küste trifft. Solche über hunderte von Kilometern stabilen Wellen treten aber nicht nur im Pazifischen Ozean auf, sondern auch in anderen Weltmeeren; zum Beispiel wurde das portugiesische Lissabon im Jahr 1755 durch ein Erdbeben und eine nachfolgende seismische Welle zerstört. Doch Solitonen wurden bislang nicht nur in turbulenten Gewässern entdeckt, sondern auch in den chaotischen Luftbewegungen der Erdatmosphäre. Solche atmosphärische Solitonen entstehen durch rasche Luftdruckwechsel, und sie können sich ebenfalls als stabile Druckwellen über hunderte von Kilometern fortbewegen.



Abbildung 9.9: Eine einzelne stabile Welle (Soliton)



Abbildung 9.10: Ein Tsunami trifft auf die Küste

Vollständig unvorhersehbar sind die entstehenden Strukturen, wenn sich zwei oder mehrere Wellen kreuzen.

Nichtlineare Systeme können sich also im Rahmen einer dynamischen Ordnungsbildung durch Iteration selbst zu bestimmten Strukturen organisieren (wie Attraktoren, Bifurkationen, Intermittenzen, Fraktale, Solitonen). Diese Entdeckungen stammen ausnahmslos von naturwissenschaftlichen Ansätzen der Chaosforschung, wohingegen die Geistes- und Sozialwissenschaften bis heute vernachlässigt werden. Ein wesentlicher Grund für die Vernachlässigung der geistes- und sozialwissenschaftlichen Chaosforschung liegt darin, dass die Chaostheorie im Rahmen der Physik aus dem Forschungsbereich der nichtlinearen Dynamik entstanden ist und sich die Chaosforschung daher zunächst mit verwandten Fachgebieten, wie der Mathematik,



Abbildung 9.11: Wellen, die sich überlagern, erzeugen ein absolutes Chaos

Chemie oder Biologie beschäftigt hat. Noch wichtiger ist jedoch der Grund, dass viele Chaosforscher es nicht einmal für möglich oder sinnvoll halten, geistes- und sozialwissenschaftliche Ansätze auf die Chaostheorie anzuwenden. Derartige Vorbehalte hat auch der Biophysiker, Chaosforscher und Philosoph Bernd-Olaf Küppers, der ursprünglich eine naturwissenschaftliche Laufbahn einschlug und nach seinem Physikstudium am Max-Planck-Institut für biophysikalische Chemie in Göttingen arbeitete. Dort entwickelte er eine Theorie zur Entstehung biologischer Information. Später widmete sich Küppers jedoch stärker den geisteswissenschaftlichen Fragen seiner Arbeit, und er hat mittlerweile einen Lehrstuhl am Philosophischen Institut der Universität Jena. Küppers vertritt die Ansicht, dass es grundsätzlich unweisbar ist, ob soziale Vorgänge tatsächlich nichtlinear und chaotisch sind. Eine vermutete Nichtlinearität lasse sich nicht endgültig beweisen, womit der wissenschaftliche Grundsatz der Verifizierbarkeit verletzt werde (Aussagen sind sinnlos, wenn sie sich grundsätzlich nicht beweisen lassen). Statt dessen könnten sie laut Küppers in Wirklichkeit auf eine komplexe Weise linear und somit berechenbar und vorhersagbar sein, auch wenn dies bislang nur noch nicht erkannt wurde. Gegen diese Möglichkeit lässt sich jedoch wiederum einwenden, dass es grundsätzlich unwiderlegbar ist, ob soziale Vorgänge linear ablaufen. Eine vermutete Linearität lässt sich also nicht endgültig widerlegen, womit der wissenschaftliche Grundsatz der Falsifizierbarkeit verletzt würde (Aussagen sind sinnlos, wenn sie sich grundsätzlich nicht widerlegen lassen). Küppers ist des weiteren der Meinung, dass es keinen wissenschaftlichen Nutzen biete, Ansätze der Chaostheorie auf geistes- und sozialwissenschaftliche Bereiche anzuwenden. Der Nutzen der Chaosforschung sei selbst für Naturwissenschaften schwierig zu belegen, wie zum Beispiel in der Meteorologie. Gegen diesen Einwand spricht aber, dass umgekehrt der wissenschaftliche Nutzen von klassischen, linearen Systemansätzen begrenzt ist. Weder das Wetter noch soziale Vorgänge können berechnet oder vorhergesagt werden, sondern es lassen sich allenfalls für begrenzte Zeiträume und Bereiche statistische Wahrscheinlichkeiten

ermitteln.

Es gibt bereits zahlreiche andere Ansätze, um die Chaostheorie auf verschiedene Bereiche der Geistes- und Sozialwissenschaften zu übertragen. Angesichts der Vorbehalte von Küppers überrascht es, dass sogar er selbst für ein bestimmtes Gebiet eine geistes- und sozialwissenschaftliche Chaosforschung anregt, und zwar für die Geschichtswissenschaft: „Von ihrer enormen physikalischen Bedeutung einmal abgesehen, stellen die chaotischen Systeme ganz offensichtlich auch ein interessantes Modell für das Phänomen der Geschichtlichkeit dar. Denn die Prozesse, die in solchen Systemen ablaufen, sind weder umkehrbar noch wiederholbar. Sie sind ebenso einzigartig wie alle geschichtlichen Vorgänge.“ Küppers sieht vor allem in der sensitiven Abhängigkeit historischer Entwicklungen von ihren Rahmenbedingungen einen Ansatz für die weitere chaostheoretische Geschichtsforschung. Der Soziologe Walter L. Bühl hat einen ähnlichen Ansatz und will die Chaostheorie zur Erklärung von sozialem Wandel nutzen. Er betont jedoch: „Die Chaos-Theorie ist so vor allem von Bedeutung für die Beschreibung und Erklärung von Krisen und Übergangszuständen, in denen entweder bisher wirksame Attraktoren plötzlich an Größe zu- oder abnehmen, sich verlagern oder völlig verschwinden, oder in denen sie ihre vorher reguläre Struktur verlieren. Die Chaos-Theorie ist aber sicher unbrauchbar, um einen gesellschaftlichen Dauerzustand zu beschreiben oder eine („kulturkritische“) universelle Beschreibung gesellschaftlicher Zustände oder Entwicklungen zu geben, wie dies noch bei den soziologischen Klassikern geschieht.“

Bei solchen chaostheoretischen Ansätzen muss aber immer beachtet werden, dass Geschichte (wie alle sozialen Vorgänge) wegen ihrer nichtlinearen Eigenschaften grundsätzlich unvorhersagbar und unberechenbar bleibt. Obwohl geschichtliche Vorgänge meist durch das absichtsvolle Handeln von Menschen bestimmt sind, können Revolutionen oder Kriege nicht vorhergesagt werden. Ein solches gesellschaftliches Chaos wird daher meist als erschreckend, unheimlich und gefährlich angesehen.

Ähnliches gilt auch für die Wirtschaftswissenschaft, in der Crash - Situationen oder Konjunkturkrisen nicht berechenbar sind. Die nichtlinearen Abhängigkeiten innerhalb ökonomischer Vorgänge gelten aber als belegt, die Ausprägung von chaotischem Verhalten wird allerdings noch weiter untersucht.

Psychologie: Es gibt auch vielversprechende Versuche, die Chaostheorie für die Psychologie zu nutzen. So verfolgt der Psychologe Rainer Höger einen Ansatz, um sprachpsychologische Befunde zum Stottern chaostheoretisch zu erklären.

Der Psychologe und Physiologe Michael Stadler von der Universität Bremen entwickelte ein chaostheoretisches Modell zur Erklärung von straffälligem Affektverhalten. Gemeinsam mit Thomas Fabian vom Bremer Institut für Gerichtspsychologie

stellt er fest, dass solche Affekttaten in der Regel zwei besondere Merkmale aufweisen:

- a) Der Anlass der Tat ist scheinbar geringfügig.
- b) das Ausmaß der Reaktion ist unverhältnismäßig gewaltsam.

In solchen Zuständen können sozial gelernte Verhaltensmuster versagen. Es existieren phylogenetisch tief verankerte Reaktionen wie Flucht, Angriff oder Totstellreflex, die in solchen Situationen - in der Terminologie der ChaosTheorie - starke Attraktoren darstellen. Instabilitäten gehen notwendigerweise mit Fluktuationen einher, was an Bifurkationspunkten zu einem Abgleiten des Verhaltens in solche Attraktoren führen kann. Genau dieses ist bei Affekt-Taten zu beobachten. [40]

9.2 Dreidimensionale Fraktale

In früheren Kapiteln tauchten schon gelegentlich dreidiemnsionale Verallgemeinerungen von ein - oder zweidimensionalen Fraktalen auf. Man kann sich zur Erinnerung noch einmal den Menger-Schwamm oder die Sierpinski-Pyramide anschauen. Sie sollen hier zusammenfassend untersucht werden. Dieser Abschnitt steht deshalb am Ende, weil diese Fraktale besonders schön sind.[43]

Ein apollinischen Dichtung ist ein fraktales Bild, das gebildet wird aus einer Sammlung von immer kleiner werdenden Kreisen, die in einem einzigen großen Kreis liegen. Jeder Kreis ist tangential zu den angrenzenden Kreisen - in anderen Worten, die Kreise in der apollinischen Dichtung berühren sich in einem einzelnen Punkt. Sie sind nach dem griechischen Mathematiker Apollonios von Perge benannt.

1982 wurden von A. Norton Algorithmen angegeben, die dreidimensionale fraktale Strukturen erzeugten und darstellten. Dabei wurden zum ersten Mal Quaternionen verwendet. Quaternionen sind eine Erweiterung der komplexen Zahlen:

$$H = \{a + b \cdot i + b \cdot j + b \cdot k, \quad i^2 = j^2 = k^2 = -1\}.$$

Corrado Segre (1860-1924) stellt in einer Arbeit bi-komplexe, tri-komplexe, ..., n-komplexe Zahlen vor. Eine bi-komplexe Zahl wird in folgender Weise definiert:

$$T = \{a + b \cdot i_1 + c \cdot i_2 + d \cdot j \quad i_1^2 = i_2^2 = -1, j^2 = 1\},$$

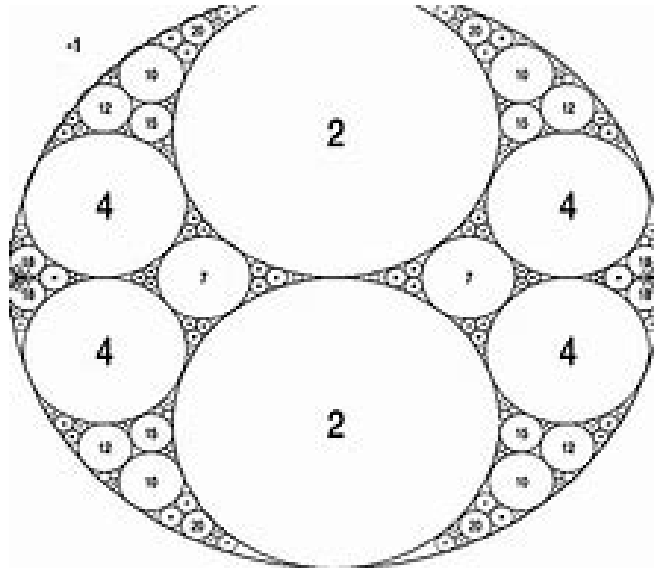


Abbildung 9.12: Eine zweidimensionale apollonische Figur

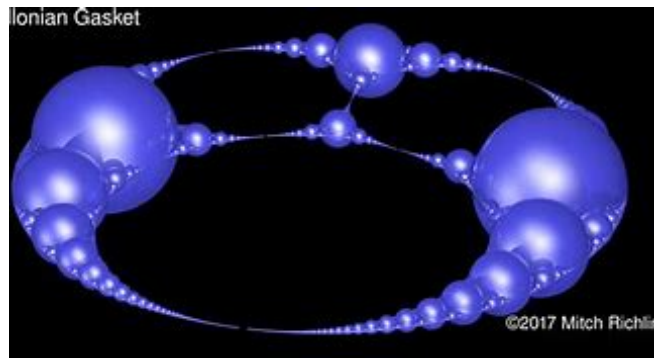


Abbildung 9.13: Ein Schmuckstück, das das apollonische Prinzip dreidimensional verwendet

wobei

$$i_2 \cdot j = j \cdot i_2 = -i_1, \quad i_1 \cdot j = j \cdot i_1 = -i_2, \quad i_2 \cdot i_1 = i_1 \cdot i_2 = j$$

gelten, wobei a, b, c, d reelle Zahlen sind. Aus dieser Definition kann man eine Reihe mathematischer Eigenschaften ableiten, die hier übergangen werden können. Wir arbeiten weiter mit dem Polynom

$$P_c(w) = w^2 + c, \quad (9.1)$$

wobei $c \in T$ und $w \in T$ gelten. Das Polynom arbeitet also mit bi-komplexen Zahlen. Genau wie früher wird dieses Polynom wiederholt auf sich selbst angewendet:

$$P_c^{n+1}(w) = P_c^n(P_c(w)). \quad (9.2)$$

Es bleibt das Prinzip erhalten, dass man das Polynom auf eine bi-komplexe Zahl anwendet; danach wird das Ergebnis der ersten Anwendung als neues Argument verwendet. Man findet die gleiche Methodik vor wie bei der Definition der Mandelbrot-

Menge und verwendet bi-komplexe Zahlen anstatt der vorher verwendeten komplexen Zahlen. Auf dieser Grundlage definiert man die Mandelbrot-Menge M_2 für bi-komplexe Zahlen:

$$M_2 = \{c \in T : P_c^n(0) \text{ ist endlich}\}. \quad (9.3)$$

Ebenso verallgemeinert man die Julia-Menge auf bi-komplexe Zahlen:

$$K_{2,c} = \{w \in T : P_c^n(w) \text{ ist endlich}\}. \quad (9.4)$$

Es sollen jetzt einige Konkretisierungen dieses allgemeinen Vorgehens dargestellt werden. Das Tetrabrot ist definiert durch

$$T = \{a + b \cdot i_1 + c \cdot i_2 d \cdot j : d = 0, P_c^n(0) \text{ ist endlich}\}. \quad (9.5)$$

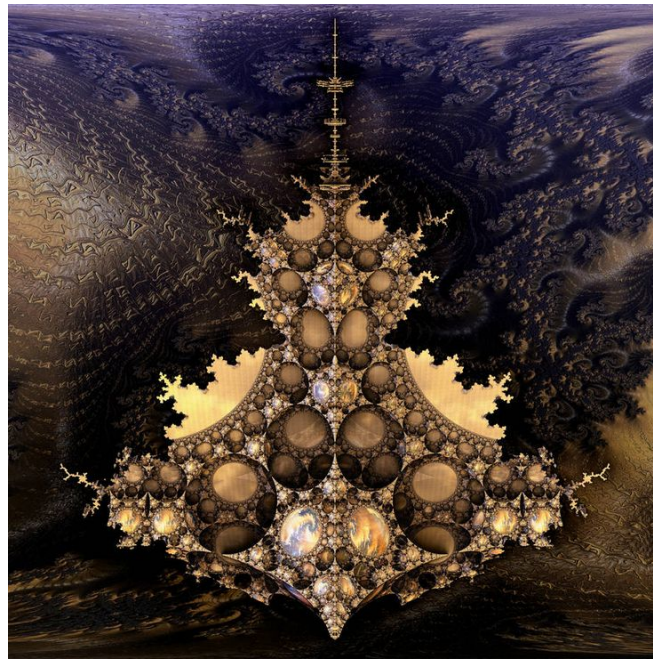


Abbildung 9.14: Eine dreidimensionale Mandelbrotstruktur

Eine bi-komplexe Julia-Menge wird in folgender Weise beschrieben:

$$L_{2,c} = \{w = a + b \cdot i_1 + c \cdot i_2 + d_j \in T : d = 0, P_c^n \text{ ist endlich}\}. \quad (9.6)$$

Unter <http://www.3dfractals.com/> kann man sich diese Bilder als Video mit Musik anschauen.

Zum Abschluss ein paar schöne, auf Fraktalen beruhende Mandalas.



Abbildung 9.15: Das Tetrabrot

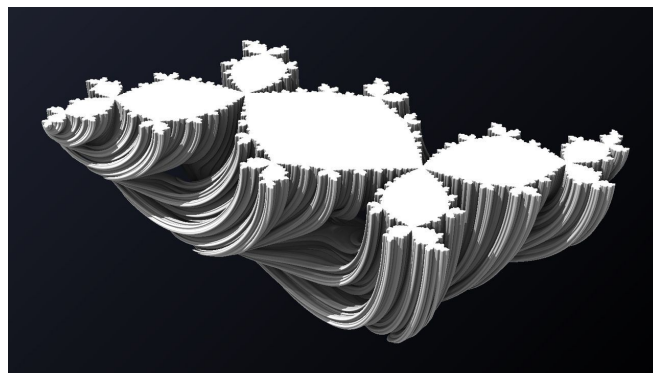


Abbildung 9.16: Eine gefüllte Julia-Menge



Abbildung 9.17: Ein Mandala



Abbildung 9.18: Ein weiteres Mandala