

Mathematik Primarstufe und Sekundarstufe I + II



Sebastian Bauer

Mathematisches Modellieren

als fachlicher Hintergrund
für die Sekundarstufe I + II

MOREMEDIA



Springer Spektrum

Mathematik Primarstufe und Sekundarstufe I + II

Reihe herausgegeben von

Friedhelm Padberg, Universität Bielefeld, Bielefeld, Deutschland

Andreas Büchter, Universität Duisburg-Essen, Essen, Deutschland

Die Reihe „Mathematik Primarstufe und Sekundarstufe I + II“ (MPS I+II), herausgegeben von Prof. Dr. Friedhelm Padberg und Prof. Dr. Andreas Büchter, ist die führende Reihe im Bereich „Mathematik und Didaktik der Mathematik“. Sie ist schon lange auf dem Markt und mit aktuell rund 60 bislang erschienenen oder in konkreter Planung befindlichen Bänden breit aufgestellt. Zielgruppen sind Lehrende und Studierende an Universitäten und Pädagogischen Hochschulen sowie Lehrkräfte, die nach neuen Ideen für ihren täglichen Unterricht suchen.

Die Reihe MPS I+II enthält eine größere Anzahl weit verbreiteter und bekannter Klassiker sowohl bei den speziell für die Lehrerbildung konzipierten Mathematikwerken für Studierende aller Schulstufen als auch bei den Werken zur Didaktik der Mathematik für die Primarstufe (einschließlich der frühen mathematischen Bildung), der Sekundarstufe I und der Sekundarstufe II.

Die schon langjährige Position als Marktführer wird durch in regelmäßigen Abständen erscheinende, gründlich überarbeitete Neuauflagen ständig neu erarbeitet und ausgebaut. Ferner wird durch die Einbindung jüngerer Koautorinnen und Koautoren bei schon lange laufenden Titeln gleichermaßen für Kontinuität und Aktualität der Reihe gesorgt. Die Reihe wächst seit Jahren dynamisch und behält dabei die sich ständig verändernden Anforderungen an den Mathematikunterricht und die Lehrerbildung im Auge.

Konkrete Hinweise auf weitere Bände dieser Reihe finden Sie am Ende dieses Buches und unter <http://www.springer.com/series/8296>

Sebastian Bauer

Mathematisches Modellieren

als fachlicher Hintergrund für die
Sekundarstufe I + II

Sebastian Bauer
Mathematisches Institut
Georg-August-Universität Göttingen
Göttingen, Niedersachsen, Deutschland

Mathematik Primarstufe und Sekundarstufe I + II
ISBN 978-3-662-61787-8 ISBN 978-3-662-61788-5 (eBook)
<https://doi.org/10.1007/978-3-662-61788-5>

Die Deutsche Nationalbibliothek verzeichnet diese Publikation in der Deutschen Nationalbibliografie; detaillierte bibliografische Daten sind im Internet über <http://dnb.d-nb.de> abrufbar.

© Springer-Verlag GmbH Deutschland, ein Teil von Springer Nature 2021

Das Werk einschließlich aller seiner Teile ist urheberrechtlich geschützt. Jede Verwertung, die nicht ausdrücklich vom Urheberrechtsgesetz zugelassen ist, bedarf der vorherigen Zustimmung des Verlags. Das gilt insbesondere für Vervielfältigungen, Bearbeitungen, Übersetzungen, Mikroverfilmungen und die Einspeicherung und Verarbeitung in elektronischen Systemen.

Die Wiedergabe von allgemein beschreibenden Bezeichnungen, Marken, Unternehmensnamen etc. in diesem Werk bedeutet nicht, dass diese frei durch jedermann benutzt werden dürfen. Die Berechtigung zur Benutzung unterliegt, auch ohne gesonderten Hinweis hierzu, den Regeln des Markenrechts. Die Rechte des jeweiligen Zeicheninhabers sind zu beachten.

Der Verlag, die Autoren und die Herausgeber gehen davon aus, dass die Angaben und Informationen in diesem Werk zum Zeitpunkt der Veröffentlichung vollständig und korrekt sind. Weder der Verlag, noch die Autoren oder die Herausgeber übernehmen, ausdrücklich oder implizit, Gewähr für den Inhalt des Werkes, etwaige Fehler oder Äußerungen. Der Verlag bleibt im Hinblick auf geografische Zuordnungen und Gebietsbezeichnungen in veröffentlichten Karten und Institutionsadressen neutral.

Planung/Lektorat: Annika Denkert

Springer Spektrum ist ein Imprint der eingetragenen Gesellschaft Springer-Verlag GmbH, DE und ist ein Teil von Springer Nature.

Die Anschrift der Gesellschaft ist: Heidelberger Platz 3, 14197 Berlin, Germany

Hinweis der Herausgeber

Dieser Band von Sebastian Bauer bietet fachliches Hintergrundwissen zum Mathematischen Modellieren für (angehende) Lehrkräfte der Sekundarstufe II. Der Band erscheint in der Reihe Mathematik Primarstufe und Sekundarstufe I + II, aus der Sie insbesondere die folgenden Bände unter mathematischen oder mathematikdidaktischen Gesichtspunkten interessieren könnten:

- A. Büchter/H.-W. Henn: Elementare Analysis – Von der Anschauung zur Theorie
- H. Humenberger/B. Schuppar: Mit Funktionen Zusammenhänge und Veränderungen beschreiben
- G. Wittmann: Elementare Funktionen und ihre Anwendungen
- B. Barzel/M. Glade/M. Klinger: Algebra und Funktionen. Fachlich und fachdidaktisch.
- T. Leuders: Erlebnis Algebra – zum aktiven Entdecken und selbstständigen Erarbeiten
- A. Filler: Elementare Lineare Algebra – Linearisieren und Koordinatisieren
- F. Padberg/A. Büchter: Elementare Zahlentheorie
- H. Kütting/M. Sauer: Elementare Stochastik
- B. Schuppar/H. Humenberger: Elementare Numerik für die Sekundarstufe
- B. Schuppar/Geometrie auf der Kugel – Alltägliche Phänomene rund um Erde und Himmel
- G. Greefrath: Anwendungen und Modellieren im Mathematikunterricht
- R. Danckwerts/D. Vogel: Analysis verständlich unterrichten
- G. Greefrath et al.: Didaktik der Analysis für die Sekundarstufe II
- H.-W. Henn/A. Filler: Didaktik der Analytischen Geometrie und Linearen Algebra
- K. Heckmann/F. Padberg: Unterrichtsentwürfe Mathematik Sekundarstufe I
- C. Geldermann et al.: Unterrichtsentwürfe Mathematik Sekundarstufe II
- A. Pallack: Digitale Medien im Mathematikunterricht der Sekundarstufe I + II

- K. Krüger et al.: Didaktik der Stochastik in der Sekundarstufe I
- V. Ulm/M. Zehnder: Mathematische Begabung in der Sekundarstufe

Bielefeld
Essen
März 2020

Friedhelm Padberg
Andreas Büchter

Vorwort

Dieses Buch richtet sich an Studierende und Lehrende des Lehramts der Sekundarstufe sowie Lehrerinnen und Lehrer der Sekundarstufe. Es ist auf der Basis einer einsemestrigen Vorlesung „Mathematisches Modellieren für Lehramtler“ entstanden, die der Autor mehrfach an der Universität Duisburg-Essen für Studierende des gymnasialen Lehramts gehalten hat. Es werden Kenntnisse vorausgesetzt, wie sie üblicherweise im Rahmen der Einführungsvorlesungen zur Analysis und linearen Algebra erworben werden. Insbesondere sollte die Differentialrechnung im $\mathbb{R}^n$ bekannt sein. Bezüglich der Integralrechnung werden keine Kenntnisse benötigt, die über den Begriff des Riemann-Integrals im Eindimensionalen hinausgehen.

Was will dieses Buch? Mathematisches Modellieren ist heutzutage in aller Munde, hat Eingang in die Standards, Richtlinien und Lehrpläne der Schulen gefunden und beansprucht zunehmend auch einen Platz unter den mathematischen Lehrveranstaltungen an den Universitäten. Dieses Buch will keinen dezidierten Beitrag zum fachdidaktischen Diskurs um das mathematische Modellieren in der Schule und die diversen Modellierungskreisläufe und Modellierungskompetenzen leisten. Es ist auch kein systematisches Fachbuch über angewandte Mathematik, in dem bestimmte mathematische Modelle auf dem aktuellen Stand der Wissenschaft präsentiert werden. Für Ersteres verweisen wir auf [15, 51] und die darin zitierte Literatur, für Letzteres gibt es einen nahezu unübersehbaren Bestand an einschlägigen Fachbüchern.

Ziel dieses Buchs ist es, anknüpfend an typische Modellierungsaufgaben, die aus der Schulpraxis bekannt sind, einige bedeutende, historisch interessante oder ästhetisch ansprechende mathematische Modelle unterschiedlicher Bezugswissenschaften wie der Physik, der Biologie und der Chemie zusammen mit der dafür benötigten mathematischen Theorie zu entwickeln. Dabei soll die Entwicklung des Modells in Verzahnung mit der Entwicklung der nötigen mathematischen Theorie „erzählt“ werden. Neben der unterschiedlichen Auswahl an Modellen und Kontexten sehen wir hierin den Hauptunterschied zu den auch als Fallstudien zum Modellieren konzipierten Büchern von W. Krabs [55], T. Sonar [84] und Ortlieb et al. [70] sowie den auf einem deutlich höheren fachlichen Niveau angesiedelten Büchern von Fowkes und Mahony [31] und Eck et al. [21].

Die Auswahl der studierten Modellierungen ist dabei zum einen der nötigen Einschränkung auf einen handhabbaren mathematischen Werkzeugkasten geschuldet, folgt zum anderen den Vorlieben und Kenntnissen des Autors und ist somit in keiner Weise zwingend.

Ein durchaus erwünschter Nebeneffekt der getroffenen Auswahl ist die elementare Behandlung von ebenen Kurven, Variationsprinzipien und gewöhnlichen Differentialgleichungen, allesamt stoffliche Inhalte, die nicht in jedem Fall in der gymnasialen Lehrerbildung vorkommen, aber nach Ansicht des Autors zu einer mathematischen Allgemeinbildung gehören.

Der Schwerpunkt des Buches liegt auf der Modellierung von Zeitverläufen mit Differenzen- und Differentialgleichungen, sodass bestimmte Phänomene und Techniken in unterschiedlichen Zusammenhängen und auf unterschiedlichem Niveau immer wieder aufgegriffen werden können. Dazu zählen auf der Seite der Phänomene das Langzeitverhalten, Bifurkationen, Hysterese-Effekte und der Übergang von Stabilität in das Chaos, auf der technischen Ebene die Nutzung der Dimensionsanalyse, der Lösungstheorie für Differenzen- und Differentialgleichungen, der Linearisierung um stationäre Punkte mit anschließender Stabilitätsanalyse und das systematische Vernachlässigen kleiner Einflussfaktoren sowie die Nutzung unterschiedlicher Skalen im Rahmen von asymptotischen Betrachtungen.

Dabei steht das qualitative Verständnis der Modellierung im Vordergrund. Die benötigten Rechenkalküle werden lediglich auf der linearen Ebene entwickelt, um ein qualitatives Verständnis auf der nichtlinearen Ebene via Linearisierung zu ermöglichen. Besonderer Wert wird auf graphische Verfahren wie die Nutzung von Phasenporträts gelegt, die ohne großen rechnerischen Aufwand die qualitativen Eigenschaften der Modellierung herausstellen können.

Auch wenn es nicht zu erwarten ist, dass die Leserin nach Lektüre des Buches in der Lage ist, anspruchsvolle Modellierungen selbstständig durchzuführen, so sollen doch neben den entwickelten mathematischen Inhalten auch einige für Modellierungsprozesse typische Fähigkeiten und Fertigkeiten wie die Dimensionsanalyse und das Einschätzen von Größenordnungen über das gesamte Buch hinweg entwickelt werden.

Das Buch hat einen Schwerpunkt auf Populationsmodellen, die über diskrete und kontinuierliche Situationen vom einfachsten Modell des exponentiellen Wachstums einer einzelnen Population bis hin zu interagierenden Populationen mit Schwellen- und Sättigungseffekten und unterschiedlichen Zeitskalen reichen. Dazu gesellen sich Nebenschauplätze wie Trassierungs- und Maximierungsproblem, Würfe, Reaktionskinetik und Elektrophysiologie.

Sowohl zur Berechnung von Lösungen als auch zu deren Darstellung werden Computer in diesem Buch mannigfaltig verwendet. Dabei setzen wir einen Schwerpunkt auf die schulübliche Software, nämlich Tabellenkalkulationsprogramme (TK) und Dynamische Geometriesoftware (DGS). Zum einen soll dadurch die Vertrautheit mit diesen hoch schulrelevanten Werkzeugen gefördert werden. Zum anderen bleibt gerade beim Einsatz von Tabellenkalkulationsprogrammen mehr von den algorithmischen

Verfahren sichtbar als beim (oft bequemerem) Einsatz eines Computer-Algebra-Systems (CAS), welches nur in Ausnahmefällen herangezogen wird.

Wie kann dieses Buch für die Planung einer Vorlesung verwendet werden? Alle aufgeführten Inhalte und Übungsaufgaben sind in Vorlesungen mehrfach erprobt (und modifiziert) worden. Dabei muss für eine vierstündige Vorlesung mit zweistündiger Übung eine Auswahl getroffen werden. Der Kern, der in jedem Fall enthalten sein sollte, besteht aus den Abschn. 1.1 zur Einführung in die Entdimensionalisierung sowie den Kap. 5 und 6. Dieser Teil beansprucht nach der Erfahrung des Autors über die Hälfte der Vorlesungszeit. Die Kap. 3 und 4 sowie die Kap. 7 und 8 bilden zwei voneinander unabhängige Blöcke, die je nach Zeit und Interesse teilweise oder vollständig behandelt werden können. Ähnliches gilt für die Abschn. 1.2, 1.3, 2.1 und 2.2, die einzeln je nach Interesse und Zeit behandelt werden können oder auch nicht. Das Buch enthält zu jedem Kapitel Aufgaben unterschiedlicher Komplexität. Ein Fundus von Lösungen ausgewählter Aufgaben soll mit der Zeit auf der Seite <https://www.uni-math.gwdg.de/bauer/mamo> entstehen.

Im Sinne der Lesbarkeit wird hauptsächlich die genderneutrale wir-Form genutzt. Wenn die weibliche Form benutzt wird, sind ebenfalls auch alle männlichen Leser gemeint, umgekehrt verhält es sich entsprechend. Ebenso sind stets auch alle Lesenden angesprochen, die keinem der beiden aufgeführten Geschlechter angehören.

Danken möchte ich Rafael Schmitt, Thomas Stenzel, Michael Schomburg und Stephan Berendonk für viele anregende Diskussionen, Lisa Hefendehl-Hebeker und Andreas Büchter für andauernde Unterstützung und Ermunterung sowie meiner Familie für eine gehörige Portion Nachsicht.

Mülheim an der Ruhr
im März 2020

Sebastian Bauer

Inhaltsverzeichnis

- 1 Hängende Kabel und Trassierungen – Steckbriefaufgaben und Entdimensionalisierungen** **1**
- 1.1 Einleitung: Stromkabel und Seilbahnen. 1
 - 1.1.1 Modell 1: stückweise linearer Verlauf 2
 - 1.1.2 An das Problem angepasste Skalierung und Simulation 5
 - 1.1.3 Modell 2: parabelförmiger Verlauf. 6
 - 1.1.4 Modell 3: Die Physiker sagen, es sei eine Kettenlinie. 8
 - 1.1.5 Verallgemeinerung: Wie sieht es bei anderen Verhältnissen von Abstand und Durchhang aus? 10
 - 1.1.6 Reflexionen zum mathematischen Modellieren I: der Modellbegriff und Gütekriterien für Modelle 11
- 1.2 Trassierungsaufgaben und der Begriff der Krümmung 14
 - 1.2.1 Kurven und ihre Krümmung 15
 - 1.2.2 Elementare Theorie ebener Kurven 16
 - 1.2.3 Straßenplanung mit Geraden, Kreisbögen und Klothoiden 22
 - 1.2.4 Reflexionen zum mathematischen Modellieren II: normative Modelle 25
- 1.3 Der Fahnenmast – vom Nutzen der Entdimensionalisierung und der Wahl geeigneter Skalen 26
 - 1.3.1 Annahmen und Vereinfachungen 26
 - 1.3.2 Das mathematische Modell 27
 - 1.3.3 Entdimensionalisierung 28
 - 1.3.4 Vom Nutzen der Entdimensionalisierung 30
- Aufgaben. 31
- 2 Optimale Geschwindigkeiten und optimale Funktionen – Optimierungen und Variationen.** **39**
- 2.1 Mit welcher Geschwindigkeit erhält man den optimalen Verkehrsfluss? Ein Optimierungsproblem 39

2.1.1	Das Modell „Halber Tacho“	40
2.1.2	Das Modell „Anhalteweg“	40
2.1.3	Bewertung	42
2.2	Woher kennen die Physiker die genaue Form des Kabels? – Die Kettenlinie als Minimierer	44
2.2.1	Das Energiefunktional	44
2.2.2	Ein Minimierer des Energiefunktionals und die Euler-Lagrange-Gleichungen	45
2.2.3	Die Euler-Lagrange-Gleichungen für das Kabel passen (noch) nicht	48
2.2.4	Die Kabellänge ist vorgegeben: Minimierung unter einer Nebenbedingung	49
2.2.5	Extrema unter Nebenbedingungen im zweidimensionalen Fall: Lagrange'sche Multiplikatoren	50
2.2.6	Minimierer der Energie unter der Nebenbedingungen vorgegebener Kabellänge	52
2.2.7	Beweis: Lagrange-Multiplikator für Extrema unter Nebenbedingungen.	53
2.2.8	Einige Bemerkungen zu den Extremalprinzipien	55
2.2.9	Reflexionen zum mathematischen Modellieren III: Ad-hoc-Modelle, Modelle nach First Principles und die „unreasonable effectiveness of mathematics in the natural science“	56
	Aufgaben.	57
3	Zinsen, Strahlung und Kaninchen – Modellieren mit linearen Differenzengleichungen	65
3.1	Guthaben und Zinsen	65
3.1.1	Guthaben	65
3.1.2	Festgeldkonto mit jährlicher Einlage	66
3.2	Der Radongehalt in der Luft eines unbelüfteten Kellerraums	66
3.2.1	Modellierung	67
3.2.2	Interpretation	68
3.3	Zusammenfassung, Systematisierung und graphische Analyse.	68
3.3.1	Definition, eindeutige Lösbarkeit und Darstellungsformel	68
3.3.2	Das Cobwebbing	68
3.4	Summen der n -ten Potenzen und Kaninchen – das Lösungsschema für lineare Gleichungen	70
3.4.1	Das lineare Lösungsschema.	70
3.4.2	Eine Anwendung: die Summe der ersten n -Kubikzahlen.	73
3.4.3	Eine zweite Anwendung: die Fibonacci-Folge	75
	Aufgaben.	76

4 Insekten zwischen Stabilität und Chaos – Modellieren mit nichtlinearen Differenzengleichungen	81
4.1 Die logistische Differenzengleichung – ein Prototyp für den Übergang von einem stabilen zu einem chaotischen Verhalten	81
4.1.1 Das lineare Modell	81
4.1.2 Das logistische Modell	82
4.1.3 Von stabilen Zuständen in das Chaos	83
4.2 Das Prinzip der Linearisierung	86
4.3 Zurück zur logistischen Differenzengleichung	89
4.3.1 Die Stabilität der stationären und 2-periodischen Punkte und das Bifurkationsdiagramm	89
4.3.2 Der Satz von Sarkovskii und die Feigenbaum-Konstante	91
4.4 Insektenpopulation mit nichtüberlappender Generationenfolge	93
4.4.1 Grundmodelle für Insektenpopulationen mit nichtüberlappender Generationenfolge	93
4.4.2 Das Hassell-Modell mit exakter Kompensation	94
4.4.3 Modelle mit Über- und Unterkompensation	96
4.4.4 Eine Bekämpfungsstrategie: sterile Insekten	99
Aufgaben	103
5 Populationsmodelle und Befischung – Modellieren mit Differentialgleichungen	109
5.1 Modellieren mit Differentialgleichungen	109
5.1.1 Das Modell des exponentiellen Wachstums	110
5.1.2 Das logistische Modell	112
5.2 Qualitatives Lösen von autonomen Differentialgleichungen	115
5.3 Populationsmodelle mit Bejagung	117
5.3.1 Das Lösungsverhalten der logistischen Differentialgleichung	117
5.3.2 Fangstrategien: Fischfang mit konstanter Fangrate und mit konstantem Aufwand	117
5.3.3 Schwellen- und Sättigungseffekte	121
5.4 Der Tannenwickler und die Balsamtanne – erster Teil	124
5.4.1 Die Modellierung der Populationsdichte der Tannenwickler	126
5.4.2 Stationäre Punkte, Bifurkationen und die kritische Kurve	128
5.5 Eine einfache Lösungstheorie für Differentialgleichungen – die Methode von der Trennung der Variablen	132
5.5.1 Die Existenz- und Eindeutigkeitsaussagen	133
5.5.2 Grobe Bestimmung der Zeitabhängigkeit und Vergleichssätze	137
5.6 Ein paar ausgewählte Lösungsverfahren	140
5.6.1 Die inhomogene lineare Differentialgleichung	141
5.6.2 Die Bernoulli'sche Differentialgleichung als Beispiel für ein Substitutionsverfahren	143

5.7	Ein einfaches Verfahren zur numerischen Lösung von Anfangswertproblemen	144
5.8	Modellierungen in Abituraufgaben – Blutalkohol	145
5.8.1	Die Formulierung der Abituraufgabe	146
5.8.2	Die Modellierung der Alkoholkonzentration	147
5.8.3	Die Berechnung der Modellfunktion	151
5.8.4	Vergleich der beiden Modellfunktionen	152
5.8.5	Reflexionen zum mathematischen Modellieren IV: deskriptive und explikative Modellierungen	154
	Aufgaben	155
6	Räuber-Beute-Modelle – Modellieren mit Differentialgleichungssystemen	167
6.1	Das Räuber-Beute-System nach Lotka-Volterra	167
6.2	Lösungstheorie: der Existenz- und Eindeigkeitsatz von Picard-Lindelöf	169
6.3	Das Lotka-Volterra-System fortgesetzt – die Methode der Lyapunov-Funktion	171
6.3.1	Elementare Diskussion des Phasenporträts	171
6.3.2	Eine numerische Annäherung mit dem expliziten Euler-Verfahren	175
6.3.3	Die Lyapunov-Funktion oder das erste Integral	176
6.3.4	Mittelwerte und Abschätzungen für die Periode	179
6.3.5	Anwendung auf die Interaktion von Luchs und Schneeschuhhase	181
6.4	Räuber-Beute-Systeme mit logistischem Wachstum – lineare Differentialgleichungssysteme und der Linearisierungssatz von Grobman-Hartman	184
6.4.1	Modellbildung und erste Diskussion des Modells	184
6.4.2	Lineare Differentialgleichungssysteme	187
6.4.3	Die Phasenporträts der linearen Differentialgleichungssysteme für $n = 2$	190
6.4.4	Der Linearisierungssatz von Grobman und Hartman	194
6.4.5	Fortsetzung: Räuber-Beute-Systeme mit logistischem Wachstum	196
6.4.6	Diskussion des Ergebnisses	198
6.4.7	Könnte es im Räuber-Beute-Modell mit logistischem Wachstum auch unbeschränkte Lösungen geben?	199
6.4.8	Reflexionen über das mathematische Modellieren V: Die Art des Schließens in der mathematischen Untersuchung	203

6.5	Räuber-Beute-Systeme mit Sättigungs- und Schwelleneffekten – stabile Grenzzyklen und der Satz von Poincaré-Bendixson	204
6.5.1	Die Modellbildung	204
6.5.2	Diskussion des Räuber-Beute-Modells mit Sättigung	206
6.5.3	Stabile Grenzzyklen und der Satz von Poincaré-Bendixson	211
	Aufgaben	216
7	Würfe, Enzymreaktionen und der Tannenwickler – asymptotische Methoden	223
7.1	Der senkrechte Wurf im Galilei'schen und im Newton'schen Modell	224
7.1.1	Das Galilei'sche Modell des senkrechten Wurfs	224
7.1.2	Das Newton'sche Gravitationsmodell	225
7.1.3	Entdimensionalisierung und Wahl der geeigneten Skalen	228
7.1.4	Die Idee der asymptotischen Entwicklung	230
7.1.5	Beweis der Approximationseigenschaft	233
7.2	Reaktionskinetik, Enzyme und die Methode des Asymptotic Matching	235
7.2.1	Das grundlegende Modell für chemische Reaktionen: Reaktionskinetik	235
7.2.2	Enzymatische Reaktionen – zwei unterschiedliche Zeitskalen und das Verfahren des Asymptotic Matching	237
7.2.3	Entdimensionalisierung und asymptotische Betrachtungen	241
7.3	Der Tannenwickler und die Balsamtanne – zweiter Teil	245
7.3.1	Die Modellbildung	246
7.3.2	Numerische Simulationen	250
7.3.3	Analyse des Systems mit asymptotischen Methoden	252
7.3.4	Bewertung des Modells	259
7.3.5	Exkurs: der Übergang vom stationären Punkt zum oszillierenden Verhalten	260
	Aufgaben	264
8	Das Modell der Reizverarbeitung in Nervenzellen nach Hodgkin und Huxley – Experimente, Modelle und Spielzeugmodelle	271
8.1	Die Nervenzelle und das Aktionspotential	272
8.2	Erklärungsansätze vor Hodgkin und Huxley	274
8.3	Die Versuchsreihe von Hodgkin und Huxley	274
8.3.1	Beschreibung des Versuchsaufbaus	274
8.3.2	Die Trennung in kapazitive Ströme und Ionenströme	279
8.3.3	Die Trennung in Natrium- und Kaliumionenströme	283
8.3.4	Der „passive“ Ionentransport und die selektive Leitfähigkeit	286
8.3.5	Die Kaliumleitfähigkeit	290
8.3.6	Die Natriumleitfähigkeit	292
8.3.7	Die qualitative Erklärung des Aktionspotentials	297

8.4	Die mathematische Modellierung und die quantitative Berechnung der Aktionspotentials	298
8.4.1	Die Modellierung der Kaliumleitfähigkeit.	299
8.4.2	Die Modellierung der Natriumleitfähigkeit	303
8.4.3	Simulationen mit Hilfe des mathematischen Modells	305
8.4.4	Die Erkenntnisse von H & H und das heutige Lehrbuchwissen	308
8.5	Toy-Models – das vereinfachte Modell nach FitzHugh-Nagumo	309
8.5.1	Ein Aktionspotential.	312
8.5.2	Ein biochemischer Schalter	313
8.5.3	Eine Folge von Aktionspotentialen: das neuronale Feuern	315
8.6	Vom Spielzeugmodell zum großen Modell: neuronales Feuern im Modell von Hodgkin und Huxley	317
	Aufgaben.	318
 Bisher erschienene Bände der Reihe MathematikPrimarstufe und Sekundarstufe I + II		 321
Literatur.		323
Stichwortverzeichnis.		329



Hängende Kabel und Trassierungen – Steckbriefaufgaben und Entdimensionalisierungen

1

Inhaltsverzeichnis

1.1	Einleitung: Stromkabel und Seilbahnen	1
1.2	Trassierungsaufgaben und der Begriff der Krümmung	14
1.3	Der Fahnenmast – vom Nutzen der Entdimensionalisierung und der Wahl geeigneter Skalen	26
	Aufgaben	31

Zu Beginn eines Buches mit dem Titel *Mathematisches Modellieren* erwartet der Lesende wohl zuerst eine Begriffsdefinition, was denn mathematisches Modellieren überhaupt sei, welchen Kriterien es zu entsprechen habe, wie es überhaupt um das Verhältnis von Realität und Mathematik bestellt sei, was denn der Unterschied zwischen Modellieren und angewandter Mathematik sei, welche Modelle es zur Beschreibung des mathematischen Modellierens gäbe und so fort.

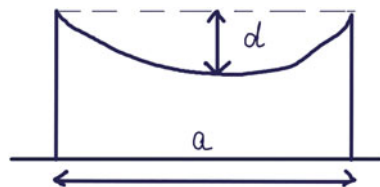
Anstatt diese Fragen zu Beginn des Buches zu thematisieren, beginnen wir lieber direkt mit einem Beispiel und versuchen Stück für Stück herauszufinden, was es mit dem mathematischen Modellieren auf sich hat. Betrachtungen über das Modellieren werden gelegentlich am Ende eines Kapitels eingestreut.

Ferner werden Dinge und Sachverhalte, die im Laufe dieses Buches immer wieder eine Rolle spielen, in einen Kasten eingerahmt.

1.1 Einleitung: Stromkabel und Seilbahnen

Zwischen zwei Masten soll ein Kabel eingehängt werden, das einen genau vorgegebenen Durchhang hat (siehe Abb. 1.1)

Abb. 1.1 Kabel zwischen zwei Masten: a = der Abstand der beiden Masten, d = der maximale Durchhang.
Beispielwerte: $a = 137.60$ m
und $d = 54.20$ m



Das Problem

Wie lang muss das Kabel sein, damit genau dieser Durchhang realisiert wird? Gesucht ist also die Länge des Kabels.

Annahmen und Idealisierungen

Zu Beginn suchen wir möglichst einfache mathematische Modellierungen und überlegen, welche Annahmen und Vereinfachungen in der angenommenen Realsituation davon impliziert werden. Wir suchen ein Modell, in dem der Verlauf des Kabels nur vom Abstand a und dem Durchhang d abhängt. Welche Vereinfachungen und Annahmen in der angenommenen Realsituation werden dadurch impliziert?

- In dem Modell steckt eine Spiegelsymmetrie zwischen den beiden Aufhängepunkten des Kabels und dem Boden. Wir verlangen, dass beide Aufhängepunkte dieselbe Höhe über dem Boden haben.
- Der Kabelverlauf soll sich alleine aus der geometrischen Situation ergeben, Materialeigenschaften des Kabels sollen keine Rolle spielen. Wir nehmen an, dass das Kabel sich in seinem Längenverlauf beliebig verformen kann.
- Die Länge des Kabels soll konstant sein. Wir nehmen also an, dass wir Ausdehnungen aufgrund des eigenen Gewichts oder aufgrund von Temperaturschwankungen vernachlässigen können.

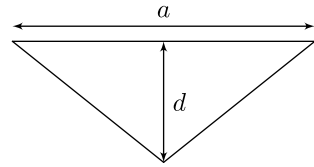
Damit sollen der Verlauf und die Länge des Kabels nur von den Größen a und d abhängen.

1.1.1 Modell 1: stückweise linearer Verlauf

Wir wollen zunächst ein möglichst einfaches Modell aufstellen: Das Kabel besteht aus stückweisen Geraden und hat an seiner tiefsten Stelle einen Knick, siehe Abb. 1.2. Mit Pythagoras erhalten wir sofort

$$l = 2\sqrt{\frac{a^2}{4} + d^2}$$

Abb. 1.2 Ein stückweise linearer Kabelverlauf



und mit unseren Beispielwerten $a = 137,60 \text{ m}$ und $d = 54,20 \text{ m}$ ergibt das

$$\begin{aligned} l &= 2\sqrt{\frac{(137,60 \text{ m})^2}{4} + (54,20 \text{ m})^2} \\ &= 2\sqrt{7671,08 \text{ m}^2} = 175,2 \text{ m}. \end{aligned} \quad (1.1)$$

Der Umgang mit dem Gleichheitszeichen

Der Umgang mit dem Gleichheitszeichen ist bestimmt erläuterungsbedürftig, klarerweise gilt ja $\sqrt{7671,08} \neq 175,2$. Natürlich sind die beiden reellen Zahlen $\sqrt{7671,08}$ und $175,2$ nicht gleich, wir müssten also $\sqrt{7671,08} \approx 175,2$ schreiben, um auszudrücken, dass $\sqrt{7671,08}$ im Intervall $[175,165; 175,175]$ liegt. In unserem Beispiel handelt es sich aber um die Maßzahl einer empirisch gemessenen kontinuierlichen Größe, die nur im Rahmen einer gewissen Messgenauigkeit, also als ein Intervall, angegeben werden kann. Die Angabe, ein Gegenstand habe die Länge $l = 5,0 \text{ m}$, trägt dabei durch die Angabe der gültigen oder signifikanten Ziffern die Messgenauigkeit in sich: Die Länge ist bis auf Dezimeter genau bestimmt. Die tatsächliche Länge liegt also irgendwo zwischen 495 cm und 505 cm . Um zu unterscheiden, ob wir es mit Maßzahlen von kontinuierlichen Größen, also eigentlich mit Intervallen zu tun haben oder mit normalen reellen Zahlen, werden wir erstere mit dem Punkt als Dezimaltrennungszeichen schreiben, letztere mit dem Komma.

Die in die Rechnung (1.1) eingehenden Größen können eine unterschiedliche Anzahl gültiger Ziffern haben. Grob gesprochen kann das Ergebnis dann nur bis auf die geringste Anzahl signifikanter Ziffern bestimmt sein. Wenn man es genauer nehmen möchte, muss untersucht werden, wie sich durch die Rechnungen „die Fehler fortpflanzen“, mehr darüber findet sich z. B. in [9, 80]. Wir gehen in diesem Buch jedoch pragmatisch vor und davon aus, dass diejenige Größe, die mit der geringsten Genauigkeit, also mit der kleinsten Anzahl signifikanter Ziffern, in das Problem eingeht, auch die Genauigkeit des Ergebnisses bestimmt.

In unserem Beispiel hat die Größe d die geringste Anzahl signifikanter Ziffern, nämlich vier. Nach unserer Daumenregel hat dann auch das Ergebnis vier signifikante Ziffern, $l = 175.2$ m.

So wie in (1.1) würde die Angelegenheit wohl im Physikunterricht der Schule notiert. Im Mathematikunterricht wird häufig anders vorgegangen:

$$l = 2\sqrt{\frac{137.60^2}{4} + 54.20^2} = 175.2 \quad l \text{ in Meter}$$

In den USA würde man stattdessen dieses finden:

$$l = 2\sqrt{\frac{451.44^2}{4} + 83.01^2} = 481.0 \quad l \text{ in feet (Fuß)}$$

Was ist hier passiert? Die Gleichung $l = 2\sqrt{\frac{a^2}{4} + d^2}$ ist eine Gleichung zwischen *Größen*, die eine *Dimension* haben, in diesem Fall die Dimension einer Länge. Um die Dimension einer Größe anzuzeigen, benutzen wir eckige Klammern, z.B. $[l] = L$. Die Dimensionen selber kürzen wir mit großen Buchstaben ab, L steht also für die Dimension *Länge*. Auch die rechte Seite hat die Dimension einer Länge, denn

$$\begin{aligned} \left[2\sqrt{\frac{a^2}{4} + d^2} \right] &= \left[\sqrt{\frac{a^2}{4} + d^2} \right] = \sqrt{\left[\frac{a^2}{4} \right] + [d^2]} \\ &= \sqrt{[a]^2 + [d]^2} = \sqrt{L^2 + L^2} = \sqrt{L^2} = L. \end{aligned}$$

Man beachte, dass für den Umgang mit Dimensionen recht eigentümliche Rechenregeln gelten, so ist z.B. $2L = L + L = L$. Ferner gelten sehr elementare Regeln für alle Modellbildungen:

Grundregeln für die Bildung mathematischer Modelle

- Konsistente Modelle dürfen nur Größen gleicher Dimension gleichsetzen.
- Es dürfen nur Größen gleicher Dimension addiert werden.
- Es dürfen Größen unterschiedlicher Dimensionen multipliziert werden.

Eine Gleichung der Form $l = a^2 + d^2$ wäre also von vornherein unsinnig, da eine Länge mit einem Längenquadrat gleichgesetzt wird. Für den Mathematikunterricht wurde nun die Gleichung $l = 2\sqrt{\frac{a^2}{4} + d^2}$ *entdimensionalisiert*. Dies geschieht durch die Einführung einer

dimensionslosen Variablen: In Deutschland nutzen wir in der Regel die Maßeinheit Meter, also

$$l = \lambda_D \cdot 1 \text{ m} \quad \text{bzw.} \quad \lambda_D = \frac{l}{1 \text{ m}}.$$

Man beachte, dass

$$[\lambda_D] = \left[\frac{l}{1 \text{ m}} \right] = \frac{L}{L} = 1,$$

dass λ_D also eine Zahl ist. In den USA dagegen hätten wir

$$l = \lambda_{\text{USA}} \cdot 1 \text{ feet} \quad \text{bzw.} \quad \lambda_{\text{USA}} = \frac{l}{1 \text{ feet}}.$$

Für λ_D gilt also:

$$\begin{aligned} \lambda_D &= \frac{l}{1 \text{ m}} = \frac{2\sqrt{\frac{a^2}{4} + d^2}}{1 \text{ m}} = 2\sqrt{\frac{1}{4} \left(\frac{a}{1 \text{ m}}\right)^2 + \left(\frac{d}{1 \text{ m}}\right)^2} \\ &= 2\sqrt{0,25 \cdot 137,60^2 + 54,20^2} = 175,2 \end{aligned}$$

Diese Entdimensionalisierung mag einem im Moment trivial und gekünstelt formalisiert erscheinen, wir werden jedoch gleich sehen, dass dieser Formalismus weiterhelfen kann.

1.1.2 An das Problem angepasste Skalierung und Simulation

Ist das Ergebnis der Modellierung nah an der Realität? Wir können das durch ein Experiment überprüfen. Müssen wir dafür ein Kabel zwischen 137.60 m entfernte Maste spannen? Gott sei Dank nicht. Durch eine geeignete *Skalierung* können wir das Ergebnis im verkleinerten Maßstab betrachten. Dazu drücken wir die Länge l nicht als Vielfaches von Metern aus, wie es die Variable λ_D macht, oder als Vielfaches von feet, wie es die Variable λ_{USA} macht, sondern als Vielfaches einer charakteristischen, dem Problem innewohnenden Länge $\bar{l}$, die wir gleich geschickt wählen werden. Wir werden dazu also nicht auf Standardmaßeinheiten zurückgreifen, die ja mit dem konkreten Problem eigentlich gar nichts zu tun haben, sondern suchen eine für das Problem typische Länge. Deshalb rechnen wir erst einmal mit einer unbestimmten Längenskala $\bar{l}$ und sagen erst am Schluss, was wir für $\bar{l}$ wählen.

$$\lambda = \frac{l}{\bar{l}} = 2\sqrt{\frac{1}{4} \left(\frac{a}{\bar{l}}\right)^2 + \left(\frac{d}{\bar{l}}\right)^2}$$

Wir wählen als charakteristische Länge $\bar{l} = a = 137,60 \text{ m}$ und erhalten die einfache Gleichung

$$\lambda_\varepsilon = 2 \cdot \sqrt{\frac{1}{4} + \varepsilon^2} \quad \text{mit} \quad \varepsilon := \frac{d}{a}.$$

ε ist also das Verhältnis von Durchhang zu Abstand und man beachte, dass $[\varepsilon] = 1$. λ_ε hängt also nur vom Verhältnis von d und a ab, aber nicht von d und a selber. In unserem Beispiel gilt $\varepsilon = 0.394$ und $\lambda_\varepsilon = 1.273$. Wir können das Experiment nun anstatt mit $a = 137.60$ m auch mit der handlichen Länge $a = 1$ m, inklusive Messgenauigkeit also $a = 1.00$ m, durchführen. Die Seillänge ist dann

$$l = 1.27 \cdot \bar{l} = 1.27 \cdot a = 1.27 \cdot 1.00 \text{ m} = 1.27 \text{ m}.$$

Wenn wir das Experiment tatsächlich durchführten, würden wir jedoch einen kleineren Durchhang messen als den vom Modell berechneten $d = \varepsilon a = 0.39$ m. Das Seil ist also zu kurz und die Modellierung ist nicht gut genug.

1.1.3 Modell 2: parabelförmiger Verlauf

Anstatt durch eine stückweise lineare Funktion versuchen wir den Höhenverlauf nun durch eine quadratische Funktion zu beschreiben, siehe Abb. 1.3. Durch geschickte Wahl des Koordinatenursprungs als den tiefsten Punkt des Seilverlaufs und durch Ausnutzung der der Situation innewohnenden Symmetrie erhalten wir als Funktionsterm

$$H(x) := kx^2. \quad (1.2)$$

Die Konstante k bestimmen wir aus der Forderung

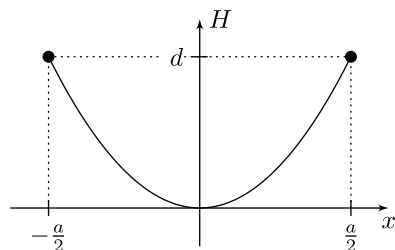
$$H\left(\frac{a}{2}\right) = d, \quad \text{also} \quad k = \frac{4d}{a^2}.$$

Wir wollen auch diese Gleichungen entdimensionalisieren. Wir haben zwei Variablen H und x , jeweils mit der Dimension einer Länge, also

$$[H] = [x] = L.$$

Wir könnten nun für die Höhe H und die horizontale Länge x unterschiedliche Skalen einführen. Da aber beide Skalen dieselbe Dimension haben, ist es zweckmäßiger, für beide

Abb. 1.3 Modellierung des Höhenverlaufs des Kabels mit einer quadratischen Funktion



Variablen ein und dieselbe Skala zu wählen. Wir führen die dimensionslosen Variablen mit den folgenden Gleichungen ein:

$$y := \frac{x}{\bar{l}}, \quad \text{also} \quad x = \bar{l}y \quad \text{sowie} \\ h(y) := \frac{H(x)}{\bar{l}} = \frac{H(\bar{l}y)}{\bar{l}}, \quad \text{also} \quad H(x) = \bar{l}h(y) = \bar{l}h\left(\frac{x}{\bar{l}}\right).$$

Die Funktion

$$h(y) = \frac{H(\bar{l}y)}{\bar{l}} = \frac{4d\bar{l}^2}{a^2\bar{l}} \cdot y^2$$

wird auf dem Intervall $I = \left[-\frac{a}{2\bar{l}}, \frac{a}{2\bar{l}} \right]$ betrachtet. Wir haben nun wieder alle Freiheiten, die Skala $\bar{l}$ zu wählen. Wir bleiben, auch um die Vergleichbarkeit mit der Modellierung 1 zu gewährleisten, bei der im Kontext leicht zu interpretierenden Wahl $\bar{l} = a$ und erhalten

$$h_\varepsilon(y) = 4\varepsilon y^2 \quad \text{sowie} \quad I = \left[-\frac{1}{2}, \frac{1}{2} \right] \quad \text{mit} \quad \varepsilon = \frac{d}{a}.$$

Wie erhalten wir die Länge des Seils? In unserem Modell entspricht die Länge des Seils der Bogenlänge der Kurve, die durch den Graphen unserer Höhenverlaufsfunktion beschrieben wird. Die Formel für die Bogenlänge eines Funktionsgraphen wird in Aufg. 1.1 thematisiert und in Abschn. 1.2 in einen größeren Zusammenhang gestellt. Für die dimensionslose Bogenlänge $\lambda = \lambda_\varepsilon$ ergibt sich

$$\lambda_\varepsilon = \int_{-\frac{1}{2}}^{\frac{1}{2}} \sqrt{1 + h'_\varepsilon(y)^2} \, dy = \int_{-\frac{1}{2}}^{\frac{1}{2}} \sqrt{1 + (8\varepsilon y)^2} \, dy, \quad (1.3)$$

und mit der Substitution $t = 8\varepsilon y$, unter Ausnutzung der Achsensymmetrie und mit Hilfe der Stammfunktion $\int^x \sqrt{1 + y^2} \, dy = \frac{1}{2} \left(x \cdot \sqrt{1 + x^2} + \operatorname{arsinh}(x) \right) + c$ folgt

$$\begin{aligned} \lambda_\varepsilon &= \frac{1}{8\varepsilon} \int_{-4\varepsilon}^{4\varepsilon} \sqrt{1 + t^2} \, dt \\ &= \frac{1}{4\varepsilon} \int_0^{4\varepsilon} \sqrt{1 + t^2} \, dt \\ &= \frac{1}{8\varepsilon} \left[4\varepsilon \cdot \sqrt{1 + 16\varepsilon^2} + \operatorname{arsinh}(4\varepsilon) \right] \\ &= \frac{1}{2} \sqrt{1 + 16\varepsilon^2} + \frac{1}{8\varepsilon} \operatorname{arsinh}(4\varepsilon). \end{aligned}$$

Die hyperbolischen Funktionen Sinus hyperbolicus und Kosinus hyperbolicus werden in Aufg. 1.2 thematisiert. Wir betrachten wieder unsere Beispielwerte $a = 137.60 \, \text{m}$ und $d = 54.20 \, \text{m}$ und haben mit $\varepsilon = \frac{d}{a} = \frac{54.20 \, \text{m}}{137.60 \, \text{m}} = 0.3939$. Für die dimensionslose Länge erhalten wir damit $\lambda_{0.3939} = 1.325$ im Vergleich zu 1.273 nach der ersten Modellierung. Das Seil

ist also ca. 4 % länger als nach der ersten Modellierung. Zurückskaliert muss das Seil nun also eine Länge von $l = 182.4$ m haben. Die Differenz in der Länge zwischen Modell 1 und Modell 2 beträgt also 7.19 m, eine Differenz, die durchaus von praktischer Relevanz sein kann, siehe auch Aufg. 1.4.

1.1.4 Modell 3: Die Physiker sagen, es sei eine Kettenlinie.

Bis hier waren unsere Modellierungen reine sogenannte *Ad-hoc*-Formulierungen. Die Physiker versichern uns nun, dass der Höhenverlauf des Kabels durch eine Kettenlinie, das heißt durch eine Funktion der Form

$$H(x) = k \cdot \cosh\left(\frac{x}{k}\right),$$

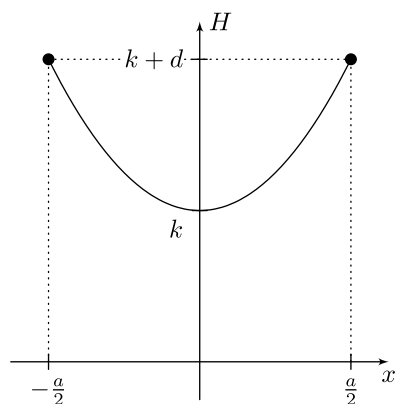
beschrieben werden soll (dabei ist $\cosh(x) = \frac{e^x + e^{-x}}{2}$ der Kosinushyperbolicus). Wir werden in Abschn. 2.2 sehen, dass dieses Modell kein Ad-hoc-Modell, sondern ein Modell aus *First Principles* ist. Um diese Modellfunktion zu nutzen, müssen wir allerdings unser Koordinatensystem verschieben, denn $H(0 \text{ m}) = k \neq 0 \text{ m}$, siehe Abb. 1.4.

Der Parameter k muss nun so gewählt werden, dass $H\left(\frac{a}{2}\right) = k + d$, also

$$k \cosh\left(\frac{a}{2k}\right) = k + d. \quad (1.4)$$

Die Gl. (1.4) lässt sich nicht explizit nach k auflösen, das k ist zunächst also nur implizit durch die Gl. (1.4) bestimmt. Wir wollen entdimensionalisieren und uns dabei direkt an das Schema gewöhnen:

Abb. 1.4 Koordinatensystem für die Modellierung des Höhenverlaufs mit dem Kosinushyperbolicus



- die entdimensionalisierten Variablen:

$$y = \frac{x}{\bar{l}}, \quad \bar{l}y = x \quad \text{und} \quad h(y) = \frac{H(x)}{\bar{l}}, \quad H(x) = \bar{l}h(y)$$

- die entdimensionalisierte Gleichung:

$$h(y) = \frac{k}{\bar{l}} \cosh\left(\frac{\bar{l}y}{k}\right)$$

- Wahl der Skala: Wir bleiben bei der bewährten Skalenwahl $\bar{l} = a$. Damit sorgen wir für eine leichtere Vergleichbarkeit mit den Modellen 1 und 2. Naheliegender wäre aber auch die Skala $\bar{l} = k$ gewesen.

Mit unserer Wahl und dem neuen dimensionslosen Verhältnis $\eta = \frac{a}{k}$ erhalten wir dann

$$h(y) = \frac{1}{\eta} \cosh(\eta y) \quad \text{in den Grenzen} \quad \left[-\frac{1}{2}, \frac{1}{2}\right].$$

Für die Bogenlänge ergibt sich

$$\lambda_\eta = \int_{-\frac{1}{2}}^{\frac{1}{2}} \sqrt{1 + \sinh^2(\eta y)} \, dy = \frac{2}{\eta} \sinh(z) \Big|_0^{\frac{\eta}{2}} = \frac{2}{\eta} \sinh\left(\frac{\eta}{2}\right). \quad (1.5)$$

Es bleibt das Problem, k und damit auch η zu bestimmen: Aus (1.4) erhalten wir durch Division durch k

$$\cosh(\eta/2) = 1 + \varepsilon \eta \quad (1.6)$$

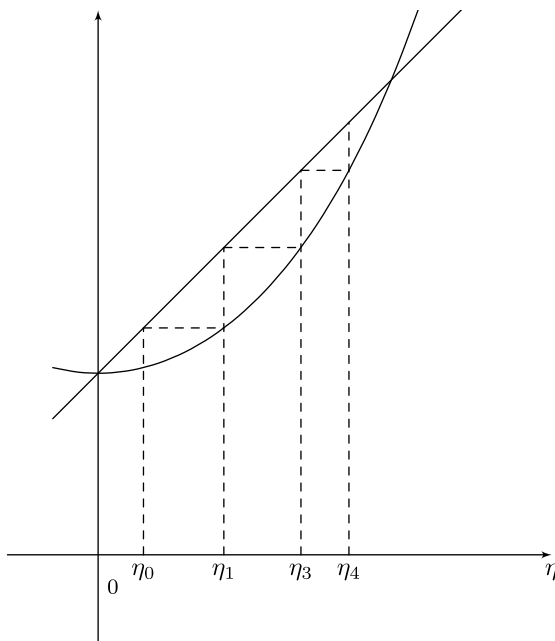
mit $\varepsilon = \frac{d}{a}$ wie vorher. Diese Gleichung lässt sich nicht analytisch, sondern nur numerisch lösen. Durch eine geeignete Kurvendiskussion sehen wir, dass die Gl. (1.6) genau eine positive Nullstelle η^* besitzt, die wir durch das folgende Iterationsschema approximieren können, siehe Abb. 1.5: Wir wählen einen Startwert η_0 und berechnen $\gamma_0 := 1 + \varepsilon \eta_0$. Dann setzen wir $\eta_1 = 2 \operatorname{arcosh}(\gamma_0)$ und berechnen $\gamma_1 = 1 + \varepsilon \eta_1$. Allgemein definieren wir

$$\gamma_n = 1 + \varepsilon \eta_n \quad \text{und} \quad \eta_{n+1} = 2 \operatorname{arcosh}(\gamma_n).$$

Die theoretische Konvergenz dieses Verfahrens soll in Aufg. 1.3 gezeigt werden. Wir implementieren diesen Algorithmus in einen Rechner, z.B. mit Hilfe eines Tabellenkalkulationsprogramms, und erhalten für $\varepsilon = \frac{54.20}{137.60}$ nach einigen Iterationsschritten den auf vier Stellen genauen Wert $\eta = 1.710$. Diesen Wert setzen wir ein und erhalten:

$$\begin{aligned} \lambda_\eta &= \frac{2}{\eta} \sinh(\eta/2) = 1.335 & k_\eta &= \frac{a}{\eta} = 50.77 \, \text{m} \\ l &= \lambda_\eta \cdot a = 183.8 \, \text{m} \end{aligned}$$

Abb. 1.5 Das Iterationsschema zur Bestimmung des Parameters η : η_0 wird als Startwert vorgegeben und für alle $n \in \mathbb{N}_0$ setzen wir $\gamma_n = 1 + \varepsilon \eta_n$ sowie $\eta_n = \operatorname{arcosh}(\gamma_n)$



Im Ergebnis unterscheiden sich die Längen nach Modellierungen 2 und 3 lediglich um circa 0.8 % bzw. 1.39 m. Dieser Unterschied würde wohl in den meisten konkreten Anwendungsfällen keine Rolle spielen.

1.1.5 Verallgemeinerung: Wie sieht es bei anderen Verhältnissen von Abstand und Durchhang aus?

Wir haben gesehen, dass das entdimensionalisierte Problem nur von dem Verhältnis ε von Durchhang d und Abstand a abhängt. Unser konkretes Problem hatte das Verhältnis $\varepsilon = 0.3939$, und zwischen Modell 2 und 3 gibt es im Ergebnis keinen für das Problem relevanten Unterschied mehr: Nach Modell 2 erhalten wir $l = 182.1$ m und nach Modell 3 $l = 183.8$ m, also einen relativen Unterschied von ca. 0.8 %. Es stellt sich die Frage, ob beide Modelle immer praktisch gleichwertig sind oder ob sich bei anderen Verhältnissen ε die Ergebnisse der Modelle wesentlich unterscheiden. In diesem Fall könnten wir durch ein Experiment entscheiden, welches der beiden Modelle das genauere ist. Für einen anschaulichen Überblick implementieren wir alle drei Modelle in eine GeoGebra-Datei und machen das Verhältnis ε mit Hilfe eines Schiebereglers variabel. Zunächst müssen wir aber den Höhenverlauf nach Modell 3 noch geeignet verschieben, damit in allen Modellen das Kabel an den Punkten $(-0,5/\varepsilon)$ und $(0,5/\varepsilon)$ aufgehängt ist und seinen tiefsten Punkt im Ursprung hat: $h_{\text{III}}(y) = \frac{1}{\eta} \cosh \eta y - \frac{1}{\eta}$. Wir sehen, dass für größer werdende ε der Unterschied

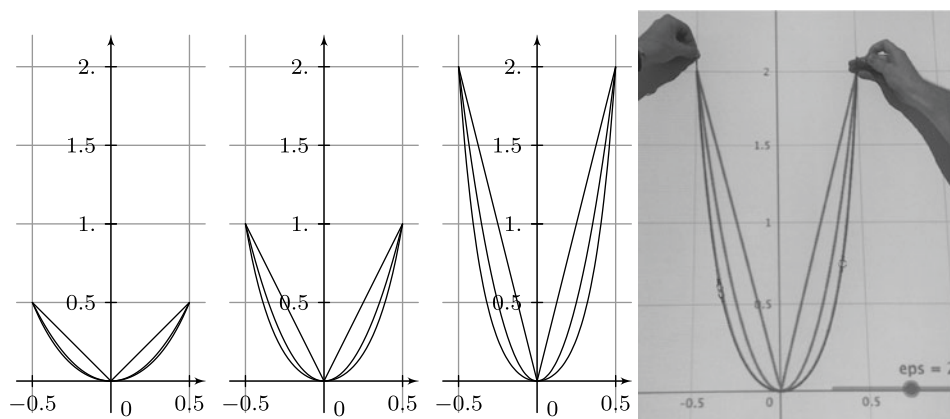


Abb. 1.6 Modell 1, 2 und 3 im Vergleich für $\varepsilon = 0,5, 1$ und 2 . Rechts: Im Experiment gibt Modell 3 den Höhenverlauf augenscheinlich „perfekt“ wieder: Bei einem Verhältnis von $\varepsilon = 2$ wird eine Kette vor den Plot der Parabel und Kettenlinien gehalten

zwischen den beiden Höhenverläufen zunimmt, siehe Abb. 1.6. Für $\varepsilon = 2$ testen wir mit einer Kugelschleife, wie man sie für Abflusstöpfen gebraucht, welcher Höhenverlauf die Form der Kette am besten wiedergibt. Die experimentelle Antwort ist eindeutig: Die Form der Kette wird durch den Graphen des hyperbolischen Kosinus (Modell 3) nach Augenschein perfekt wiedergegeben, siehe Abb. 1.6. Es bleibt die Frage, woher die Physiker das wussten. Haben sie viel experimentiert oder sehr schlaue Gerate? Wir werden in Abschn. 2.2 sehen, dass sie erstaunlich erfolgreiche *Prinzipien* haben.

1.1.6 Reflexionen zum mathematischen Modellieren I: der Modellbegriff und Gütekriterien für Modelle

Was ist nun das *mathematische Modellieren*? Wir wollen hier keine wissenschaftliche Untersuchung über die Genese dieses Begriffs anstellen, sondern lassen zu Beginn den von Ortlieb in [69] als *Geburtshelfer* des modernen Modellbegriffs bezeichneten Physiker Heinrich Hertz zu Worte kommen, zunächst zum Modellbegriff und den Zielen des Modellierens und später zu den Gütekriterien für ein Modell, siehe [39, S. 1 ff.]:

„Wir machen uns innere Scheinbilder oder Symbole der äußeren Gegenstände, und zwar machen wir sie von solcher Art, daß die denotwendigen Folgen der Bilder stets wieder die Bilder seien von den naturnotwendigen Folgen der abgebildeten Gegenstände. Damit diese Forderung überhaupt erfüllbar sei, müssen gewisse Übereinstimmungen vorhanden sein zwischen der Natur und unserem Geiste. Die Erfahrung lehrt uns, dass diese Forderung erfüllbar ist und dass also solche Übereinstimmungen in der Tat bestehen. Ist es uns einmal geglückt, aus der angesammelten bisherigen Erfahrung Bilder von der verlangten Beschaffenheit abzuleiten,

so können wir an ihnen, wie an Modellen, in kurzer Zeit die Folgen entwickeln, welche in der äußeren Welt nur in längerer Zeit oder als Folgen unseres eigenen Eingreifens auftreten werden; wir vermögen so den Tatsachen vorauszueilen und können nach der gewonnenen Einsicht unsere gegenwärtigen Beschlüsse richten. Die Bilder, von denen wir reden, sind unsere Vorstellungen von den Dingen; sie haben mit den Dingen die eine wesentliche Übereinstimmung, welche in der Erfüllung der genannten Forderung liegt, es ist aber für ihren Zweck nicht nötig, dass sie irgend eine weitere Übereinstimmung mit den Dingen haben.“

Was Hertz mit „inneren Scheinbildern oder Symbolen“ meint, nennt man heute mathematische Modelle, und wir wollen für die Belange dieses Buches unter dem mathematischen Modellieren die Untersuchung realer Objekte und die Bearbeitung damit verbundener Problemstellungen durch die Auswahl (oder ggf. Konstruktion), Untersuchung, Interpretation und Validierung mathematischer Modelle verstehen. Dabei ist ein mathematisches Modell ein für die konkreten Zwecke ausgesuchter Teilbereich der Mathematik. Mit Auswählen meinen wir, dass ein bereits bekanntes Teilgebiet der Mathematik als mathematisches Modell ausgewählt wird, mit Konstruieren, dass ein noch nicht bekanntes Teilgebiet der Mathematik neu entdeckt (oder konstruiert) werden muss. Der Einfachheit halber sprechen wir jetzt nur noch vom Auswählen, meinen aber ggf. auch das Konstruieren. Die Beziehungen der benutzten Begriffe lassen sich gut in einem Diagramm wie in Abb. 1.7 darstellen, wobei Übersetzungstätigkeiten zwischen Realität und Mathematik, das Auswählen und Interpretieren nicht nur einmalig erfolgen, sondern den gesamten Bearbeitungsprozess durchsetzen können. So ist ja durch die Wahl eines Modells noch nicht festgelegt, mit welchen Methoden das mathematische Problem bearbeitet werden soll. Dies wird man im Hinblick auf die ursprüngliche Problemstellung hin auswählen. Genauso wird nicht nur das Endergebnis der mathematischen Untersuchung interpretiert werden, sondern häufig auch bereits das gewählte Modell selbst oder etwaige Zwischenschritte und Zwischenergebnisse: Beim mathematischen Bearbeitungsprozess wird die Herkunft des Problems in der Regel mitbedacht.

Wir gehen somit von einer Dichotomie bestehend aus „der Realität“ und „der Mathematik“ aus, in der nach gewissen Kriterien einem Teilbereich der Realität ein Teilbereich der Mathematik zugeordnet wird. Es wird also nicht der Ausschnitt der Realität mit dem Ausschnitt der Mathematik identifiziert, sondern lediglich eine Zuordnung getroffen, die natürlich nicht beliebig, aber auch keineswegs eindeutig ist. Ein hängendes Kabel „ist“ keine Parabel oder Kettenlinie, sondern wird dadurch nur mehr oder weniger gut beschrieben.

Welche Kriterien sind nun bei der Auswahl eines mathematischen Modells wesentlich? Zunächst wieder Hertz [39, S. 2 ff.]:

„Verschiedene Bilder derselben Gegenstände sind möglich und diese Bilder können sich nach verschiedenen Richtungen unterscheiden. Als unzulässig sollten wir von vorneherein solche Bilder bezeichnen, welche schon einen Widerspruch gegen die Gesetze unseres Denkens in sich tragen, und wir fordern also zunächst, daß alle Bilder logisch zulässige oder kurz zulässige seien. Unrichtig nennen wir zulässige Bilder dann, wenn ihre wesentlichen Beziehungen den Beziehungen der äußeren Dinge widersprechen, das heißt, wenn sie jener ersten Grundfor-

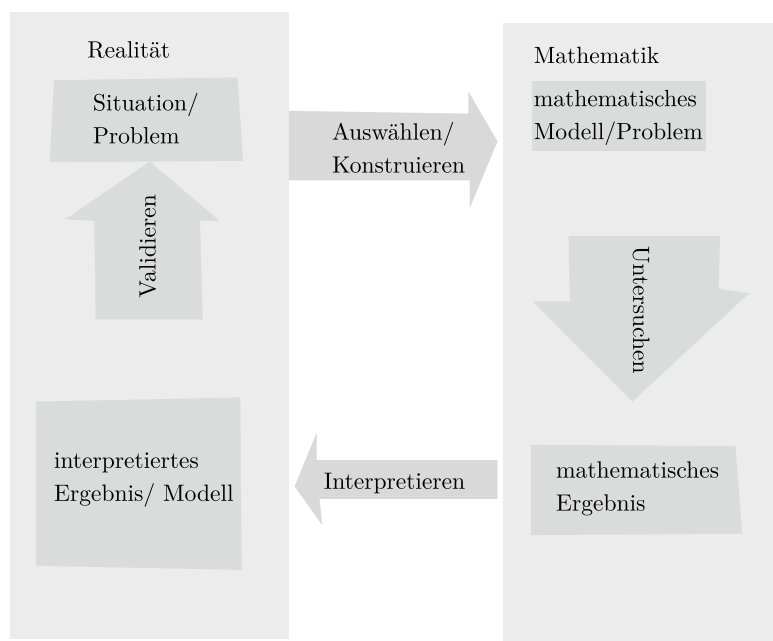


Abb. 1.7 Der Prozess des mathematischen Modellierens. Weiter differenzierte Modellierungskreisläufe finden sich in [27]

derung nicht genügen. Wir verlangen demnach, dass unsere Bilder richtig seien. Aber zwei zulässige und richtige Bilder derselben äußeren Gegenstände können sich noch unterscheiden nach der Zweckmäßigkeit. Von zwei Bildern desselben Gegenstandes wird dasjenige das zweckmäßigere sein, welches mehr wesentliche Beziehungen des Gegenstandes widerspiegelt als das andere; welches, wie wir sagen wollen, das deutlichere ist. Bei gleicher Deutlichkeit wird von zwei Bildern dasjenige zweckmäßiger sein, welches neben den wesentlichen Zügen die geringere Zahl überflüssiger oder leerer Beziehungen enthält, welches also das einfachere ist.“

Wir fassen zusammen: Modelle kann man beurteilen . . .

- nach ihrer mathematischen Wohldefiniertheit (bei komplexen Modellierungen kann das durchaus problematisch sein).
- nach ihrer Richtigkeit, also der Übereinstimmung des interpretierten mathematischen Ergebnisses mit der zu untersuchenden Eigenschaft.
- nach dem Gültigkeitsbereich des Modells. Unter welchen Bedingungen und in welchen Eigenschaften stimmen die interpretierten Ergebnisse mit der Realsituation überein?
- nach der Einfachheit des mathematischen Modells.

Betrachten wir vor diesem Hintergrund unsere drei Modelle, so sind sicherlich alle drei wohldefiniert. Ob sie *richtig* (in Bezug auf die Kabellänge) sind, kann nicht absolut beantwortet werden. Das kommt darauf an, mit welcher Genauigkeit das Ergebnis benötigt wird. Es gibt eine klare Reihenfolge: Modell 3 ist richtiger als Modell 2 ist richtiger als Modell 1. Zudem hat Modell 3 den größten Gültigkeitsbereich, liefert in mehr Situationen unterschiedlichen Durchhangs und Abstands relativ richtige Ergebnisse. Am einfachsten ist zweifellos das erste Modell.

Welches Modell ist nun das beste? Diese Frage lässt sich nicht absolut beantworten. Braucht man ganz schnell eine grobe Längenbestimmung in einer konkreten Situation, ist wohl Modell 1 das Modell der Wahl; benötigt man ein möglichst präzises bei unterschiedlichen Verhältnissen von Durchhang und Abstand gültiges Modell, wird man Modell 3 wählen.

1.2 Trassierungsaufgaben und der Begriff der Krümmung

In der Schule gehören Trassierungsaufgaben zu den typischen Steckbriefaufgaben. Der Prototyp sieht folgendermaßen aus: Zwei in einem Koordinatensystem durch Gleichungen beschriebene Straßenstücke sollen miteinander verbunden werden. Dazu soll in dem gegebenen Koordinatensystem eine ganzrationale Funktion gesucht werden, welche die beiden gegebenen Stücke knickfrei und „krümmungsruckfrei“ verbindet: Die Funktionswerte, die Ableitungen und die zweiten Ableitungen an den Nahtstellen sollen übereinstimmen. Das führt dann auf hinreichend viele Forderungen, welche die gesuchte ganzrationale Funktion eindeutig bestimmen.

Aber wie werden Straßen- oder auch Schienenverläufe tatsächlich geplant? Im Allgemeinen spielen auch Höhenunterschiede bei der Trassenplanung eine Rolle, von denen wir aber hier absehen wollen, um das mathematische Setting zweidimensional zu halten. Für die Planung der Trassierung gehen wir von den erwünschten Befahrungseigenschaften aus. Damit die Straße gut und schnell befahrbar ist, muss der Straßenverlauf so beschaffen sein, dass keine schnellen oder unregelmäßigen Lenkbewegungen und keine zu großen Lenkradeinschläge nötig sind. Folgende Forderungen werden aufgestellt:

1. Entweder soll der Lenkradeinschlag konstant gehalten werden können, das läuft auf ein Geradenstück oder auf einen Kreisbogen hinaus, oder
2. das Lenkrad soll beim Durchfahren mit konstantem Fahrtempo mit konstantem Drehtempo von einem Lenkradeinschlag zum nächsten gedreht werden können.

Ein solches Zeit-Lenkradeinschlags-Diagramm ist in Abb. 1.8 gezeigt. Vermieden werden sollen Straßenverläufe, bei denen quasi instantan von einem Lenkradeinschlag zum nächsten gewechselt wird, oder sich das Tempo des Lenkradeinschlags ändern muss. (Als Autofahrer

Lenkradeinschlag nach

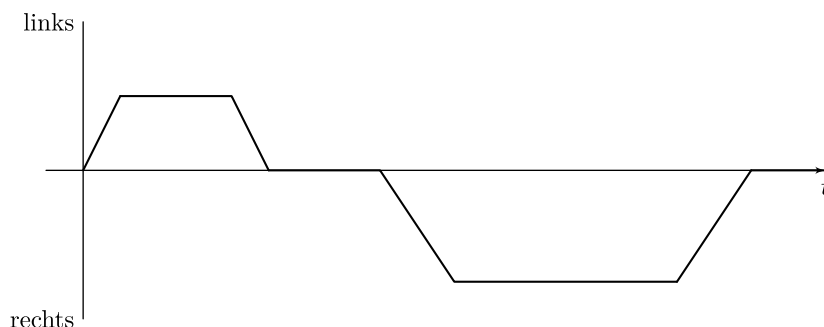


Abb. 1.8 Damit eine Straße gut befahrbar ist, soll die Geschwindigkeit des Lenkradeinschlags konstant sein

fallen einem hin und wieder Autobahnauffahrten und -kreuze auf, bei denen diese Forderungen nicht erfüllt sind, was zu unangenehmen Situationen führen kann.)

Der zugehörige Straßenverlauf besteht aus Geradenstücken und Kreisbögen (konstanter Lenkradeinschlag) sowie Stücken, bei denen das Lenkrad mit konstantem Tempo gedreht werden muss. Doch wie sehen solche Straßenstücke genau aus?

1.2.1 Kurven und ihre Krümmung

Je stärker man das Lenkrad eingeschlagen hat, desto stärker ist die Kurve, die durchfahren wird, „gekrümmt“. Wir wollen also eine mathematische Modellierung für die anschauliche Vorstellung von der Krümmung einer Kurve finden.

Schulisch wird häufig davon gesprochen, dass die zweite Ableitung einer Funktion das Krümmungsverhalten des zugehörigen Funktionsgraphen beschreibt. Schauen wir uns den Graphen der Funktion $f : x \mapsto x^2$ an, sprich die Parabel, so ist anschaulich klar, dass diese an ihrem Scheitelpunkt stärker gekrümmt ist, als weit außen an ihren Ästen. Stellt man sich vor, man würde sie durchfahren, wäre das Lenkrad am Scheitelpunkt maximal eingeschlagen, in einer großen Entfernung vom Scheitelpunkt fast ohne Einschlag. Die zweite Ableitung der Funktion f ist jedoch konstant, sie kann also nicht (direkt) die Mathematisierung der Anschauung von der Krümmung einer Kurve liefern. Die Mathematisierung sollte mindestens das Folgende liefern: Geraden sollen die Krümmung null erhalten (das liefert die zweite Ableitung). Ein Kreisbogen soll eine von null verschieden konstante Krümmung haben, die umso größer ist, je kleiner der Radius des zugehörigen Kreises ist (das erfüllt die zweite Ableitung nicht).

Wir werden im nächsten Abschnitt eine elementare Kurventheorie zur Verfügung stellen, die einen Krümmungsbegriff enthält, der das Gewünschte liefert. In der Darstellung und in

den Bezeichnungen orientieren wir uns an [56, Kap. 2], beziehen aber alles auf den Kontext „Fahrt auf einer Straße“.

1.2.2 Elementare Theorie ebener Kurven

In unserem Kontext stellen wir uns unter $\gamma(t)$ den Ort eines Autos zur Zeit t vor. Wir starten mit einer Definition:

Definition 1.1

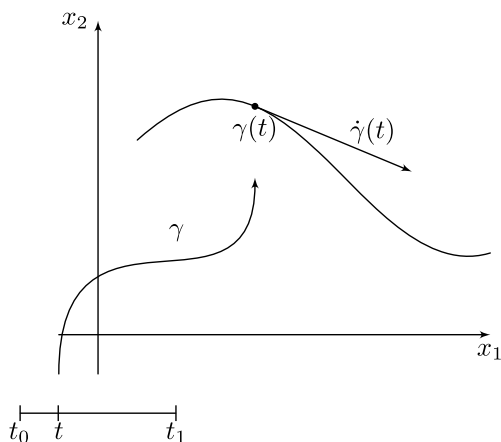
Unter einer regulären parametrisierten Kurve (einer regulären Fahrt) verstehen wir eine stetig differenzierbare injektive Abbildung γ von einem Intervall $[t_0, t_1] \subset \mathbb{R}$ in den $\mathbb{R}^2$ mit $\dot{\gamma}(t) = \frac{d\gamma}{dt}(t) \neq 0$ für alle $t \in [t_0, t_1]$. ♦

Der Vektor $\dot{\gamma}(t) \in \mathbb{R}^2$ gibt den *Geschwindigkeitsvektor* zur Zeit t an. Dieser ist als lineare Approximation tangential an die Kurve, siehe Abb. 1.9. (Im Gegensatz zu den *Verschiebungsvektoren* der analytischen Geometrie, welche eine Verschiebung der gesamten Ebene beschreiben, sind die *Geschwindigkeitsvektoren* an ihren Aufpunkt angeheftet.) Sie haben die Dimension Geschwindigkeit. Die Forderung $\dot{\gamma} \neq 0$ besagt im Kontext, dass das Auto zu keiner Zeit steht. Die zu einer Parametrisierung gehörende Kurve (im Anwendungskontext: die Straße) bildet sich aus den Bildpunkten:

$$C_\gamma := \{\gamma(t) : t \in [t_0, t_1]\}$$

Ein und dieselbe Straße kann man auf viele Arten und Weisen regulär in derselben Richtung durchfahren, entsprechend gibt es viele reguläre Parametrisierungen, welche dieselbe Kurve

Abb. 1.9 Regulär parametrisierte Kurve und ein Geschwindigkeitsvektor



(Straße) erzeugen. Wir fassen Parametrisierungen, welche dieselbe Kurve in der gleichen Richtung durchlaufen, in Äquivalenzklassen zusammen:

Definition 1.2

Zwei reguläre Parametrisierungen $\gamma : [t_0, t_1] \rightarrow \mathbb{R}^2$ und $\tilde{\gamma} : [\tau_0, \tau_1] \rightarrow \mathbb{R}^2$ heißen äquivalent, wenn es eine stetig differenzierbare Funktion $\varphi : [\tau_0, \tau_1] \rightarrow [t_0, t_1]$ mit $\dot{\varphi} > 0$ gibt (die *Parametertransformation*), sodass $\tilde{\gamma} = \gamma \circ \varphi$ gilt, siehe Abb. 1.10. ♦

Die Forderung $\dot{\varphi} > 0$ sorgt dafür, dass die Kurve in derselben Richtung durchfahren wird. In Aufg 1.6 soll nachgerechnet werden, dass es sich dabei tatsächlich um eine Äquivalenzrelation handelt. Für die Parametertransformation muss dann $\varphi = \gamma^{-1} \circ \tilde{\gamma}$ gelten. Wir definieren jetzt ganz offiziell eine reguläre orientierte Kurve $\mathcal{C}$ als Äquivalenzklasse einer regulären Parametrisierung unter dieser Äquivalenzrelation. Die Kurve $\mathcal{C}$ beschreibt also die Straße mit Auswahl einer Richtung, jedoch unabhängig davon, wie sie durchfahren wird.

Eine Straße hat eine Länge und auch bei Autofahrten ist es interessant, welche Strecke bereits zurückgelegt wurde: $\dot{\gamma}$ gibt die Geschwindigkeit an, also die momentane Richtung und das momentane Tempo. Das Tempo alleine wird durch den Betrag $|\dot{\gamma}|$ beschrieben. Aus dem Tempo einer Fahrt lässt sich die zurückgelegte Strecke durch Integration rekonstruieren: Wir definieren

$$s(t) = \int_{t_0}^t |\dot{\gamma}(\tau)| d\tau. \quad (1.7)$$

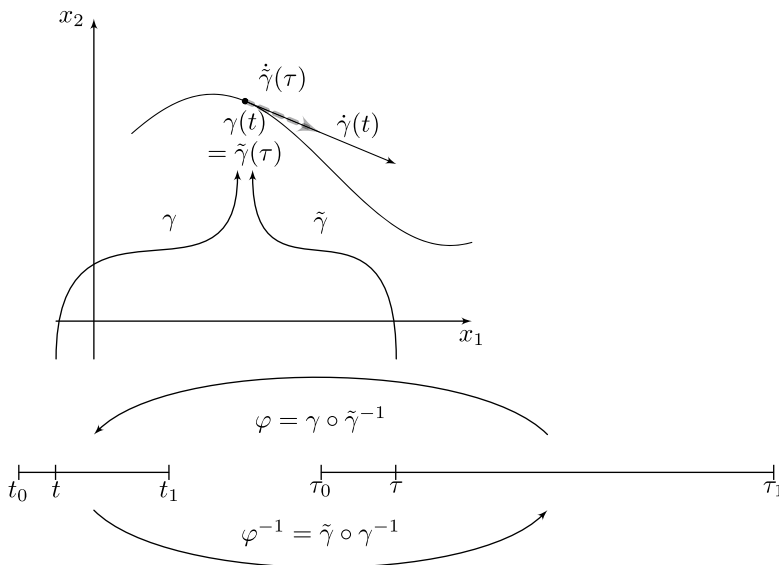


Abb. 1.10 Zwei äquivalente Parametrisierungen mit zugehörigen Transformationsabbildungen

Im Kontext hat s als Integral einer Geschwindigkeit über einer Zeit die Dimension einer Länge. Kontextunabhängig bezeichnet man s als die Bogenlänge. Wir definieren die Länge der Kurve $\mathcal{C}$ als

$$L(\mathcal{C}) = \int_{t_0}^{t_1} |\dot{\gamma}(\tau)| d\tau,$$

wobei γ eine beliebige Parametrisierung von $\mathcal{C}$ ist. Die Länge einer Straße sollte unabhängig sein von der Art und Weise, in der man sie durchfährt: In Aufg. 1.7 soll nachgerechnet werden, dass die angegebene Definition der Bogenlänge tatsächlich unabhängig vom gewählten Repräsentanten ist.

Der Strecken- oder Bogenlängenparameter s gibt uns die Möglichkeit die Kurve auch anders, quasi normiert, darzustellen. Dazu definieren wir die *Parametrisierung nach Bogenlänge*:

$$c : [0, L] \longrightarrow \mathbb{R}^2 \quad s \mapsto \gamma(t(s)), \quad (1.8)$$

wobei $s \mapsto t(s)$ die Umkehrfunktion von $t \mapsto \int_{t_0}^t |\dot{\gamma}(\tau)| d\tau$ ist, siehe die Bemerkung unten. Der zugehörige Ableitungsvektor ist dann normiert und dimensionslos, er ist quasi „nur Richtung“:

$$c'(s) = \dot{\gamma}(t(s)) \cdot t'(s) = \frac{\dot{\gamma}(t(s))}{\dot{s}(t(s))} = \frac{\dot{\gamma}(t(s))}{|\dot{\gamma}(t(s))|} \quad (1.9)$$

In Aufg. 1.7 soll ebenfalls nachgerechnet werden, dass die Parametrisierung nach Bogenlänge für eine reguläre orientierte Kurve eindeutig ist, also nicht vom gewählten Repräsentanten abhängt. Zur Beschreibung der Kurve ist der Bogenlängenparameter in besonderer Art und Weise geeignet, wie wir gleich sehen werden.

Zur Bezeichnung von Variablen und Funktionen

In (1.9) kommen Ausdrücke wie $\dot{s}(t(s))$ vor. Hier nehmen die Buchstaben s und t offenbar unterschiedliche Rollen ein: Einmal sind sie Variablen, die als Argumente in Funktionen eingesetzt werden können, andererseits bezeichnen sie aber auch Funktionen: So ist ursprünglich s in (1.7) als Funktion eingeführt worden. Als Bezeichnung für das Argument wurde t gewählt, da im Kontext eine Zeit gemeint ist. In Formel (1.9) taucht auf einmal der Ausdruck $t(s)$ auf, hier ist nun t der Name der Funktion und s das Argument, das in die Funktion eingesetzt wird. In den Einstiegsvorlesungen zur Analysis würde man so etwas strikt vermeiden und jeder Funktion einen neuen Namen geben, z.B. $\tilde{t} := s^{-1}$, und der Bogenlängenparameter würde vielleicht den Namen σ erhalten. Aus (1.8) würde

$$c(\sigma) = \gamma(\tilde{t}(\sigma))$$

und aus der Gleichungskette (1.9) entstünde

$$c'(\sigma) = \dot{\gamma}(\tilde{t}(\sigma)) \cdot \tilde{t}'(\sigma) = \frac{\dot{\gamma}(\tilde{t}(\sigma))}{\dot{s}(\tilde{t}(\sigma))} = \frac{\dot{\gamma}(\tilde{t}(\sigma))}{|\dot{\gamma}(\tilde{t}(\sigma))|}.$$

Im Kontext der Differentialgeometrie und auch in den Anwendungswissenschaften ist es üblich und praktisch, nicht auf diese Art zwischen Variablen- und Funktionennamen zu unterscheiden. s und t sind hier nicht bedeutungslose Buchstaben, sondern sie stehen für geometrische oder physikalische Größen, die in einem funktionalen Zusammenhang stehen. Unser Fall bezieht sich auf eine Bewegung, in der die Zeit t und die zurückgelegte Strecke s in einem eindeutigen funktionalen Zusammenhang stehen, die Strecke also als Funktion der Zeit ausgedrückt werden kann, $s(t)$, oder auch umgekehrt die Zeit als Funktion der Strecke, $t(s)$, beschrieben werden kann. Wir könnten auch so schöne Verkettungen wie $s(t(s))$ oder $t(s(t))$ aufstellen. Damit meinen wir dann die identische Funktion auf dem Längen- bzw. auf dem Zeitgrößenbereich.

Ableitungen nach der Zeit und Ableitungen nach der Länge

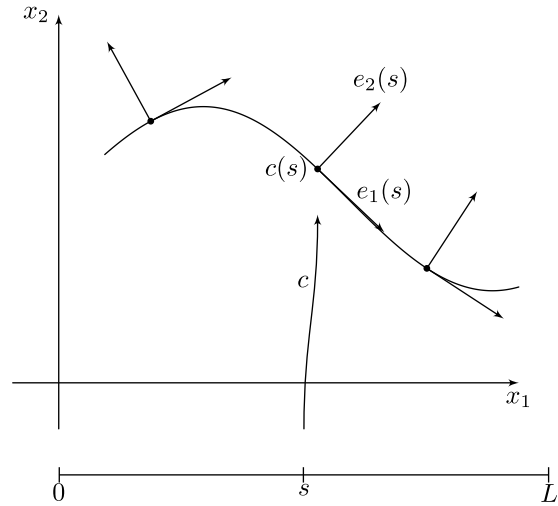
In diesem Buch werden wir prinzipiell für Ableitungen nach einer Zeit den Punkt verwenden, z. B. $\dot{s}(t)$, und für Ableitungen nach einer Länge den Strich, z. B. $t'(s)$.

Wir nutzen jetzt den Tangenteneinheitsvektor c' , um längs der Kurve ein sogenanntes Frenet'sches Zweibein einzuführen, das uns helfen wird, den Begriff der Krümmung mathematisch zu fassen. Dazu bezeichnen wir den Tangenteneinheitsvektor $c'(s) = e_1(s)$ und bezeichnen den um 90° gedrehten Vektor mit $e_2(s)$. In den kanonischen Koordinaten des $\mathbb{R}^2$ gilt also $e_2(s) = \begin{pmatrix} -c'_2(s) \\ c'_1(s) \end{pmatrix}$. Die Vektoren $e_1(s)$, $e_2(s)$ bilden das Frenet'sche Zweibein am Punkt $c(s)$, siehe Abb. 1.11. Ferner sei $\langle \cdot, \cdot \rangle$ das übliche Standardskalarprodukt im $\mathbb{R}^2$, also $\langle v, w \rangle = v_1 w_1 + v_2 w_2$. Wir können dann jeden Vektor in die Richtungen $e_1(s)$ und $e_2(s)$ zerlegen:

$$v = \langle v, e_1(s) \rangle e_1(s) + \langle v, e_2(s) \rangle e_2(s)$$

Für die weiteren Überlegungen benötigen wir etwas mehr Regularität: Im weiteren Verlauf gehen wir davon aus, dass alle auftretenden parametrisierten Kurven zweimal stetig differenzierbar sind. Weil $e_1(s)$ und $e_2(s)$ für alle s zum einen Einheitsvektoren und zum anderen orthogonal zueinander sind, erhalten wir durch Differentiation nach s drei Gleichungen:

Abb. 1.11 Eine reguläre Kurve mit Frenet'schem Zweibein



$$\begin{aligned}
 0 &= \frac{d}{ds} (\langle e_1(s), e_1(s) \rangle) = 2\langle e'_1(s), e_1(s) \rangle, \\
 0 &= \frac{d}{ds} (\langle e_2(s), e_2(s) \rangle) = 2\langle e'_2(s), e_2(s) \rangle, \\
 0 &= \frac{d}{ds} (\langle e_1(s), e_2(s) \rangle) = \langle e'_1(s), e_2(s) \rangle + \langle e_1(s), e'_2(s) \rangle.
 \end{aligned} \tag{1.10}$$

Zerlegen wir nun e'_1 in die Richtungen e_1 und e_2 , so erhalten wir aus der ersten Gleichung, dass e'_1 ein lineares Vielfaches von e_2 ist, der skalare Vorfaktor liefert uns gerade die gesuchte Krümmung:

$$e'_1 = \langle e'_1, e_2 \rangle e_2 = \kappa e_2 \quad \text{mit} \quad \kappa(s) = \langle e'_1(s), e_2(s) \rangle. \tag{1.11}$$

Die Krümmung gibt also die Rate an, mit der sich der Tangenteneinheitsvektor e_1 in Richtung des Normaleneinheitsvektors e_2 verändert. Hier erhalten wir die Belohnung für die Mühe, die Kurve nach Bogenlänge zu parametrisieren: Hätten wir irgendeine Parametrisierung genommen, hätte die Ableitung des Geschwindigkeitsvektors nach der Zeit auch eine Komponente in Richtung dieses Geschwindigkeitsvektors und nicht nur normal dazu. Die Krümmung hat übrigens die Dimension $1/L$, denn s hat die Dimension Länge und $e_1 = c'$ ist dimensionslos. Aus der zweiten und dritten Gleichung in (1.10) erhalten wir auch eine Gleichung für die Ableitung des Normaleneinheitsvektors e_2 nach der Bogenlänge:

$$e'_2 = \langle e'_2, e_1 \rangle e_1 + \langle e'_2, e_2 \rangle e_2 = -\langle e'_1, e_2 \rangle e_1 = -\kappa e_1 \tag{1.12}$$

Die beiden Gl. (1.11) und (1.12) zusammen nennt man die Frenet'schen Gleichungen in der Ebene nach dem Mathematiker Jean Frédéric Frenet (1816–1900), der auch erstmalig das Frenet'sche Zweibein benutzt hat.

Haben wir tatsächlich die Krümmung gefunden, die wir suchten? Dann sollten Strecken und Kreisbögen die einzigen Kurven mit konstanter Krümmung sein. Das soll in Aufg. 1.8 bestätigt werden. Muss, um die Krümmung zu berechnen, immer erst die Parametrisierung nach Bogenlänge bestimmt werden oder lässt sich die Krümmung der Kurve auch bezüglich einer beliebigen regulären Parametrisierung bestimmen? Dieser Frage wird in Aufg. 1.9 nachgegangen und eine Krümmungsformel für beliebige Parametrisierungen berechnet.

Unser Trassierungsproblem lautet nun folgendermaßen: Wir wissen, wie sich die Krümmung der Straße entwickeln soll, nämlich proportional zur Straßenlänge s , aber wir kennen den Straßenverlauf noch nicht. Übersetzt bedeutet das: Gegeben sei eine Krümmungsfunktion $s \mapsto \kappa(s)$, bestimme dazu eine Kurve mit dieser Krümmung. Diese Aufgabe hat eine große Ähnlichkeit zur Aufgabe, eine Funktion aus ihrer Ableitung zu rekonstruieren oder eine Bewegung aus ihrer Beschleunigung, und ist dementsprechend durch eine geeignete Integration lösbar.

Wir starten rückwärts und versuchen eine gegebene Parametrisierung nach Bogenlänge durch die Krümmung auszudrücken. Den gefundenen Ausdruck nutzen wir anschließend, um zu einer vorgegebenen Krümmungsfunktion eine zugehörige Kurve anzugeben. Den Tangenteneinheitsvektor $c'(s) = e_1(s)$ können wir auch durch den Winkel $\alpha(s)$ charakterisieren, den dieser mit der x_1 -Achse einschließt:

$$c'(s) = e_1(s) = \begin{pmatrix} \cos(\alpha(s)) \\ \sin(\alpha(s)) \end{pmatrix}$$

Der Winkel $\alpha(s)$ ist nur bis auf Addition von 2π eindeutig bestimmt und wir können ihn so wählen, dass die Funktion $s \mapsto \alpha(s)$ stetig differenzierbar auf dem Intervall $[0, L]$ ist. Nach Definition des Frenet'schen Zweibeins gilt

$$e_2(s) = \begin{pmatrix} -\sin(\alpha(s)) \\ \cos(\alpha(s)) \end{pmatrix}$$

und mit der ersten Frenet'schen Gleichung folgt

$$e_1'(s) = \alpha'(s) \begin{pmatrix} -\sin(\alpha(s)) \\ \cos(\alpha(s)) \end{pmatrix} = \kappa(s) \begin{pmatrix} -\sin(\alpha(s)) \\ \cos(\alpha(s)) \end{pmatrix}$$

für alle $s \in [0, L]$, also $\kappa = \alpha'$ und

$$\alpha(s) = \alpha(0) + \int_0^s \kappa(\sigma) d\sigma. \quad (1.13)$$

Diese Gleichung drückt aus, dass die Krümmung die Ableitung des Tangentenwinkels nach der Bogenlänge ist. Wir können damit schreiben:

$$c(s) = c(0) + \int_0^s c'(\sigma) d\sigma = c(0) + \int_0^s \begin{pmatrix} \cos(\alpha(0) + \int_0^\sigma \kappa(\tau) d\tau) \\ \sin(\alpha(0) + \int_0^\sigma \kappa(\tau) d\tau) \end{pmatrix} d\sigma \quad (1.14)$$

Eine gegebene, nach Bogenlänge parametrisierte Kurve lässt sich also durch ihren Startpunkt $c(0)$, den Winkel $\alpha(0)$, den sie zu Beginn mit der x -Achse einschließt, und ihre Krümmung ausdrücken.

Jetzt muss nur noch nachgerechnet werden, dass durch den Ausdruck auf der rechten Seite von (1.14) bei gegebener stetiger Funktion $\kappa : [0, L] \rightarrow \mathbb{R}$, gegebenem Anfangspunkt $c(0)$ und gegebenem Anfangswinkel $\alpha(0)$ eine Abbildung $c : [0, L] \rightarrow \mathbb{R}^2$ definiert wird, welche die Gleichungen $|c'| = 1$ und $e'_1 = c'' = \kappa c' = \kappa e_1$ erfüllt. Diese Aufgabe überlassen wir der Leserin.

Wir formen die rechte Seite von (1.14) mit den Additionstheoremen der trigonometrischen Funktionen weiter um und erhalten

$$c(0) + \begin{pmatrix} \cos(\alpha(0)) & \sin(\alpha(0)) \\ \sin(\alpha(0)) & \cos(\alpha(0)) \end{pmatrix} \int_0^s \begin{pmatrix} \cos\left(\int_0^\tau \kappa(\tau) d\tau\right) \\ \sin\left(\int_0^\tau \kappa(\tau) d\tau\right) \end{pmatrix} d\sigma. \quad (1.15)$$

Wir können uns die resultierende Abbildung also folgendermaßen zusammensetzen: Wir berechnen zunächst die Kurve mit vorgegebener Krümmung, die im Ursprung startet und in Richtung positiver x_1 -Achse „losläuft“. Diese ist gegeben durch den Integralterm in (1.15). Die Matrix beschreibt die Drehung um den Ursprung gegen den Uhrzeigersinn um den Winkel $\alpha(0)$. Dadurch ist die Kurve so gedreht worden, dass sie in die durch den Winkel $\alpha(0)$ vorgegebene Richtung losläuft. Schließlich wird die gedrehte Kurve durch die Addition von $c(0)$ verschoben, sodass sie im gewünschten Punkt startet.

Mit Blick auf unser Trassierungsproblem reicht es also, das gesuchte Trassierungselement für den Fall zu berechnen, dass es im Ursprung beginnt und in x_1 -Richtung losläuft. Anschließend kann es passend gedreht und verschoben werden.

1.2.3 Straßenplanung mit Geraden, Kreisbögen und Klothoiden

Wir berechnen jetzt das gesuchte Trassierungselement, einen Straßenverlauf, der bei gleichmäßiger Drehung des Lenkrads mit konstantem Tempo durchfahren werden kann: Die Krümmung κ der gesuchten Kurve soll proportional zur durchfahrenen Strecke s sein. Die Krümmung κ hat die Dimension $1/L$, die Proportionalitätskonstante muss also die Dimension $1/L^2$ haben. Wir führen als Parameter eine Länge A ein und fordern die Gleichung

$$\kappa(s) = \frac{s}{A^2}. \quad (1.16)$$

Weil die Krümmung für positive Streckenlängen positiv gewählt ist, ergibt das eine sich nach links (gegen den Uhrzeigersinn) eindrehende Kurve. Die Wahl $\kappa(s) = -\frac{s}{A^2}$ würde eine sich nach rechts (im Uhrzeigersinn) eindrehende Kurve ergeben. Das gesuchte linksdrehende Trassierungselement, die sogenannte linksdrehende *Klothoide* oder *Spinkurve* mit Formparameter A , hat dann die Form

$$K_A(s) = \int_0^s \begin{pmatrix} \cos\left(\int_0^\sigma \frac{\tau}{A^2} d\tau\right) \\ \sin\left(\int_0^\sigma \frac{\tau}{A^2} d\tau\right) \end{pmatrix} d\sigma = \int_0^s \begin{pmatrix} \cos\left(\frac{\sigma^2}{2A^2}\right) \\ \sin\left(\frac{\sigma^2}{2A^2}\right) \end{pmatrix} d\sigma = \sqrt{2}A \begin{pmatrix} C\left(\frac{s}{\sqrt{2}A}\right) \\ S\left(\frac{s}{\sqrt{2}A}\right) \end{pmatrix} \quad (1.17)$$

mit den dimensionslosen Funktionen

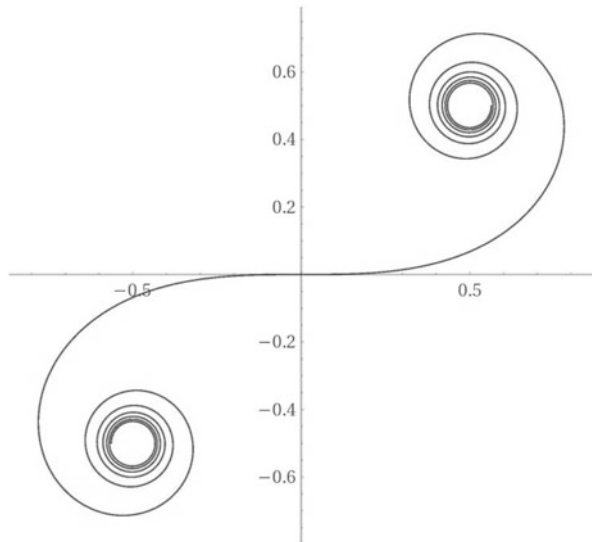
$$C(\sigma) = \int_0^\sigma \cos(\tau^2) d\tau = \sum_{j=0}^{\infty} (-1)^j \frac{\sigma^{4j+1}}{(4j+1)(2j)!}$$

$$S(\sigma) = \int_0^\sigma \sin(\tau^2) d\tau = \sum_{j=0}^{\infty} (-1)^j \frac{\sigma^{4j+3}}{(4j+3)(2j+1)!}.$$

Die Funktionen C und S sind nicht elementar durch trigonometrische oder andere gängige Funktionen ausdrückbar und müssen numerisch durch den Potenzreihenausdruck, der natürlich durch Integration aus den Potenzreihen der trigonometrischen Funktionen entsteht, approximiert werden. (Im Endeffekt sind uns ja auch die Werte der trigonometrischen Funktionen nur über Approximationen durch die Potenzreihe gegeben.) Das Bild der Standardklothoide $\sigma \mapsto (C(\sigma), S(\sigma))^T$ ist in Abb. 1.12 wiedergegeben.

Jetzt sollen mit Geraden, Klothoiden und Kreisen Straßen trassiert werden. Als einfache Standardaufgabe wird hier der Übergang von einer Geraden in einen Kreisbogen mit Radius R vorgestellt. Ein komplexere Trassierungsaufgabe soll in Aufg. 1.11 angegangen werden. Wir wählen das Koordinatensystem so, dass das Geradenstück auf der negativen x_1 -Achse liegt und der Übergang von der Geraden in die Klothoide sich gerade am Ursprung vollzieht. Die folgenden Größen sind relevant:

Abb. 1.12 Die Standardklothoide für Parameterwerte $-5 \leq s \leq 5$



- der Radius R des Kreisbogens, in den überführt werden soll,
- der Klothoidenparameter A ,
- die Bogenlänge L des Klothoidenstücks,
- der Winkel α , um den der Tangenteneinheitsvektor während des Übergangs von der Geraden zum Kreisbogen gedreht werden soll.

Diese Größen stehen in den folgenden Zusammenhängen:

$$A^2 = RL \quad \text{und} \quad \alpha = \frac{L^2}{2A^2}. \quad (1.18)$$

Der Radius R ist nämlich gerade das Inverse der Krümmung, welche die Klothoide am Ende des Klothoidenstücks hat: $\kappa(L) = \frac{L}{A^2} = 1/R$, siehe (1.16). Damit haben wir ebenfalls einen orientierten Radius: Bei einem positiven Radius wird der Kreisbogen gegen den Uhrzeigersinn durchlaufen, bei einem negativen Radius im Uhrzeigersinn. Andererseits gilt ebenfalls $\alpha(L) = \alpha(0) + \int_0^L \kappa(\sigma) d\sigma = 0 + \frac{L^2}{2A^2}$, siehe (1.13). Zwei dieser vier Parameter können also frei gewählt werden, die anderen beiden sind dann bereits fixiert. Nach den *Richtlinien für die Anlage von Landstraßen* ist $R/3 \leq A \leq R$ einzuhalten: die erste Ungleichung, damit der Übergang in die Kurve rechtzeitig vom Fahrer wahrgenommen wird, und die zweite Ungleichung, damit das Lenkrad nicht zu schnell gedreht werden muss, siehe [76]. Zu vorgegebenem R und A soll jetzt der Mittelpunkt des Kreises berechnet werden, in den durch das Klothoidenstück überführt werden soll, siehe Abb. 1.13. Für die Bogenlänge L und den Winkel α gilt somit

$$L = \frac{A^2}{R} \quad \text{und} \quad \alpha = \frac{A^2}{2R^2}.$$

Mit den Bezeichnungen aus der Figur in Abb. 1.13 sind die Koordinaten (x_P, y_P) des Punktes P bekannt, nämlich

$$\begin{aligned} x_P &= \sqrt{2}AC \left(\frac{A}{\sqrt{2}R} \right) \\ y_P &= \sqrt{2}AS \left(\frac{A}{\sqrt{2}R} \right). \end{aligned}$$

Die Tangente t an die Klothoide im Punkt P ist gleichzeitig die Tangente an den Kreis im Punkt P . Aufgrund des rechten Winkels zwischen Tangente und Radius sind die Dreiecke ABP und MCP ähnlich und die Winkelgröße α erhalten wir sowohl am Winkel $\angle BAP$ als auch am Winkel $\angle CMP$. Mit den trigonometrischen Beziehungen im Dreieck $CM P$ ergibt sich

$$\begin{aligned} \cos(\alpha) &= \frac{y_M - y_P}{R} \quad \text{also} \quad y_M = R \cos \left(\frac{A^2}{2R^2} \right) + y_P \\ \sin(\alpha) &= \frac{x_P - x_M}{R} \quad \text{also} \quad x_M = x_P - R \sin \left(\frac{A^2}{2R^2} \right). \end{aligned} \quad (1.19)$$

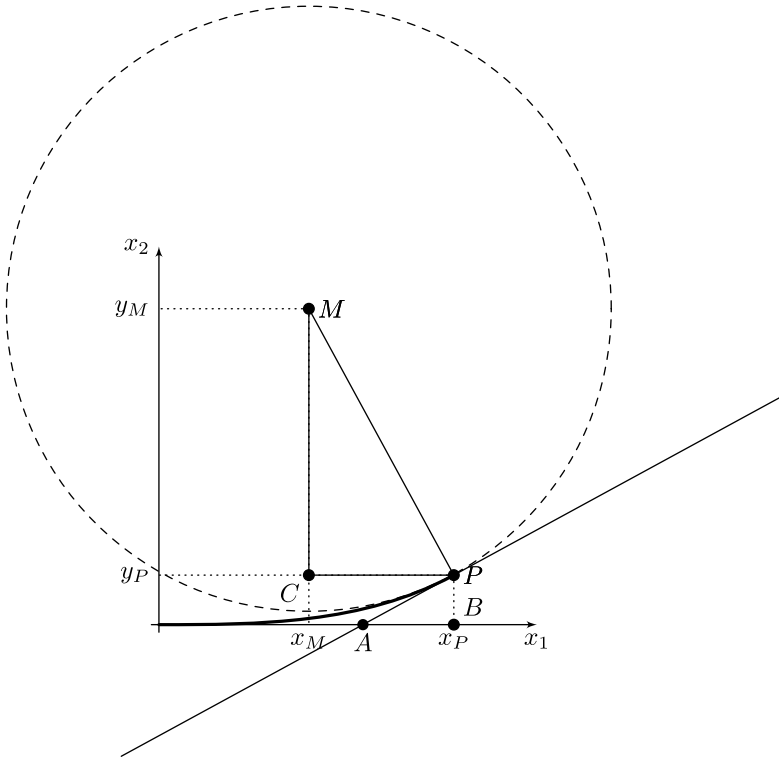


Abb. 1.13 Der Übergang von einem Geradenstück in einen Kreisbogen mit Radius R mit einem Klothoidenstück mit Parameter A . In dieser Abbildung gilt $A = R$

Damit ist der einfachste Übergang von einer Geraden auf einen Kreisbogen mit Radius R konstruiert. Für den Übergang von einem Kreis mit Radius R_1 zu einem Kreis mit Radius R_2 muss das Klothoidenstück zwischen den Krümmungen $\frac{1}{R_1}$ und $\frac{1}{R_2}$ eingepasst werden. Bei der Verwendung einer Klothoide mit Formparameter A ist dies das Stück von K_A mit s zwischen $\frac{A^2}{R_1}$ und $\frac{A^2}{R_2}$. Durch geeignet gewählte Drehmatrizen und Verschiebungsvektoren kann dieses Stück dann an die richtige Stelle gedreht und verschoben werden.

1.2.4 Reflexionen zum mathematischen Modellieren II: normative Modelle

Es besteht ein gravierender Unterschied zwischen der Modellierung eines eingespannten Kabels durch eine hyperbolische Funktion und der Modellierung eines Straßenverlaufs durch Geraden, Kreise und Klothoiden: Während sich ein Kabel von Natur aus auf irgendeine Art und Weise hingängt, wird der Straßenverlauf durch den Menschen geplant. Es

wird also nicht ein mathematisches Modell ausgesucht, um einen Teilbereich der Realität zu beschreiben, sondern ein Teilbereich der Realität wird erst nach der Maßgabe eines mathematischen Modells erschaffen. Man spricht in diesem Zusammenhang von *normativen Modellen*. Ähnlich gelagert sind z. B. Zins- und Steuermodelle, aber auch Maschinen oder Computerprogramme sind Teile der Realität, die zumindest in bestimmten Teilen nach mathematischen Modellen erschaffen worden sind. Die zunehmende Bedeutung der Mathematik für die moderne Welt hängt also nicht nur damit zusammen, dass immer mehr Bereiche der Natur und der Gesellschaft durch mathematische Modelle beschrieben werden, sondern vor allem auch damit, dass immer mehr gesellschaftlich relevante Bereiche überhaupt erst nach mathematischen Modellen erschaffen werden.

1.3 Der Fahnenmast – vom Nutzen der Entdimensionalisierung und der Wahl geeigneter Skalen

Vielleicht ist der Leser noch nicht vom Nutzen des Entdimensionalisierens überzeugt und denkt sich, dass man die maßstabsgetreue Verkleinerung auch ohne solch großen Aufwand bewerkstelligen könne. Wir wollen nun ein Beispiel ausführen, in dem deutlicher zu erkennen ist, dass die Dimensionsanalyse tatsächlich so etwas wie eine *magic wand* des Modellierens ist: Ohne tatsächlich Lösungen des Modells ausrechnen zu müssen, lässt sich bereits viel über die Lösungen aussagen. Allerdings müssen wir, um den Nutzen der Entdimensionalisierung eindrucksvoll erleben zu können, auch anspruchsvollere Mathematik betrachten. Dafür soll das Verhalten eines Fahnenmastes bei einem Erdbeben modelliert werden. Die Darstellung orientiert sich an [31]. Wir betrachten einen Fahnenmast zunächst einmal als stabförmig mit der Länge l , ein plausibler Beispielwert ist z. B. $l = 10.00$ m. Wir nutzen ein zweidimensionales Koordinatensystem (u, x) , wobei u die horizontale und x die vertikale Koordinate angibt. Ohne Erdbeben nimmt der Fahnenmast die Punkte der Form $u = 0$ und $0 < x < l$ ein. Durch das Erdbeben bewegt sich nun der Erdboden in einer bestimmten Art und Weise und der Fahnenmast reagiert darauf, indem er sich aus der Ruhelage auslenkt. Wir wollen mit $u(t, x)$ die Auslenkung zur Zeit t in der Höhe x bezeichnen. Dafür müssen wir wieder einige vereinfachende Annahmen treffen.

1.3.1 Annahmen und Vereinfachungen

Wir ordnen die Annahmen und Vereinfachungen in unterschiedliche Gruppen, zunächst betrachten wir das Erdbeben: Für das Erdbeben nehmen wir an, dass sich der Erdboden näherungsweise sinusförmig hin- und herbewegt und dabei den Fußpunkt der Fahnenstange mitnimmt. Damit ergibt sich für die Auslenkung der Fahnenstange am Boden die Bewegung

$$u(0, t) = a \sin(\omega t). \quad (1.20)$$

Der Parameter a hat dabei die Dimension einer Länge und gibt die Amplitude des Erdbebens an, der Parameter ω ist das Inverse einer Zeit, hat also die Dimension $1/T$, und gibt die Frequenz des Bebens an.

Für die Fahnenstange nehmen wir an, dass es sich um einen Hohlzylinder aus Stahl handelt mit der Länge l und dem äußeren Radius r .

1.3.2 Das mathematische Modell

Nun benötigen wir ein mathematisches Modell, das diese Situation gut beschreibt. Mit der Modellierung dieser komplexen Situation sind wir überfordert, aber Physiker und Bauingenieure behaupten, dass die folgende partielle Differentialgleichung eine gute Beschreibung der Situation liefert (eine Herleitung der Gleichungen findet sich in [81]):

$$\rho \frac{\partial^2 u(t, x)}{\partial t^2} = -Ek^2 \frac{\partial^4 u(t, x)}{\partial x^4} \quad \text{für } 0 < x < l \quad (1.21)$$

Wir haben also drei weitere Parameter ρ , E und k , deren Bedeutung wir gleich besprechen werden. Der Term $\frac{\partial^2 u(t, x)}{\partial t^2}$, also die zweifache partielle Ableitung nach der Zeit t , gibt die Beschleunigung des Fahnenstangenteils in der Höhe x zur Zeit t an und der Term $-Ek^2 \frac{\partial^4 u(t, x)}{\partial x^4}$, also die vierfache partielle Ableitung nach dem Ort x , die elastische Kraft, die auf die Fahnenstange an dieser Stelle aufgrund der Verformung der Fahnenstange wirkt. Die physikalisch vorgebildete Leserin kann hier vielleicht die Newton-Gleichung $F = m \cdot a$ wiedererkennen. Gl. (1.21) besitzt sehr viele Lösungen, z. B. wäre die Nullfunktion $u(t, x) = 0$ eine. Um nun eine eindeutige Lösung zu erhalten, müssen weitere Bedingungen hinzugefügt werden. Weil das Erdbeben als periodisch mit der Periode $T = \frac{2\pi}{\omega}$ angesetzt wurde, ist es plausibel, dasselbe auch von der Auslenkung zu erwarten, wir suchen also eine T -periodische Funktion in der Zeitvariablen:

$$u(t + T, x) = u(t, x) \quad \text{für alle Zeiten } t \text{ und alle Höhen } 0 < x < l. \quad (1.22)$$

Ferner sind die folgenden vier sogenannten Randbedingungen plausibel: Für alle Zeiten t gilt:

$$\begin{aligned} u(0, t) &= a \sin(\omega t), & \frac{\partial u}{\partial x}(0, t) &= 0, \\ \frac{\partial^2 u}{\partial x^2}(l, t) &= 0, & \frac{\partial^3 u}{\partial x^3}(l, t) &= 0. \end{aligned} \quad (1.23)$$

Dabei sagt die erste Gleichung in (1.23) aus, dass der Fuß der Fahnenstange genau die Bewegung des Erdbebens mitmacht, die Stange also nicht abbricht. Die zweite Gleichung sagt aus, dass die Stange am Fuß senkrecht stehen bleibt, also nicht abknickt. Die dritte und vierte Gleichung beschreiben das Verhalten der Fahnenstange am oberen Ende: Das obere Ende muss nicht mehr senkrecht nach oben zeigen, aber die Fahnenstange ist am Ende

ungekrümmt. Es ist nun die Aufgabe der auf diesen Bereich spezialisierten Mathematiker, zu beweisen, dass ein solches sogenanntes periodisches Randwertproblem zu einer partiellen Differentialgleichung, bestehend aus den Gl. (1.21) und (1.23), genau eine T -periodische Lösung hat.

Wir wollen jetzt auf die verwendeten Parameter eingehen. Der Parameter ρ gibt die Dichte des Materials an und hat in diesem Kontext die Dimension $\frac{\text{Masse}}{\text{Volumen}} = \frac{M}{L^3}$. Für das Material Stahl haben wir $\rho_{\text{Stahl}} = 7.8 \times 10^3 \text{ kg m}^{-3}$. Auch der sogenannte Young'sche Elastizitätsmodul E ist eine Eigenschaft des verarbeiteten Materials und hat die Dimension einer Kraft, $\frac{\text{Masse}}{\text{Länge} \cdot \text{Zeit}^2} = \frac{M}{LT^2}$. Für Stahl gilt $E_{\text{Stahl}} = 2.0 \times 10^7 \text{ kg m}^{-1} \text{ s}^{-2}$. Der Parameter k bezieht sich auf die Geometrie des Querschnitts der Fahnenstange und wird Trägheitsradius genannt. Die Bauingenieure behaupten, dass für einen dünnwandigen Hohlzylinder mit Radius d die Beziehung $k = \alpha r$ mit $\alpha \approx \frac{1}{\sqrt{2}}$ gilt, siehe [11] (an dieser Stelle müssen wir tatsächlich das Approximationszeichen verwenden). Die Approximation ist umso besser, je kleiner das Verhältnis von innerem Radius zu äußerem Radius des Hohlzylinders ist. Der Trägheitsradius hat damit tatsächlich die Dimension einer Länge.

Auch wenn wir die einfließenden physikalischen Überlegungen nicht im Detail oder vielleicht auch gar nicht verstehen, können wir zumindest überprüfen, ob die Modellgleichung (1.21) von den Dimensionen her konsistent ist, und überlegen, welche Dimensionen die partiellen Ableitungen haben: Die Auslenkung $u(t, x)$ hat die Dimension einer Länge. Die (erste) partielle Ableitung nach t entsteht als Limes eines Differenzenquotienten $\frac{u(t+\Delta t, x) - u(t, x)}{\Delta t}$ für $\Delta t \rightarrow 0$, hat also die Dimension $\frac{L}{T}$, die zweite partielle Ableitung nach der Zeit dann entsprechend die Dimension $\frac{L}{T^2}$. Damit haben wir auf der linken Seite von (1.21) insgesamt die Dimension

$$\left[\rho \frac{\partial^2 u}{\partial t^2} \right] = \frac{M}{L^3} \cdot \frac{L}{T^2} = \frac{M}{L^2 T^2}.$$

Für die rechte Seite erhalten wir mit entsprechenden Überlegungen

$$\left[Ek^2 \frac{\partial^4 u}{\partial x^4} \right] = \frac{M}{LT^2} \cdot L^2 \cdot \frac{L}{L^4} = \frac{M}{L^2 T^2},$$

die Gleichung ist also in Bezug auf die Dimensionen konsistent.

1.3.3 Entdimensionalisierung

Wir wollen das Problem nun entdimensionalisieren und notieren zunächst einmal die vorkommenden Variablen und Parameter:

Variablen: Zeit t , Ort x und Auslenkung u .

Parameter: die Erdbebenamplitude a , die Erdbebenfrequenz ω , die Materialparameter ρ und E sowie die geometrischen Parameter d und l .

Für die Variablen führen wir dimensionslose Variablen ein, indem wir die noch unbestimmten Skalen $\bar{t}$, $\bar{x}$ und $\bar{u}$ benutzen:

$$\begin{aligned} \tau &= \frac{t}{\bar{t}}, & \xi &= \frac{x}{\bar{x}}, & v &= \frac{u}{\bar{u}} \text{ bzw.} \\ t &= \bar{t}\tau, & x &= \bar{x}\xi, & u &= \bar{u}v, \end{aligned}$$

zusammen mit den Beziehungen

$$v(\xi, \tau) = \frac{u}{\bar{u}}(\bar{x}\xi, \bar{t}\tau) \text{ bzw.} \quad u(t, x) = \bar{u}v\left(\frac{x}{\bar{x}}, \frac{t}{\bar{t}}\right).$$

Wir berechnen nun, welche Differentialgleichung die dimensionslose Funktion v löst, und beginnen mit der doppelten Zeitableitung:

$$\begin{aligned} \frac{\partial^2 v}{\partial \tau^2}(\xi, \tau) &= \frac{\partial^2}{\partial \tau^2} \left(\frac{1}{\bar{u}} u(\bar{x}\xi, \bar{t}\tau) \right) = \frac{\bar{t}^2}{\bar{u}} \frac{\partial^2 u}{\partial t^2}(\bar{x}\xi, \bar{t}\tau) = \frac{\bar{t}^2}{\bar{u}} \frac{\partial^2 u}{\partial t^2}(x, t) \\ &= \frac{\bar{t}^2}{\bar{u}} \left(-\frac{Ek^2}{\rho} \frac{\partial^4 u}{\partial x^4}(x, t) \right) = -\frac{\bar{t}^2 Ek^2}{\bar{u}\rho} \frac{\partial^4}{\partial x^4} \left(\bar{u}v\left(\frac{x}{\bar{x}}, \frac{t}{\bar{t}}\right) \right) \\ &= -\frac{\bar{t}^2 Ek^2}{\bar{u}\rho} \cdot \frac{\bar{u}}{\bar{x}^4} \frac{\partial^4 v}{\partial \xi^4}(v, \tau) \end{aligned} \quad (1.24)$$

Dabei haben wir für das zweite und sechste Gleichheitszeichen zweimal bzw. viermal die Kettenregel angewendet (bei jedem Ableiten erhält man als innere Ableitung $\bar{t}$ bzw. $1/\bar{x}$) und für das vierte Gleichheitszeichen mit Hilfe der Differentialgleichung (1.21) die zweite Zeitableitung ersetzt.

Es gibt eine Kontrollmöglichkeit, ob wir richtig gerechnet haben: Bei korrekter Rechnung muss die entstandene Parameter- und Skalenkombination $\frac{\bar{t}^2 Ek^2 \bar{u}}{\bar{u}\rho \bar{x}^4}$ dimensionslos sein, wovon wir uns leicht überzeugen können. Wir entdimensionalisieren ebenfalls die Randbedingungen (1.23) und erhalten mit Hilfe entsprechender Rechnungen

$$\begin{aligned} v(0, \tau) &= \frac{a}{\bar{u}} \sin(\omega \bar{t}\tau), & \frac{\partial v}{\partial \xi}(0, \tau) &= 0, \\ \frac{\partial^2 v}{\partial \xi^2} \left(\frac{l}{\bar{x}}, \tau \right) &= 0, & \frac{\partial^3 v}{\partial \xi^3} \left(\frac{l}{\bar{x}}, \tau \right) &= 0. \end{aligned} \quad (1.25)$$

Wie sollen die Skalen $\bar{x}$, $\bar{t}$ und $\bar{u}$ jetzt gewählt werden? Die aus der Situation vorgegebenen Skalen sind $\bar{x} = l$, die Länge der Fahnenstange, $\bar{u} = a$ und $\bar{t} = \frac{1}{\omega} = T$, welche das Erdbeben beschreiben. Wählen wir diese Skalen, so bleibt in der Gl. (1.24) der dimensionslose Parameter

$$J = \frac{Ek^2}{\rho\omega l^4} \quad (1.26)$$

und wir erhalten das Problem

$$\begin{aligned} \frac{\partial^2 v}{\partial \tau^2} &= -J \frac{\partial^4 v}{\partial \xi^4} & v \text{ periodisch in } \tau \text{ mit Periode } 2\pi \\ v(0, \tau) &= \sin(\tau), & \frac{\partial v}{\partial \xi}(0, \tau) &= 0, \\ \frac{\partial^2 v}{\partial \xi^2}(1, \tau) &= 0, & \frac{\partial^3 v}{\partial \xi^3}(1, \tau) &= 0. \end{aligned} \quad (1.27)$$

1.3.4 Vom Nutzen der Entdimensionalisierung

Welche Vorteile bringt uns das jetzt?

- Das Problem hängt nur noch von einem Parameter J statt von den Parametern a , ω , E , ρ , k und l ab. Das macht die Arbeit für den Analytiker einfacher, das Lösungsverhalten hängt schließlich nur von J ab. Es macht auch die Arbeit für den Numeriker einfacher, denn es reicht, den einen Parameter zu variieren anstatt der sechs.
- Man kann das Verhalten eines solchen Systems im Kleinen simulieren und so überprüfen, ob die Lösungen des mathematischen Modells sich so verhalten wie der Gegenstand, den es modellieren soll.
- Man kann geeignete verkleinerte Realmodelle bauen, die das Verhalten des großen Systems simulieren.

Den letzten Punkt wollen wir genauer erläutern. Dazu legen wir zunächst plausible Werte für die Fahnenstange fest: Wir gehen von einer Stange mit der Länge $l = 10.0$ m und einem Durchmesser $r = 0.15$ m aus. Für den Trägheitsradius erhalten wir dann $k = 5.3 \times 10^{-2}$ m. Für das Erdbeben setzen wir eine Periode von $T = 2.0$ s an, also $\omega \approx \frac{\pi}{s}$. Für einen Stahlmast gilt $\rho_{\text{Stahl}} = 7.8 \times 10^3 \text{ kg m}^{-3}$ und $E = 2.0 \times 10^7 \text{ kg m}^{-1} \text{ s}^{-2}$. Mit diesen Angaben berechnet sich der dimensionslose Parameter J gemäß (1.26) zu

$$J = 0.78 \times 10^{-2}.$$

Für die Simulation der Fahnenstange können wir ein anderes System wählen, das im Bereich desselben J liegt. Möglich wäre ein Holzstab, sagen wir mit der Länge $l_H = 2$ m. Für die Materialkonstanten gilt $\rho_{\text{Holz}} = \frac{\rho_{\text{Stahl}}}{13}$ und $E_{\text{Holz}} = \frac{E_{\text{Stahl}}}{20}$. Wir könnten auch die Erdbebenfrequenz ändern, bleiben da aber bei $T = 2$ s. Richtig einstellen müssen wir nun den Durchmesser des Holzstabs, sodass wir den gleichen dimensionslosen Parameter erhalten. Für den Trägheitsradius k_{Holz} benötigen wir also

$$k_{\text{Holz}}^2 = \frac{13}{20 \cdot 5^4} k_{\text{Fahnenstange}}^2.$$

Wir müssen dabei beachten, dass wir es bei dem Holzstab mit einem Voll- und keinem Hohlzylinder zu tun haben. Mit der nötigen Anpassung erhalten wir schlussendlich $r_{\text{Holz}} = 1 \text{ cm}$.

Gemäß dem mathematischen Modell verhält sich also ein 2 m langer Holzstab mit einem Durchmesser von 1 cm bei einem Beben der Periode 2 s genauso wie ein 10 m hoher Hohlzylinder-Stahlmast mit Durchmesser 15 cm.

Aufgaben

1.1 Die Bogenlänge eines Funktionsgraphen

Für eine stetig differenzierbare Funktion $f : [a, b] \rightarrow \mathbb{R}$ wird die Bogenlänge des Graphen definiert als

$$\int_a^b \sqrt{1 + (f'(x))^2} dx.$$

Erklären Sie, warum diese Definition sinnvoll ist.

Hinweis: Teilen Sie das Intervall zunächst einmal in n gleich große Teilstücke und approximieren Sie durch Sekantenstücke, siehe Abb. 1.14, ...

1.2 Die hyperbolischen Funktionen

Die Funktionen Sinushyperbolicus und Kosinushyperbolicus sind definiert als

$$\sinh : \mathbb{R} \longrightarrow \mathbb{R} \quad \text{und} \quad \cosh : \mathbb{R} \longrightarrow \mathbb{R}$$

$$x \longmapsto \frac{e^x - e^{-x}}{2} \quad \quad \quad x \longmapsto \frac{e^x + e^{-x}}{2}$$

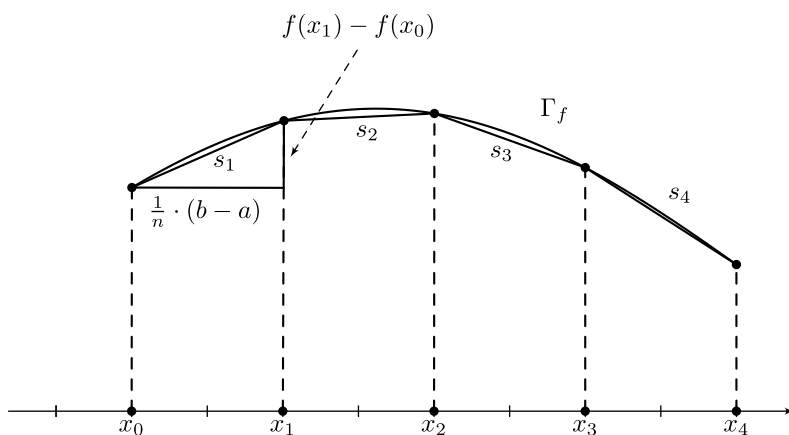


Abb. 1.14 Approximation der Bogenlänge durch die Länge von Sekantenstücken

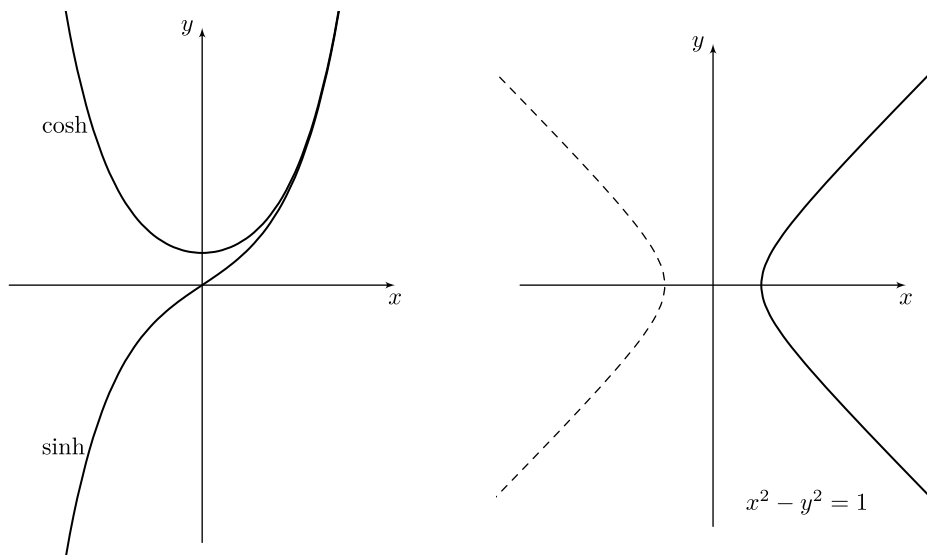


Abb. 1.15 Die Graphen der hyperbolischen Funktionen und die Hyperbel $x^2 - y^2 = 1$

Zeigen Sie:

- Die Funktionen sind ungerade bzw. gerade.
- $\sinh' = \cosh$ und $\cosh' = \sinh$
- $\cosh^2 - \sinh^2 = 1$
- $\sinh(x) = \sum_{j=0}^{\infty} \frac{x^{2j+1}}{(2j+1)!}$ und $\cosh(x) = \sum_{j=0}^{\infty} \frac{x^{2j}}{(2j)!}$
- Der Sinushyperbolicus hat eine Umkehrfunktion, genannt Areasinushyperbolicus, $\operatorname{arsinh} : \mathbb{R} \rightarrow \mathbb{R}$. Für diese gilt $\operatorname{arsinh}'(x) = \frac{1}{\sqrt{1+x^2}}$ für alle $x \in \mathbb{R}$.
- Die Einschränkung des Kosinushyperbolicus auf das Intervall $[0, \infty)$ hat eine Umkehrfunktion, genannt Areakosinushyperbolicus, $\operatorname{arcosh} : [1, \infty] \rightarrow [0, \infty)$. Für diese gilt $\operatorname{arcosh}'(x) = \frac{1}{\sqrt{x^2-1}}$ für alle $x > 1$.
- Die Kurve $s \mapsto \begin{pmatrix} \cosh(s) \\ \sinh(s) \end{pmatrix}$ durchläuft den rechten Zweig der Hyperbel $x^2 - y^2 = 1$. Daher rührt die Bezeichnung *hyperbolische Funktionen*, siehe Abb. 1.15.

1.3 Das numerische Verfahren zur Bestimmung von η

In dieser Aufgabe soll die numerische Lösung der Gleichung

$$\cosh(\eta) - 1 - 2\varepsilon\eta = 0 \quad (1.28)$$

genauer untersucht werden.

- a) Beweisen Sie, dass die Gleichung (1.28) genau eine positive Lösung hat.
- b) Bestimmen sie für $a = 100.00$ m und $d = 10.00$ m bzw. $d = 50.00$ m die positive Nullstelle η^* der Gleichung bis auf fünf Dezimalstellen genau. Verwenden Sie zum einen das vorgestellte Verfahren

$$\eta_{n+1} = \operatorname{arcosh}(1 + 2\varepsilon\eta_n) \quad (1.29)$$

und zum anderen das Newton-Verfahren zur Bestimmung von Nullstellen. Implementieren Sie beide Verfahren mit einem Tabellenkalkulationsprogramm und vergleichen Sie die Schnelligkeit der Annäherung.

- c) Beweisen Sie mit den Mitteln der Analysis die Konvergenz des Verfahrens (1.29).

1.4 Numerischer Vergleich der Kabelmodelle für verschiedene Werte des Verhältnisses von Durchhang zu Abstand

Wie groß wird der Unterschied für die Vorhersagen der Kabellänge nach den Modellen 1, 2 und 3 in Abschn. 1.1 für unterschiedliche Verhältnisse $\varepsilon = \frac{d}{a}$? Erstellen Sie eine GeoGebra-Datei, die in einer entdimensionalisierten Situation die Graphen nach Modell 1, 2 und 3 in einem gemeinsamen Koordinatensystem darstellt (ε mit Schieberegler einbinden, die nötige Iteration für η lässt sich mit der TK-Funktion von GeoGebra durchführen) und jeweils die entdimensionalisierte Länge des Kabels berechnet.

1.5 Planung einer Seilbahn

Eine Seilbahn soll auf einer horizontalen Entfernung a einen vertikalen Höhenunterschied d überwinden. Das verwendete Kabel soll eine Länge l haben.

- a) Erstellen Sie eine Steckbriefaufgabe, die die Höhenverlaufsfunktion des Seils bestimmt. Verwenden Sie als Modellfunktion geeignete Verschiebungen des Kosinushyperbolicus.
- b) Versuchen Sie das durch die Steckbriefaufgabe entstandene Gleichungssystem zu lösen. Falls Ihnen keine algebraisch-analytische Lösung gelingt, versuchen Sie für die Werte $a = 2000$ m, $d = 2000$ m und $l = 3000$ m auf irgendeine Art und Weise eine numerische Approximationslösung zu finden (das darf ruhig hemdsärmelig geschehen).

1.6 Reguläre orientierte Kurven als Äquivalenzklasse von Parametrisierungen

Rechnen Sie nach, dass mit der Definition 1.2 tatsächlich eine Äquivalenzrelation definiert wird.

1.7 Die Bogenlänge und die Parametrisierung nach Bogenlänge

Es sei $f : \mathbb{R}^2 \longrightarrow \mathbb{R}$ eine glatte Funktion und $\mathcal{C}$ eine regulär parametrisierbare Kurve.

- a) Zeigen Sie, dass das Kurvenintegral

$$\int_C f \, dl := \int_a^b f(\gamma(t)) |\dot{\gamma}(t)| \, dt \quad (1.30)$$

wohldefiniert ist, d. h. unabhängig von der Wahl der Parametrisierung γ ist.

- b) Klären Sie die Beziehung zwischen den Bogenlängendefinitionen in Gleichung (1.3) und Gleichung (1.7).
 c) Zeigen Sie, dass es zu C eine *Parametrisierung nach Bogenlänge* gibt, d. h. eine Parametrisierung

$$c : [0, L] \longrightarrow \mathbb{R}^2 \quad \text{mit} \quad |c'(s)| = 1 \quad \text{für alle} \quad s \in [0, L].$$

Ist diese Parametrisierung eindeutig?

1.8 Kurven mit konstanter Krümmung

Zeigen Sie, dass genau die Strecken und die Kreisbögen Kurven mit konstanter Krümmung sind. Betrachten Sie dazu für ein $\kappa(s_0) \neq 0$ den Punkt $c(s) + \frac{1}{\kappa(s_0)} e_2(s)$. Unter welchen Umständen ist dieser Punkt konstant, wenn der Bogenlängenparameter s variiert wird?

1.9 Die Krümmung in beliebigen Parametrisierungen

Wir betrachten eine reguläre orientierte Kurve C mit Parametrisierung $\gamma : [t_0, t_1] \longrightarrow \mathbb{R}^2$ und der Parametrisierung nach Bogenmaß $c : [0, L] \longrightarrow \mathbb{R}^2$.

- a) Zeigen Sie, dass für die Krümmung als Funktion vom Parameter t die folgende Beziehung gilt:

$$\kappa(t) = \frac{\ddot{\gamma}_2(t)\dot{\gamma}_1(t) - \ddot{\gamma}_1(t)\dot{\gamma}_2(t)}{|\dot{\gamma}(t)|^3} \quad (1.31)$$

Anleitung: Nutzen Sie die bekannte Formel für die Krümmung und betrachten Sie $\kappa(s) = \kappa(s(t(s)))$. Bezüglich der Parametrisierung γ ist die Krümmung dann durch $\kappa(t) = \kappa(s(t))$ gegeben. Bezüglich des Gebrauchs von Variablen- und Funktionsbezeichnungen siehe auch den ersten Kasten in Abschn. 1.2.2.

- b) Berechnen Sie die Formel für die Krümmung des Graphen einer Funktion f im Punkt $(t, f(t))$.

- c) Wir definieren den Berührungskreis an den Punkt $c(s_0)$ einer Kurve als den Kreis mit Radius $1/|\kappa(s_0)|$ um den Punkt $M(s_0) = c(s_0) + \frac{1}{\kappa(s_0)} e_2(s_0)$ (c sei die Parametrisierung nach Bogenlänge). Zeigen Sie, dass dieser Kreis die Kurve C *in zweiter Ordnung berührt*.

Anleitung: Nehmen Sie zunächst an, dass $e_1(s_0)$ in Richtung der positiven x -Achse zeigt, und überlegen Sie sich danach, warum Sie das dürfen. Überlegen Sie sich, was wohl *in zweiter Ordnung berühren* bedeuten könnte. (Hinweis: Die Tangente berührt in erster Ordnung.)

1.10 Die logarithmische Spirale

In Polarkoordinaten ist die *logarithmische Spirale* gegeben durch $r(t) = \exp(t)$ und $\varphi(t) = at$ mit einer Konstanten a . Zeigen Sie:

- Die Länge der logarithmischen Spirale über dem Intervall $(-\infty, t]$ ist proportional zum Radius $r(t)$.
- Der Ortsvektor der logarithmischen Spirale hat einen konstanten Winkel mit der Tangente.

1.11. Planung einer Schnellstraße

Eine neue Schnellstraße soll geplant werden, die im Groben von der Geraden, die im gewählten Koordinatensystem durch die Gleichung $y = 0$ bestimmt wird, in einer Linkskurve auf die Gerade, die im Koordinatensystem durch die Gleichung $y = \tan(\pi/3)(x - 300)$ gegeben ist, wechselt, siehe Abb. 1.16. Das Koordinatensystem ist in Meter skaliert. Die Straße soll mit einer Höchstgeschwindigkeit von 100 km h^{-1} befahren werden können. Der minimale Krümmungsradius muss dafür größer als 450 m sein. Planen Sie einen fachmännischen Straßenverlauf und plotten Sie ihn mit GeoGebra. Können Sie den Verlauf so planen, dass der im Bild zu sehende Hof nicht abgerissen werden muss?

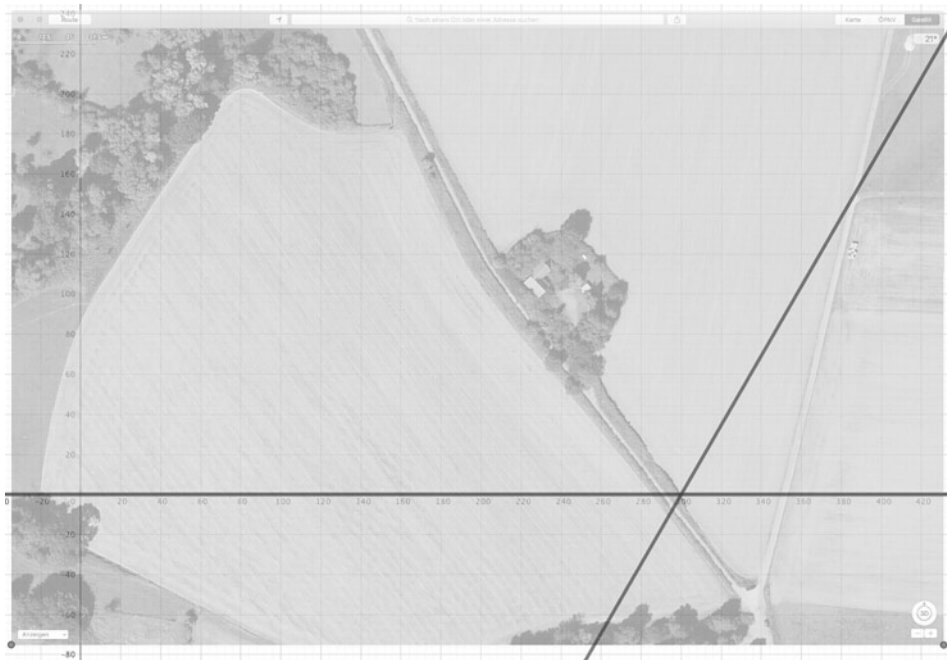


Abb. 1.16 Vorgabe für einen zu planenden Straßenverlauf, Abbildung mit GeoGebra

1.12 Die Planung von Schiffen

In dieser Aufgabe geht es, wie im Abschn. 1.3, mehr darum, den Vorgang der Entdimensionalisierung und das Finden geeigneter Skalen zu üben, als darum, die hinter dieser Aufgabe stehende Physik vollständig zu verstehen.

Es soll modelliert werden, wie eine Flüssigkeit einen Gegenstand umspült, z. B. Wasser einen Schiffsrumpf. Dazu stellen wir uns den Gegenstand als ruhend vor und die Flüssigkeit in Bewegung, allerdings so, dass der Raumbereich, den die Flüssigkeit einnimmt, sich nicht ändert. Man kann sich z. B. einen begradigten Fluss vorstellen. Die Bewegung der Flüssigkeit wird nun beschrieben, indem zur Zeit t jedem Punkt x , der von der Flüssigkeit eingenommen wird, ein Geschwindigkeitsvektor $v = v(t, x)$ zugeordnet wird, der die Geschwindigkeit der Flüssigkeit im Punkt x zur Zeit t beschreibt. Nach den Modellen der Fluidodynamik soll das Geschwindigkeitsvektorfeld v dann den folgenden Gleichungen genügen:

- Für große x , das heißt für Punkte, die weit entfernt von dem störenden Körper sind, soll die Geschwindigkeit konstant sein. Mathematisch ausgedrückt heißt das

$$\lim_{|x| \rightarrow \infty} v(t, x) = V \in \mathbb{R}^3$$

für alle Zeiten t , wobei V ein vorgegebener Geschwindigkeitsvektor ist. Wenn wir uns den begradigten Fluss vorstellen, würde V richtungsmäßig einfach den Fluss hinab zeigen und hätte einen gewissen Betrag.

- Bei Vernachlässigung aller weiteren von außen einwirkenden Kräfte gelten die Navier-Stokes-Gleichungen:

$$\begin{aligned} \varrho (\partial_t v + (v \cdot \nabla) v) &= -\nabla p + \mu \Delta v \\ \nabla \cdot v &= 0 \end{aligned}$$

Dabei sind

$$\begin{aligned} \nabla \cdot v &= \sum_{i=1}^3 \frac{\partial v_i}{\partial x_i} \in \mathbb{R} \quad \text{die Divergenz des Vektorfelds } v, \\ \Delta v &= \sum_{i=1}^3 \frac{\partial^2 v}{\partial x_i^2} \in \mathbb{R}^3 \quad \text{der Laplace-Operator angewendet auf } v, \\ (v \cdot \nabla) v &= \left(\sum_{i=1}^3 v_i \frac{\partial v_j}{\partial x_i} \right)_{j=1,2,3} \in \mathbb{R}^3 \quad \text{der Konvektionsterm,} \\ \nabla p &= \left(\frac{\partial p}{\partial x_i} \right)_{i=1,2,3} \in \mathbb{R}^3 \quad \text{der Gradient von } p. \end{aligned}$$

Ferner ist ϱ die konstante Massendichte der Flüssigkeit, $[\varrho] = \frac{M}{L^3}$. Die sogenannte *dynamische Viskosität* der Flüssigkeit wird mit μ bezeichnet. Sie ist ein Maß für die

Zähflüssigkeit der Flüssigkeit. Zum Beispiel gilt $\frac{\mu_{\text{Honig}}}{\mu_{\text{Wasser}}} \gg 1$. Der Druck, unter dem die Flüssigkeit zur Zeit t an der Stelle x steht, wird mit $p(t, x)$ bezeichnet. Den Durchmesser des Gegenstandes, der vom Wasser umströmt wird, bezeichnen wir mit l .

- Bestimmen Sie die Dimensionen aller auftretenden Parameter und Variablen.
- Entdimensionalisieren Sie das Problem. Wählen Sie dabei dem Problem angepasste Skalen und bestimmen Sie geeignete dimensionslose Parameter.
- Beschreiben Sie, wie Sie vorgehen können, wenn Sie das Verhalten eines 300 m langen Ozeanriesen in einer 50 m langen Halle untersuchen wollen.

1.13 Die Berechnung der Planck-Größen aus den Naturkonstanten

In der modernen Physik betrachtet man die folgenden Größen als fundamentale Naturkonstanten:

- die Lichtgeschwindigkeit im Vakuum $c = 299\,792\,458 \text{ m s}^{-1}$,
- die Gravitationskonstante $G = 6.67408 \times 10^{-11} \text{ m}^3/\text{kg s}^2$,
- das reduzierte Planck'sche Wirkungsquantum $\hbar = 1.054\,571\,800 \times 10^{-34} \text{ J s}$ (es gilt $1 \text{ J} = 1 \text{ m}^2 \text{ kg s}^{-2}$).

Ermitteln Sie aus diesen drei Größen die sogenannte Planck-Länge l_P , die Planck-Masse m_P und die Planck-Zeit t_P , indem sie passende Kombinationen der Naturkonstanten bilden.



Optimale Geschwindigkeiten und optimale Funktionen – Optimierungen und Variationen

2

Inhaltsverzeichnis

2.1 Mit welcher Geschwindigkeit erhält man den optimalen Verkehrsfluss? Ein Optimierungsproblem.....	39
2.2 Woher kennen die Physiker die genaue Form des Kabels? – Die Kettenlinie als Minimierer	44
Aufgaben	57

2.1 Mit welcher Geschwindigkeit erhält man den optimalen Verkehrsfluss? Ein Optimierungsproblem

In diesem Kapitel wird ein Klassiker unter den Modellierungsaufgaben für die Schule vorgestellt, siehe z.B. [37, 82]. Das Problem lautet: Unter welchen Bedingungen kann eine Straße am meisten Autoverkehr aufnehmen? Das muss zunächst präzisiert werden: Gemeint ist nicht, wann die meisten Autos auf der Straße sind, sondern unter welchen Bedingungen eine Straße den größten Verkehrsfluss zulässt, oder anders formuliert: Unter welchen Bedingungen können am meisten Autos pro Zeit eine bestimmte Stelle der Straße passieren? Uns interessiert hier die Situation, dass der Verkehr sehr dicht ist und durch Geschwindigkeitsvorgaben optimiert werden soll. Klar ist: Je schneller ein Auto fährt, desto weniger Zeit benötigt es, um einen Streckenabschnitt zu passieren, insofern wären hohe Geschwindigkeiten günstig. Andererseits benötigen Autos bei hohen Geschwindigkeiten einen größeren Abstand voneinander, es passen also weniger Autos zur gleichen Zeit auf einen bestimmten Streckenabschnitt, das würde für eine niedrige Geschwindigkeit sprechen. Irgendwo in der Mitte sollte die optimale Geschwindigkeit liegen. Um überhaupt von *einer* Geschwindigkeit sprechen zu können, nehmen wir an, dass sich alle Autofahrer präzise an die vorgegebene Geschwindigkeit v halten und auch alle den gleichen Sicherheitsabstand s zum Vordermann einhalten. Wir stellen uns vor, dass Auto A gerade eine Stelle passiert hat, und berechnen die

Zeit, die mindestens vergehen muss, bis das folgende Auto B dieselbe Stelle passieren kann, wenn der Sicherheitsabstand eingehalten wird: Auto A muss die Strecke $s + l$ zurückgelegt haben, wobei l die Länge von Auto A ist. Dafür benötigt es die Zeit $t = \frac{s+l}{v}$. Diese Zeit muss minimiert werden, damit möglichst schnell das nächste Auto dieselbe Stelle passieren kann. Der Sicherheitsabstand s soll von der Geschwindigkeit abhängen, hohe Geschwindigkeiten benötigen einen größeren Abstand als geringe Geschwindigkeiten. Welcher funktionale Zusammenhang zwischen v und s bietet sich an? Aus der Praxis sind zwei Modelle bekannt.

2.1.1 Das Modell „Halber Tacho“

In den Fahrschulen wird die „Halber Tacho“-Regel gelehrt: Die Maßzahl des in Meter gemessenen Abstands soll die Hälfte der Maßzahl der in km/h gemessenen Geschwindigkeit betragen. Der Zusammenhang soll also linear sein, $s = \alpha v$, und die angesprochene Regel führt zu der Konstanten

$$\alpha = \frac{\text{m}}{2\text{km/h}} = 1.8 \text{ s.}$$

Mit anderen Worten: Als Sicherheitsabstand wird die Strecke vorgeschlagen, die bei der vorherrschenden Geschwindigkeit in 1,8 s zurückgelegt wird. Für die zu minimierende Zeit ergibt sich

$$t(v) = \frac{\alpha v + l}{v} = \alpha + \frac{l}{v}.$$

Diese Funktion ist streng monoton fallend. Es sollte also so schnell wie möglich gefahren werden. Damit alle Verkehrsteilnehmer dieselbe Geschwindigkeit fahren können, muss man sich am Langsamsten orientieren. Nach diesem Modell würde also eine Geschwindigkeit von 100 km/h ohne LKWs oder 80 km/h mit LKWs zum optimalen Verkehrsfluss führen.

Welche Verkehrsflüsse lassen sich damit pro Fahrspur erreichen? Gehen wir von einem Verkehr mit LKWs, also der Geschwindigkeit $v = 80 \text{ km/h}$ aus, benötigt ein Auto der Länge l mitsamt dem angehängten Sicherheitsabstand eine Zeit von $t = 1.8 \text{ s} + \frac{3.6 \cdot l}{80} \frac{\text{s}}{\text{m}}$. Schätzen wir die durchschnittliche Länge aller Verkehrsteilnehmer auf 10 m, ergibt sich ein durchschnittlicher Fluss von einem Auto pro 2.25 s und Spur, oder anders ausgedrückt, von 0.44 Autos pro Sekunde und Spur. Eine Erhöhung der gemeinsamen Geschwindigkeit auf 100 km/h würde den Verkehrsfluss lediglich geringfügig auf 0.46 Autos pro Sekunde erhöhen.

2.1.2 Das Modell „Anhalteweg“

Plausibel wäre es auch, den Sicherheitsabstand so zu wählen, dass der Hinterherfahrende im Notfall auf dieser Strecke zum Stehen kommen könnte, man würde also den sogenannten *Anhalteweg* wählen. Nach den üblichen Modellen der Verkehrsphysik setzt sich dieser Weg zusammen aus dem *Reaktionsweg* und dem *Bremsweg*. Der Reaktionsweg ist die Strecke,

die in der Zeit mit der ursprünglichen Geschwindigkeit v zurückgelegt wird, bis die Fahrerin realisiert, dass vor ihr ein Hindernis ist, und die Bremsung einleitet. Üblicherweise wird für diese Zeit eine „Schrecksekunde“ angesetzt, also $t_R = 1$ s. Der Reaktionsweg s_R hat somit die Länge $s_R(v) = t_R v$. Für den Bremsweg wird üblicherweise das Modell der gleichmäßig beschleunigten (hier abgebremsten) Bewegung gewählt. Eigentlich wollen wir die Strecke berechnen, die benötigt wird, um einen Gegenstand bei gleichmäßiger Bremsbeschleunigung $-a$ von der Geschwindigkeit v auf die Geschwindigkeit null zu bringen. Die Physik erlaubt es aber, die Situation quasi rückwärts zu betrachten: Es handelt sich um die gleiche Strecke, die benötigt wird, um einen Gegenstand bei gleichmäßiger Beschleunigung a aus der Ruhe auf die Geschwindigkeit v zu bringen. Es ergibt sich die Formel $s_B(v) = \frac{v^2}{2a}$ und damit das Modell

$$t(v) = \frac{l + vt_R + \frac{v^2}{2a}}{v} = \frac{l}{v} + t_R + \frac{v}{2a},$$

siehe Aufg. 2.1. Um die Funktion t über dem Intervall $(0, \infty)$ zu minimieren, berechnen wir die Ableitung

$$t'(v) = -\frac{l}{v^2} + \frac{1}{2a}$$

mit der einzigen positiven Nullstelle $v_e = \sqrt{2al}$. Wegen der Grenzwertbeziehungen $\lim_{v \rightarrow 0} t(v) = \lim_{v \rightarrow \infty} t(v) = \infty$ liegt an dieser Stelle das absolute Minimum vor. Die optimale Geschwindigkeit hängt also zum einen von der erreichbaren Bremsverzögerung a ab. Dieser Wert variiert mit der Straßen- und Reifenbeschaffenheit. Bei trockener Straße und verkehrstauglichen Reifen wird hier ein Wert von $a = 5 \text{ m/s}^2$ genannt, siehe z. B. [24]. Das Ergebnis – bei guter Bremsleistung sind höhere Geschwindigkeiten besser – deckt sich wohl mit der Alltagsintuition.

Schwieriger zu verstehen ist die Abhängigkeit von der Fahrzeuglänge l . Bei längeren Verkehrsteilnehmern ist das Verhältnis der geschwindigkeitsunabhängigen Fahrzeuglänge zum geschwindigkeitsabhängigen Anhalteweg größer als bei kurzen Autos. Für einen reinen LKW-Verkehr wäre nach diesem Modell eine höhere Geschwindigkeit optimal als für einen reinen PKW-Verkehr. In Mischverkehren müsste man sich im Sinne der Sicherheit nach den kürzesten Autos richten.

Setzen wir für einen PKW eine Mindestlänge von 4 m an und rechnen mit dem erwähnten Wert $a = 5 \text{ m/s}^2$ erhalten wir als optimale Geschwindigkeit $v_e = \sqrt{40 \text{ m}^2 \text{ s}^{-2}} = 6.32 \text{ m/s} = 22.8 \text{ km/h}$, eine erschreckend geringe Geschwindigkeit. Wir gehen wieder davon aus, dass der Durchschnitt über die Fahrzeuglängen aller Verkehrsteilnehmer 10 m beträgt. Dann benötigt ein Auto im Schnitt die Zeit

$$t = \frac{10}{6.23} \text{ s} + 1 \text{ s} + \frac{6.32}{10} \text{ s} = 3.21 \text{ s},$$

was einem Fluss von 0.31 Auto pro Sekunde und Spur entspricht.

Für einen reinen LKW-Verkehr mit Mindestlänge $l = 10$ m ergibt sich $v_e = 36$ km/h. Gehen wir von einer durchschnittlichen LKW-Länge von 15 m aus, erhalten wir eine Zeit von

$$t = \frac{15}{10} \text{ s} + 1 \text{ s} + \frac{10}{10} \text{ s} = 3.5 \text{ s},$$

was einen Fluss von ungefähr 0.29 LKWs pro Sekunde und Spur ergibt.

2.1.3 Bewertung

Die beiden Modelle liefern stark unterschiedliche Aussagen, an welches sollte man sich also halten? Das physikalische Modell des Anhaltewegs ist empirisch stark abgesichert und liefert insofern eine gute Prognose für den Anhalteweg.

In der Praxis ist es jedoch nur äußerst selten nötig, den gesamte Anhalteweg zur Verfügung zu haben, weil der Vordermann ja in der Regel nicht instantan zum Stehen kommt, sondern auch einen Bremsweg zurücklegt, den der Nachfolger ebenfalls zum Bremsen nutzen kann. Aus dieser Perspektive wäre nur der reine Reaktionsweg nötig und die „Halber-Tacho“-Regel stellt sogar einen Reaktionsweg für eine Reaktionszeit von 1.8 s zur Verfügung. Andererseits sind damit bei sehr abrupten Bremsmanövern des direkten Vordermanns Auffahrunfälle unvermeidlich: Betrachten wir die „Halber-Tacho“-Regel, ergibt sich, dass bei Veranschlagung einer Schrecksekunde lediglich 0,8 s verbleiben, um auf eine Störung zu reagieren. Legen wir wieder das physikalische Modell der gleichmäßig beschleunigten Bewegung mit $a = \frac{5 \text{ m}}{\text{s}^2}$ zugrunde, findet in 0.8 s lediglich eine Geschwindigkeitsreduktion um $v_d = 4 \text{ m s}^{-1} = 14.6 \text{ km/h}$ statt. Der Autofahrer passiert also die Stelle, an der die Störung stattgefunden hat, mit nur einer geringfügig verringerten Geschwindigkeit. Zusammenfassend lässt sich sagen, dass ein recht hoher Fluss erreichbar ist, dieser aber ziemlich anfällig gegenüber abrupten Bremsmanövern ist.

Das Modell des vollen Anhaltewegs liefert sehr viel Sicherheitsreserve: Man käme selbst dann noch rechtzeitig zum Stehen, wenn der Vordermann instantan stehen bliebe, was in der Praxis gar nicht möglich ist, es sei denn, es fiele eine massive Mauer vom Himmel auf ihn drauf. Der erreichte Fluss von 0.31 Fahrzeugen pro Sekunde und Spur ist bedeutend niedriger als im anderen Modell. Allerdings sind nach diesem Modell Unfälle sehr unwahrscheinlich.

Praktikabel erscheint die Variante, dass man den Anhalteweg nicht auf den direkten Vordermann, sondern auf den Vorvorgänger bezieht. Die erreichbaren Verkehrsflüsse sollen in Aufg. 2.2 berechnet werden.

Das Modell „Anhalteweg“ liefert allerdings auch den Hinweis, dass es günstig sein könnte, sehr große Verkehrsströme, die mehrere Fahrbahnen benutzen, nach der Fahrzeuglänge auf die verschiedenen Spuren aufzuteilen und für jede Spur die Geschwindigkeit einzeln zu optimieren. Es ist allerdings äußerst fraglich, ob ein solches Modell ausreichend Verständnis und Akzeptanz bei den Verkehrsteilnehmern erzielen würde.

Eine letzte Bemerkung betrifft die empirisch gemessenen erreichten Verkehrsflüsse. Diese scheinen bei einer Geschwindigkeit von ca. 90 km/h am höchsten zu sein und Flüsse von 5000 Autos pro Stunde auf einer dreispurigen Autobahn zu erreichen (Autobahn A3 mit Tempolimit), das entspricht einem Fluss von 0.46 Autos pro Spur und Sekunde, siehe [94]. In den empirischen Daten in Abb. 2.1 lassen sich zwei „Datenwolken“ erkennen. Die Datenpunkte in der oberen scheinen für einen gut fließenden Verkehr zu sprechen, während in den Datenpunkten, die diagonal verteilt sind, von sehr dicht, bis zäh fließendem und stockendem Verkehr zu sprechen ist. Aus den Geschwindigkeits- und Verkehrsflussdaten können wir auf eine Verkehrsdichte, also Autos pro Strecke, schließen. Der Fluss ϕ (in Abb. 2.1 wird er Verkehrsstärke genannt und mit q bezeichnet) hat die Dimension Autos pro Zeit, die Geschwindigkeit v hat die Dimension Länge pro Zeit: Auch ohne Intuition und Anschauung kann man aus den Dimensionen erkennen, dass sich die Verkehrsdichte ρ nach der Formel $\rho = \frac{\phi}{v}$ berechnen lässt. Die in Abb. 2.1 dargestellten empirischen Daten deuten darauf hin, dass für den betrachteten Autobahnabschnitt mit drei Fahrspuren bei den bei der Messung vorliegenden Witterungs- und Lichtverhältnissen ungefähr die Dichte von $\rho = 4000/80$ Autos pro Kilometer, also 50 Autos pro Kilometer, kritisch ist. Bei dieser Dichte kippt der Verkehrsfluss von einem frei fließenden zu einem mehr oder weniger stockenden Verkehr um.

Wahrscheinlich kennen Sie das Phänomen eines „Staus aus dem Nichts“. Eine didaktisch aufbereitete Modellierung dieser Situation nach [66] findet sich in [70].

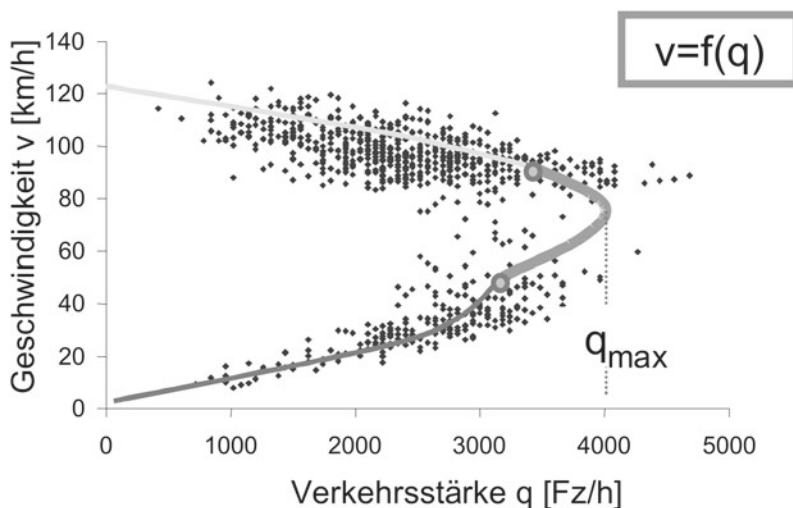


Abb. 2.1 Empirische Messdaten des Verkehrsflusses auf der A3 (Abbildung aus [94], ©(2000) Kirschbaum)

2.2 Woher kennen die Physiker die genaue Form des Kabels? – Die Kettenlinie als Minimierer

Es ist auf den ersten Blick sehr erstaunlich, dass der Graph des hyperbolischen Kosinus die Form einer aufgehängten Kette so gut modelliert, und es stellt sich die Frage, woher die Physiker wissen, dass es genau diese Funktion sein muss. Haben sie das durch experimentelles Probieren herausgefunden oder stehen theoretische Überlegungen dahinter? Letzteres ist der Fall. Die Physiker haben dazu ein Prinzip sehr konsequent angewendet, nämlich das sogenannte *Minimierungsprinzip*: Die Energie eines stationären Zustands ist minimal. Was soll das bedeuten? Hängen wir ein Kabel an zwei Stellen x_1 und x_2 auf den Höhen H_1 bzw. H_2 ein und rütteln daran, kann es alle möglichen Höhenverläufe annehmen. Wir übernehmen die Sprechweise der Physiker und werden im Folgenden von „Zuständen“ sprechen. Hören wir jedoch mit dem Rütteln auf, wird das Kabel immer im selben Zustand zur Ruhe kommen, einem sogenannten *stationären Zustand*. In jedem möglichen Zustand, den es annehmen kann, ist mit dem Kabel eine bestimmte Energie verbunden. Das Minimierungsprinzip sagt aus, dass der schlussendlich angenommene stationäre Zustand die minimale Energie unter allen möglichen zugänglichen Zuständen hat.

Wir werden im Folgenden, solange nicht explizit etwas anderes angegeben ist, stillschweigend davon ausgehen, dass alle auftretenden Funktionen *glatt*, d. h. nach allen Variablen beliebig oft differenzierbar sind.

2.2.1 Das Energiefunktional

Was sind nun die möglichen Zustände, die das Kabel annehmen kann, und wie sieht die zugehörige Energie aus? Mit Zuständen meinen wir Höhenverläufe des Kabels und wir wollen annehmen, dass nur glatte Zustände möglich sind, das Kabel also keine scharfen Knicke hat. Einen möglichen Zustand können wir also durch eine stetig differenzierbare Funktion über dem Intervall $[x_1, x_2]$ beschreiben, welche die Einspannbedingung erfüllt. Wir fassen solche Funktionen zur Menge der zugänglichen (*accessible*) Zustände zusammen:

$$\mathcal{A} = \{ G : [x_1, x_2] \longrightarrow \mathbb{R} : G(x_1) = H_1, G(x_2) = H_2 \} \quad (2.1)$$

Nun müssen wir uns überlegen, welche Energie mit einem Zustand $G \in \mathcal{A}$ verbunden ist. Dazu denken wir uns das Kabel zunächst in viele kleine Einzelteile zerlegt, wir können z. B. an die Kugeln der Kugelmkette aus dem Experiment in **Kap. 1.1.5** denken, siehe Abb. 1.6. Aus dem Physikunterricht ist vielleicht noch die potentielle Energie oder auch Lageenergie bekannt, die eine Kugel hat, die sich in der Höhe H befindet: Diese beträgt mgH , wobei m die Masse der Kugel und g die Erdbeschleunigung ist. Eine andere Energieform spielt in dieser Situation keine Rolle. Um die Energie der gesamten Kugelmkette zu berechnen, müssen wir einfach über die Energie aller einzelnen Kugeln summieren. In unserem Modell hat das Kabel aber einen kontinuierlichen Verlauf und die Masse ist zu einer Massendichte pro

Länge ϱ verschmiert, das Kabel hat also z. B. eine Massendichte von 1 kg/m. Anstatt über einzelne Kugeln zu summieren, müssen wir also über den Höhenverlauf integrieren und dabei noch die Länge des Kabelstückchens berücksichtigen, über das wir gerade integrieren: Ein Kabelstückchen über dem kleinen Intervall $[x, x+h]$, das genau waagrecht verläuft, hat eine geringere Länge als ein Kabelstückchen über demselben Intervall, das schräg verläuft. Diese Frage wurde bereits in Aufg. 1.1 diskutiert. Insgesamt erhalten wir für einen Zustand G die zugehörige Energie

$$E(G) = \varrho g \int_{x_1}^{x_2} G(x) \sqrt{1 + G'(x)^2} dx. \quad (2.2)$$

Mathematisch gesprochen ist die Energie E eine Abbildung aus der Menge der zugänglichen Zustände $\mathcal{A}$ in die reellen Zahlen (eigentlich in den skalaren Größenbereich Energie), $E : \mathcal{A} \rightarrow \mathbb{R}$. Man spricht bei solchen Abbildungen, die ganze Funktionen auf reelle Zahlen abbilden, auch von Funktionalen, hier also von dem Energiefunktional. Das Minimierungsprinzip sagt nun aus, dass der physikalisch angenommene Zustand H ein Minimierer des Funktional E über der Menge $\mathcal{A}$ ist, dass also $E(H) \leq E(G)$ für alle $G \in \mathcal{A}$ gilt.

2.2.2 Ein Minimierer des Energiefunktional und die Euler-Lagrange-Gleichungen

Minimierungsaufgaben kennen wir aus der Schule und der Analysis-Vorlesung: An einer Minimalstelle im Inneren des Definitionsbereichs der Funktion muss die Ableitung verschwinden. Das Problem ist aber nun, dass unsere Variablen selbst schon Funktionen sind. Wir können also nicht einfach einen Differenzenquotienten aufstellen, dazu müssten wir ja durch eine Funktion teilen. Aber bereits bei der Differentialrechnung im Mehrdimensionalen wird das Konzept einer Richtungsableitung benutzt, das sich hier verallgemeinern lässt: Wir probieren *in die Richtung einer Funktion* abzuleiten. Dafür fixieren wir eine Richtungsfunktion $\varphi \in C^1([x_1, x_2]; \mathbb{R})$, diese wird üblicherweise *Testfunktion* genannt, und betrachten die Ableitung von E in Richtung φ . In diesem Zusammenhang haben sich ein neues Wort und ein neues Symbol eingebürgert, wir sprechen von der *Variation des Funktional E an der Stelle G in Richtung φ* :

$$\lim_{s \rightarrow 0} \frac{E(G + s\varphi) - E(G)}{s} =: \delta E(G, \varphi) \quad (2.3)$$

Wir haben allerdings eine Sache übersehen: Der leicht geänderte Zustand $G + s\varphi$ muss die Einspannbedingung erfüllen. Wir müssen also $\varphi(x_1) = \varphi(x_2) = 0$ verlangen. Wir fassen nun alle möglichen Richtungen, in die wir ableiten können, wieder in einer Menge zusammen und werden diese Menge die Menge der Testfunktionen nennen:

$$\mathcal{T} := \{ \varphi : [x_1, x_2] \rightarrow \mathbb{R} : \varphi(x_1) = \varphi(x_2) = 0 \}$$

Stillschweigend setzen wir im Moment voraus, dass der auftretende Grenzwert existiert. Wir werden aber gleich zeigen, dass das tatsächlich der Fall ist. Jetzt betrachten wir den physikalisch angenommenen stationären Zustand H des Systems, der das Energiefunktional E minimiert. Für diesen Zustand und eine beliebige Testfunktion φ gilt $E(H) \leq E(H + s\varphi)$ für alle $s \in \mathbb{R}$. Also hat die Funktion $\Phi(s) = E(H + s\varphi)$ ein Minimum in $s = 0$ und, falls Φ in Null differenzierbar ist, gilt notwendigerweise $\Phi'(0) = \delta E(H, \varphi) = 0$. Für einen Minimierer H des Funktionals E über $\mathcal{A}$ gilt also:

$$\delta E(H, \varphi) = 0 \quad \text{für alle Testfunktionen } \varphi \in \mathcal{T}. \quad (2.4)$$

Können wir aus dieser Information schlussfolgern, dass H ein Kosinushyperbolicus ist? Dazu wollen wir $\delta E(H, \varphi)$ genauer auszurechnen. Die folgende Rechnung hängt zunächst einmal nicht von der konkreten Definition des Energiefunktionals ab und wir wollen sie so allgemein wie möglich durchführen. In das Funktional gehen die Funktion G und ihre Ableitung G' ein. Wir können also schreiben:

$$E(G) = \int_{x_1}^{x_2} \mathcal{L}(G(x), G'(x)) \, dx \quad \text{mit} \quad \mathcal{L}(p, q) = \rho g p \sqrt{1 + q^2}. \quad (2.5)$$

Die Funktion $\mathcal{L} : \mathbb{R}^2 \rightarrow \mathbb{R}$ nennt man die *Lagrange-Dichte* des Funktionals E . Wir berechnen jetzt für eine allgemeine Lagrange-Dichte und eine beliebige Testfunktion $\varphi \in \mathcal{T}$ die Variation und beginnen dafür mit dem Differenzenquotienten. Für $s \neq 0$ gilt:

$$\frac{E(G + s\varphi) - E(G)}{s} = \int_{x_1}^{x_2} \frac{1}{s} \left(\mathcal{L}(G(x) + s\varphi(x), G'(x) + s\varphi'(x)) - \mathcal{L}(G(x), G'(x)) \right) dx \quad (2.6)$$

Um den Grenzübergang $s \rightarrow 0$ übersichtlich durchführen zu können, entwickeln wir zunächst die Lagrange-Dichte nach Taylor um eine Stelle (p, q) bis zur ersten Ordnung und verwenden dabei die sogenannte Lagrange-Restglieddarstellung:

$$\begin{aligned} \mathcal{L}(p + h_1, q + h_2) &= \mathcal{L}(p, q) + \partial_p \mathcal{L}(p, q) h_1 + \partial_q \mathcal{L}(p, q) h_2 \\ &\quad + \frac{1}{2} \partial_p^2 \mathcal{L}(*) h_1^2 + \partial_p \partial_q \mathcal{L}(*) h_1 h_2 + \frac{1}{2} \partial_q^2 \mathcal{L}(*) h_2^2, \end{aligned} \quad (2.7)$$

wobei $*$ = $(p + \tau h_1, q + \tau h_2)$ und $\tau = \tau(p, q, h_1, h_2) \in (0, 1)$. Für eine Herleitung der mehrdimensionalen Taylor-Formel aus der eindimensionalen Taylor-Formel siehe Aufg. 2.5. Wir verwenden nun diese Entwicklung mit $p = G(x)$, $h_1 = s\varphi(x)$, $q = G'(x)$ und $h_2 = s\varphi'(x)$ und formen die Gl. (2.6) weiter um:

$$\begin{aligned} &= \int_{x_1}^{x_2} \partial_p \mathcal{L}(G(x), G'(x)) \varphi(x) + \partial_q \mathcal{L}(G(x), G'(x)) \varphi'(x) \, dx \\ &\quad + s \int_{x_1}^{x_2} \frac{1}{2} \partial_p^2 \mathcal{L}(**) \varphi^2(x) + \partial_p \partial_q \mathcal{L}(**) \varphi(x) \varphi'(x) + \frac{1}{2} \partial_q^2 \mathcal{L}(**) \varphi'(x)^2 \, dx, \end{aligned}$$

wobei $** = (G(x) + \tau s \varphi(x), G'(x) + \tau s \varphi'(x))$ und $\tau = \tau(G(x), G'(x), \varphi(x), \varphi'(x), s) \in (0, 1)$. Der erste Summand ist unabhängig von s . Der Integrand des zweiten Integrals, der über die Argumente $**$ von s abhängt, ist im Betrag beschränkt, unabhängig von $|s| \leq 1$. Das liegt daran, dass die stetigen Funktionen $\partial_p^2 \mathcal{L}$, $\partial_p \partial_q \mathcal{L}$ und $\partial_q^2 \mathcal{L}$ für die Integration nur über einer kompakten Mengen im $\mathbb{R}^2$ ausgewertet werden müssen, und über kompakten Mengen sind stetige Funktionen stets beschränkt. Damit können wir den Grenzübergang $s \rightarrow 0$ problemlos durchführen und erhalten

$$\delta E(G, \varphi) = \int_{x_1}^{x_2} \partial_p \mathcal{L}(G(x), G'(x)) \varphi(x) + \partial_q \mathcal{L}(G(x), G'(x)) \varphi'(x) dx. \quad (2.8)$$

Für den Minimierer H verschwindet dieser Ausdruck gemäß Gl. (2.4) für jede Testfunktion $\varphi \in \mathcal{T}$. Können wir aus dieser Information die Funktion H bestimmen? Das wesentliche Werkzeug, das an dieser Stelle benutzt werden muss, nennt sich *Fundamentallemma der Variationsrechnung* und erklärt auch ein bisschen den Ausdruck *Testfunktion*.

Satz 2.1 (Das Fundamentallemma der Variationsrechnung)

Für eine stetige Funktion $G : [x_1, x_2] \rightarrow \mathbb{R}$ gelte

$$\int_{x_1}^{x_2} G(x) \varphi(x) dx = 0 \quad \text{für alle Test funktionen } \varphi \in \mathcal{T}.$$

Dann ist G die Nullfunktion.

Nach einem Beweis dieser Version des Fundamentallemmas wird in Aufg. 2.6 gefragt.

Satz 2.1 ist eine Version des Fundamentallemmas, die genau auf unsere momentanen Zwecke zugeschnitten ist. Das Lemma gilt allerdings in viel weitreichenderen Zusammenhängen: So kann das Intervall $[x_1, x_2]$ durch ein beliebiges Gebiet $\Omega \subset \mathbb{R}^n$ ersetzt werden. Es können sehr viel mehr Funktionen G getestet werden, üblich sind hier die lokal Lebesgue-integrierbaren Funktionen (streng genommen handelt es sich dabei gar nicht um Funktionen, sondern um bestimmte Äquivalenzklassen von Funktionen). Es reicht, mit sehr viel weniger Funktionen zu testen, der üblicherweise benutzte Testraum besteht aus den beliebig oft differenzierbaren Funktionen mit kompaktem Träger.

Zurück zu unserem Problem: Wir können das Fundamentallemma noch nicht auf Gl. (2.8) anwenden, weil hier einmal gegen φ und einmal gegen φ' integriert wird. Hier hilft die gute alte Schulformel von der partiellen Integration: $\int_a^b f'(x)g(x) dx = f(x)g(x)|_a^b - \int_a^b f(x)g'(x) dx$. Die Rolle von f nimmt die Funktion φ ein, die Rolle von g nimmt die Funktion $x \mapsto \partial_q \mathcal{L}(G(x), G'(x))$ ein. Weil φ an den Rändern verschwindet, verschwindet auch der Randterm und wir erhalten

$$\delta E(G, \varphi) = \int_{x_1}^{x_2} \left(\partial_p \mathcal{L}(G(x), G'(x)) - \frac{d}{dx} (\partial_q \mathcal{L}(G(x), G'(x))) \right) \cdot \varphi(x) dx. \quad (2.9)$$

Setzen wir jetzt für das G den Minimierer H , folgt aus (2.4) und dem Fundamentallemma der Variationsrechnung, dass H die sogenannten *Euler-Lagrange-Gleichungen* löst. Das soll in Form eines Satzes zusammengefasst werden:

Satz 2.2

Gegeben sei ein Intervall $[x_1, x_2] \subset \mathbb{R}$, zwei Zahlen $H_1, H_2 \in \mathbb{R}$, eine Lagrange-Dichte $\mathcal{L} : \mathbb{R}^2 \rightarrow \mathbb{R}$. Die Funktion $H : [x_1, x_2] \rightarrow \mathbb{R}$, sei ein Minimierer des Funktional

$$E(G) = \int_{x_1}^{x_2} \mathcal{L}(G(x), G'(x)) dx$$

über der Menge $\mathcal{A}$ aus (2.1). Dann löst H die Euler-Lagrange-Gleichungen

$$\partial_p \mathcal{L}(G(x), G'(x)) - \frac{d}{dx} (\partial_q \mathcal{L}(G(x), G'(x))) = 0 \quad \text{für alle } x \in [x_1, x_2]. \quad (2.10)$$

2.2.3 Die Euler-Lagrange-Gleichungen für das Kabel passen (noch) nicht

Jetzt spezifizieren wir endlich wieder auf unsere konkrete Lagrange-Dichte

$$\mathcal{L}(p, q) = \rho g p \sqrt{1 + q^2} \quad (2.11)$$

und rechnen mit einer etwas länglichen (und fehleranfälligen) Rechnung nach, dass der Minimierer H des Energiefunktional (2.2) die Gleichung

$$H'' H = 1 + H'^2 \quad (2.12)$$

löst, siehe Aufg. 2.7. Diese Gleichung stiftet eine Beziehung zwischen der Funktion H und ihren ersten beiden Ableitungen. Man spricht von einer Differentialgleichung zweiter Ordnung – zweite Ordnung, weil die höchste auftretende Ableitung die zweite ist. Die *Theorie der gewöhnlichen Differentialgleichungen* beschäftigt sich mit der Lösbarkeit und dem konkreten Berechnen von Lösungen solcher Gleichungen. Wir werden einiges davon in den Kapiteln 5 und 6 kennenlernen. Wir können zwar mit unseren Mitteln keine Lösung für (2.12) berechnen, aber überprüfen, ob die Modellfunktionen aus Abschn. 1.1 diese Differentialgleichung lösen. Für den parabelförmigen Höhenverlauf soll das in Aufg. 2.7 untersucht (und verneint) werden, hier betrachten wir den hyperbolischen Kosinus mit seinen Verschiebungen, Streckungen und Stauchungen

$$H(x) = k \cosh\left(\frac{x - \bar{x}}{k}\right) + H_0 \quad \text{mit } \bar{x}, k, H_0 \in \mathbb{R} \quad \text{und } k \neq 0 \quad (2.13)$$

und rechnen nach:

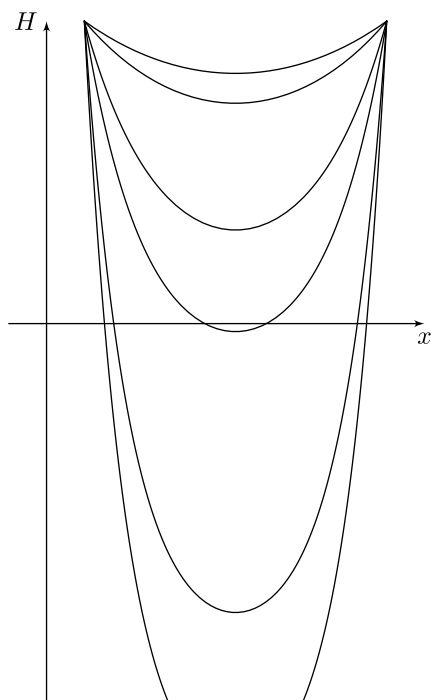
$$(H''H - (1 + H'^2))(x) = \cosh\left(\frac{x - \bar{x}}{k}\right) \cdot \left(\cosh\left(\frac{x - \bar{x}}{k}\right) + H_0\right) - \cosh^2\left(\frac{x - \bar{x}}{k}\right)$$

Dieser Ausdruck ist aber genau dann null, wenn $H_0 = 0$ ist. Das hieße aber, dass der Kabelverlauf nur dann ein Minimierer sein könnte, wenn sein Tiefpunkt sich auf der Höhe null befindet. Das ergibt überhaupt keinen Sinn, weil wir ja das Nullniveau der Höhe nach eigenem Belieben festlegen können. Wir müssen also irgendwo einen Denk- oder einen Rechenfehler begangen haben.

2.2.4 Die Kabellänge ist vorgegeben: Minimierung unter einer Nebenbedingung

Der Denkfehler ist schon am Anfang passiert: Das in (2.2) angegebene Funktional hat überhaupt keinen Minimierer in der Menge der zugänglichen Zustände $\mathcal{A}$. Wir können nämlich zu jedem beliebigen negativen Energieniveau E_0 einen Höhenverlauf aus $\mathcal{A}$ angeben, dessen Energie kleiner ist. Wie das geht, ist in Abb. 2.2 angedeutet, der Höhenverlauf muss sich nur über eine hinreichende Strecke in einer hinreichend negativen Höhe befinden.

Abb. 2.2 Das Funktional (2.2) ist über $\mathcal{A}$ aus (2.1) nach unten unbeschränkt, es gibt keinen Minimierer



Solche Höhenverläufe haben natürlich nichts mit der physikalischen Situation zu tun, die wir eigentlich vor Augen haben. Sie unterscheiden sich dadurch, dass die Kabellängen, also die Bogenlängen der Höhenverlaufsfunktion, beliebig groß werden müssen. Es liegt aber eine physikalische Situation vor, in der die Kabellänge fest vorgegeben ist. Realisiert wird der Minimierer des Energiefunktional E unter allen Höhenverläufen G , welche die Nebenbedingung $L(G) = L_0$ erfüllen, wobei $L(G)$ gerade die Bogenlänge ist:

$$L(G) = \int_{x_1}^{x_2} \sqrt{1 + G'^2(x)} dx \quad (2.14)$$

Das stellt uns aber vor ein neues Problem: Der neue Raum der zulässigen Funktionen

$$\tilde{\mathcal{A}} := \{G \in \mathcal{A} : L(G) = L_0\}$$

ist kein affin-linearer Raum mehr. Addieren wir nämlich zu einer Funktion $G \in \tilde{\mathcal{A}}$ eine Testfunktion $\varphi \in \mathcal{T}$, so wird die Summe in der Regel kein Element von $\tilde{\mathcal{A}}$ sein, weil sich die Bogenlänge im Allgemeinen ändert. Die Euler-Lagrange-Maschine lässt sich nicht ohne Weiteres anwenden. Um Anregungen für eine Lösung dieses Problems zu erhalten, ist es sinnvoll, sich den endlich-dimensionalen Fall und hier am einfachsten den zweidimensionalen Fall anzuschauen.

2.2.5 Extrema unter Nebenbedingungen im zweidimensionalen Fall: Lagrange'sche Multiplikatoren

Tatsächlich sind Probleme solcher Art schon aus der Schule bekannt: Man sucht z. B. den Zylinder mit dem vorgegebenen Volumen V_0 , dessen Oberfläche minimales Oberflächenmaß hat. Mit den Variablen r für den Radius der kreisförmigen Grundfläche und der Höhe h möchte man also die Funktion $O(r, h) = 2\pi r^2 + 2\pi r h$ unter der Nebenbedingung $V(r, h) = \pi r^2 h = V_0$ minimieren. In der Schule wird nun in der Regel die Nebenbedingung genutzt, um die eine Variable mit Hilfe der anderen auszudrücken, also hier z. B. $h = V_0/(\pi r^2)$, und das Resultat in die zu minimierende Funktion einzusetzen, also $O(r) = 2\pi r^2 + 2V_0/r$, sodass eine zu minimierende Funktion in einer Variablen übrig bleibt. Im Allgemeinen ist dieses explizite Auflösen nach einer Variablen nicht möglich. Dann hilft die Methode der *Lagrange'schen Multiplikatoren*.

Zur Illustration der Methode der Lagrange'schen Multiplikatoren stellen wir uns unter der Funktion $f = f(x, y)$ einen Höhenverlauf vor, $f(x, y)$ könnte also z. B. die Höhe über dem Meeresspiegel am Punkt (x, y) sein. Eine zweite Funktion $g = g(x, y)$ könnte z. B. die Temperatur am Punkt (x, y) sein. Gesucht ist jetzt z. B. der höchste Punkt unter den Punkten, an denen die Temperatur $g = c$ herrscht, sagen wir z. B. $c = 7^\circ\text{C}$ (das Beispiel dient nur der Illustration, es soll nicht behauptet werden, dass so eine Fragestellung tatsächlich jemanden

interessiert). Auch hier können wir nicht einfach die kritischen Punkte von f berechnen, dort wird in der Regel nicht die Nebenbedingung gelten.

Den Höhenverlauf in der x - y -Ebene kann man durch Höhenlinien veranschaulichen, wie man sie von topographischen Landkarten kennt, und auch die Menge der Punkte, an denen die Temperatur $c = 7^\circ\text{C}$ vorliegt, wird eine glatte Kurve in der x - y -Ebene bilden, siehe Abb. 2.3. In dem uns interessierenden Punkt $P(x_P, y_P)$, in dem die Höhe f minimal unter der Nebenbedingung $g = c$ ist, müssen die durch diesen Punkt verlaufende Höhenlinie von f und die Kurve gleicher Temperatur $g = c$, Isotherme genannt, tangential zueinander sein. Wäre dem nicht so, würde sich ja die Höhe vergrößern, wenn man die Isotherme $g = c$ noch ein Stück in die richtige Richtung weiterverfolgen würde. Die Gradienten der Funktionen f und g stehen aber stets senkrecht auf den Höhenlinien bzw. den Isothermen, siehe Aufg. 2.8. Folglich müssen sie in dem gesuchten Extrempunkt unter Nebenbedingung ein lineares Vielfaches voneinander sein. Eine Einschränkung hat allerdings auch diese Methode: In dem Punkt $P(x_P, y_P)$ muss die Nebenbedingung g erst einmal einen nicht verschwindenden Gradienten besitzen, wir fordern also zusätzlich $\nabla g(x_P, y_P) \neq 0$.

Satz 2.3

Der Punkt $P(x_P, y_P)$ sei ein Extremum der Funktion f unter der Nebenbedingung $g = c$ und es gelte $\nabla g(x_P, y_P) \neq 0$. Dann gibt es ein $\lambda \in \mathbb{R}$, sodass $\nabla f(x_P, y_P) = \lambda \nabla g(x_P, y_P)$ gilt.

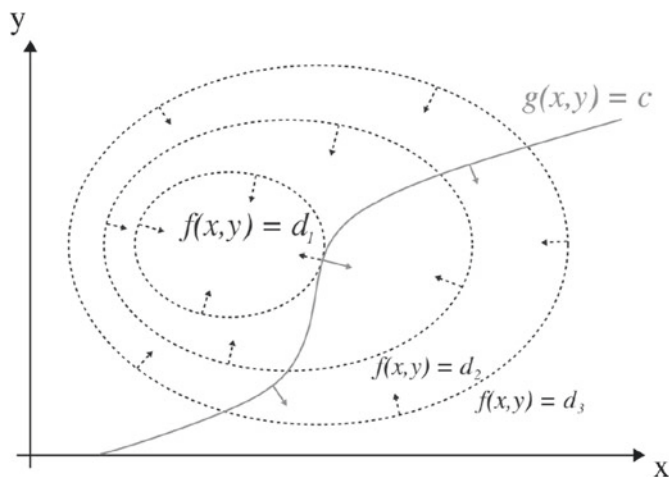


Abb. 2.3 In einem Extrempunkt unter Nebenbedingung sind die Gradienten der zu minimierenden Funktion und der die Nebenbedingung definierenden Funktion linear abhängig

Der Faktor λ ist der sogenannte *Lagrange-Multiplikator*: In Aufg. 2.9 sollen die hier gemachten anschaulichen Überlegungen auf ein solides mathematisches Fundament gestellt werden. Der gesuchte Punkt P muss also die folgenden Bedingungen erfüllen:

- Er erfüllt die Nebenbedingung $g(x_P, y_P) = g_0$.
- Es gibt ein $\lambda \in \mathbb{R}$, sodass $\nabla(f - \lambda g)(x_P, y_P) = 0$ gilt.

Wir wollen nun dieses Verfahren auf den uns interessierenden Fall von Funktionalen entsprechend übertragen.

2.2.6 Minimierer der Energie unter der Nebenbedingungen vorgegebener Kabellänge

Im endlich-dimensionalen Fall verschwinden alle Richtungsableitungen von $f - \lambda g$ im Punkt P . Wir übertragen diese Bedingungen ganz frech auf den uns interessierenden Fall und behaupten, dass die Variation in jede beliebige Testfunktionsrichtung verschwindet.

Satz 2.4

Die Funktion $H \in \mathcal{A}$ sei ein Minimierer des Funktionals E unter der Nebenbedingung $L = L_0$ und es gebe eine Testfunktion $\chi \in \mathcal{T}$ mit $\delta L(H, \chi) \neq 0$, dann gibt es eine Konstante $\lambda \in \mathbb{R}$, sodass $\delta(E - \lambda L)(H, \varphi) = 0$ für alle $\varphi \in \mathcal{T}$.

Die Voraussetzung, dass es ein $\chi \in \mathcal{T}$ mit $\delta L(H, \chi) \neq 0$ gibt, ist hier das Pendant zu der Voraussetzung $\nabla g(x_P, y_P) \neq 0$ im zweidimensionalen Fall. Die Analogiebildung zum zweidimensionalen Fall dient hier als Heuristik, um einen Satz zu vermuten, den Beweis dieser Vermutung kann sie natürlich nicht liefern. Der Beweis ist gut mit den Mitteln typischer Analysis-I/II-Vorlesungen durchführbar, als weitere Zutat geht der Satz über Umkehrabbildungen im Mehrdimensionalen ein, der für gewöhnlich in einer Analysis-II-Vorlesung enthalten ist. Für die interessierte Leserin (und für diejenigen, denen Lücken ein schlechtes Gefühl bereiten) wird der Beweis in Abschn. 2.2.7 vorgestellt, er ist aber nicht nötig für das Verständnis des weiteren Stoffs.

Wir wenden den Satz auf unser Problem an. Mit dem neuen Funktional $E - \lambda L$ haben wir es mit der Lagrange-Dichte

$$\tilde{\mathcal{L}}(p, q) = \rho g p \sqrt{1 + q} - \lambda \sqrt{1 + q^2} \quad (2.15)$$

zu tun. Wir schreiben $\lambda = \rho g H_0$ mit einer Konstante H_0 (das ist lediglich eine Umbenennung, die gleich ganz praktisch sein wird). Entscheidend ist, dass wir jetzt die bekannte Variationsmaschinerie wieder benutzen dürfen: Gilt für eine Funktion $H \in \mathcal{A}$, dass die Variation des erweiterten Funktionals verschwindet, also $\delta(E - \rho g H_0 L) = 0$, dann löst

H die Euler-Lagrange-Gleichungen zu der Lagrange-Dichte $\tilde{L}$. Hier führen längliche aber elementare Rechnungen zu der Gleichung

$$1 + H'^2 = (H - H_0)H'', \quad (2.16)$$

die tatsächlich genau durch Funktionen der Bauart $H(x) = k \cosh\left(\frac{x-\bar{x}}{k}\right) + H_0$ gelöst wird. Die Parameter $\bar{x}$, k und H_0 müssen nun so bestimmt werden, dass die Einspannbedingungen $H(x_1) = H_1$, $H(x_2) = H_2$ und die Kabellängenbedingung $L = L_0$ erfüllt werden.

Eine kleine Lücke bleibt: Um Satz 2.4 anwenden zu dürfen, darf der gefundene Minimierer H kein kritischer Punkt des Bogenlängenfunktional L sein. Es muss also eine Testfunktion $\chi \in \mathcal{T}$ geben, sodass $\delta L(H, \chi) \neq 0$ gilt. Solch eine Testfunktion können wir im Nachhinein (a posteriori) zu der jetzt bekannten Funktion H suchen und damit die Anwendung von Satz 2.4 legitimieren. Mit dieser Frage beschäftigt sich Aufg. 2.10.

2.2.7 Beweis: Lagrange-Multiplikator für Extrema unter Nebenbedingungen

Für den Beweis von Satz 2.4 benötigen wir den Satz über implizite Funktionen. Dabei geht es darum, dass sich durch eine Nebenbedingung eine Variable als Funktion der anderen Variablen ausdrücken lässt. Man sagt auch, dass man nach der einen Variablen *auföst*. In der Schule sind Extremwertaufgaben unter Nebenbedingungen in der Regel so gestellt, dass diese Auflösung explizit möglich ist. Vom theoretischen Standpunkt her gesehen ist das Entscheidende aber, dass diese Auflösung prinzipiell möglich ist, auch wenn sich die Auflösungsfunktion eventuell nicht explizit angeben lässt. Dabei kann die Auflösung in der Regel nur *lokal* geschehen, wie sich an der Kreislinie $K = \{(x, y) \in \mathbb{R}^2 : g(x, y) = 1\}$ mit $g(x, y) = x^2 + y^2$ klarmachen lässt. Global lässt sich die Kreislinie weder als Graph einer Funktion der Variablen $x \in [-1, 1]$ noch als Graph einer Funktion der Variablen $y \in [-1, 1]$ darstellen. Der untere (obere) Halbkreis ohne die Punkte $P(-1, 0)$ und $Q(1, 0)$ lässt sich aber als Graph der Funktion $y(x) = -\sqrt{1 - x^2}$ bzw. $y(x) = +\sqrt{1 - x^2}$ darstellen. Analoges gilt für den linken und den rechten Halbkreis. Hier lässt sich x als Funktion von y darstellen. Man bemerke, dass in den kritischen Punkten P und Q gerade $\partial_y g(x, y) = 0$ gilt. In einer Umgebung dieser Punkte ist keine Auflösung der Form $y = y(x)$ möglich.

Satz 2.5 (Satz über implizite Funktionen)

Gegeben sei eine Funktion $g : \mathbb{R}^2 \rightarrow \mathbb{R}$ und ein Punkt $(\hat{x}, \hat{y}) \in \mathbb{R}^2$ mit $\partial_y g(\hat{x}, \hat{y}) \neq 0$. Ferner sei $c = g(\hat{x}, \hat{y})$. Dann gibt es ein $\varepsilon > 0$ und eine Funktion $h : (-\varepsilon, \varepsilon) \rightarrow \mathbb{R}$ mit den folgenden Eigenschaften: Für alle $(x, y) \in (\hat{x} - \varepsilon, \hat{x} + \varepsilon) \times (\hat{y} - \varepsilon, \hat{y} + \varepsilon)$ gilt:

- i) $g(x, y) = c \Leftrightarrow y = h(x)$
- ii) $h'(x)\partial_y g(x, y) + \partial_x g(x, y) = 0$

Die Niveaulinie $g = c$ lässt sich also in einer Umgebung um den Punkt $(\hat{x}, \hat{y})$ als Graph der der Funktion $x \mapsto g(x, h(x))$ darstellen. Aus dieser Darstellung und der Kettenregel folgt auch sofort ii) in Satz 2.5. Der Beweis der Existenz einer solchen Auflösungsfunktion lässt sich nicht mit elementaren Mitteln führen, sondern nutzt in einer modernen Darstellung den Banach'schen Fixpunktsatz. Deshalb verzichten wir an dieser Stelle auf einen Beweis und verweisen stattdessen auf [30].

Wir wenden jetzt Satz 2.5 an, um Satz 2.4 zu beweisen.

Beweis von Satz 2.4 Im Allgemeinen gilt mit einer beliebigen Testfunktion φ nicht die Nebenbedingung $L(H + s\varphi) = L_0$ für $s \neq 0$. Um dieses Problem zu umgehen, wählen wir irgendeine Testfunktion $\chi \in \mathcal{T}$ mit $\delta L(H, \chi) \neq 0$ und betrachten für eine vorgelegte Testfunktion $\varphi \in \mathcal{T}$ den Ausdruck $H + s\varphi + t\chi$. Der Trick besteht jetzt darin, das t immer passend zum s zu wählen. Das ermöglicht der Satz über die implizite Funktion. Wir definieren

$$\Phi(s, t) := L(H + s\varphi + t\chi). \quad (2.17)$$

Es gilt $\Phi(0, 0) = L_0$ und $\partial_t \Phi(0, 0) = \delta L(H, \chi) \neq 0$. Nach dem Satz über die impliziten Funktionen, Satz 2.5, gibt es ein $\varepsilon > 0$ und eine Funktion $h : (-\varepsilon, \varepsilon) \rightarrow \mathbb{R}$, sodass $h(0) = 0$ sowie $\Phi(s, h(s)) = L_0$ und somit $H + s\varphi + h(s)\chi \in \tilde{\mathcal{A}}$ für alle $-\varepsilon < s < \varepsilon$. Weil H das Energiefunktional E über $\tilde{\mathcal{A}}$ minimiert, folgt

$$\lim_{s \rightarrow 0} \frac{E(H + s\varphi + h(s)\chi) - E(H)}{s} = 0. \quad (2.18)$$

Jetzt müssen wir noch zeigen, dass der Grenzwert in (2.18) gerade die Variation des Funktionals $E - \lambda L$ an der Stelle H in Richtung φ ist, wenn wir λ geeignet (und unabhängig von φ) wählen. Dazu definieren wir

$$\Psi(s, t) = E(H + s\varphi + t\chi).$$

Der Grenzwert in (2.18) ist dann gerade $\frac{d}{ds}(\Psi(s, h(s)))$ ausgewertet an der Stelle $s = 0$. Mit der Kettenregel folgt:

$$\frac{d}{ds}(\Psi(s, h(s))) = \partial_s \Psi(s, h(s)) + \partial_t \Psi(s, h(s)) \cdot h'(s)$$

Wir werten die einzelnen Bausteine einzeln an der Stelle $s = 0$ aus. Es gilt $\partial_s \Psi(0, 0) = \delta E(H, \varphi)$. Wegen $h(0) = 0$ erhalten wir $\partial_t \Psi(0, h(0)) = \delta E(H, \chi)$. Schließlich gilt nach Satz 2.5 ii) $h'(0) = -\frac{\delta L(H, \varphi)}{\delta L(H, \chi)}$. Zusammenfassend mit (2.18) erhalten wir:

$$0 = \delta E(H, \varphi) - \frac{\delta E(H, \chi)}{\delta L(H, \chi)} \cdot \delta L(H, \varphi)$$

Der Lagrange'sche Multiplikator ist somit $\lambda = \frac{\delta E(H, \chi)}{\delta L(H, \chi)}$ unabhängig von φ .

2.2.8 Einige Bemerkungen zu den Extremalprinzipien

In der Physik gehören Extremalprinzipien wohl zu den wichtigsten Modellierungswerkzeugen. Ganze grundlegende Theorien lassen sich aus Extremalprinzipien entwickeln, sogenannte Modellierungen aus *First Principles*. Die Modellbildungsaufgabe besteht dann i.W. darin, die richtige Lagrange-Dichte aufzustellen, der Rest ist mehr oder weniger Kalkül. Die Wahl der richtigen Lagrange-Dichte wird vor allem durch Symmetrieüberlegungen gesteuert: Die Symmetrien, die in der betrachteten physikalischen Situation vorherrschen, müssen auch für die Lagrange-Dichte gelten. Häufig schränkt das den Kandidatenkreis weitgehend ein.

Minimierer eines Funktionals müssen die Euler-Lagrange-Gleichungen lösen. Umgekehrt müssen aber Lösungen der Euler-Lagrange-Gleichungen nicht unbedingt Minimierer des zugehörigen Funktionals sein, es gibt hier ebenfalls so etwas wie Sattelpunktphänomene, wie man sie auch von den Funktionen einer Variablen kennt. In der physikalischen Situation muss dann entschieden werden, ob man auch an solchen Lösungen interessiert ist oder ob es tatsächlich ein Minimierer sein muss. Die letztere Eigenschaft müsste dann ggf. überprüft werden. Auch das kennt man schon aus dem eindimensionalen Fall.

Wir haben in Abschn. 2.2.3 ein Beispiel betrachtet, in dem ein Funktional keinen Minimierer über einer Menge zulässiger Funktionen $\mathcal{A}$ hat. Tatsächlich ist das Funktional in diesem Fall nach unten unbeschränkt, sodass die Abwesenheit eines Minimierers auch nicht verwunderlich ist. Es tritt aber häufig der Fall auf, dass ein beschränktes Funktional über einer Menge $\mathcal{A}$ keinen Minimierer hat. Als eindruckliches Beispiel wollen wir unter den Kurven, die durch drei vorgegebene Punkte verlaufen, diejenige mit der kleinsten Länge suchen. Wir wählen zum Beispiel die Punkte $P_1(0, 1)$, $P_2(1, 0)$ und $P_3(2, 1)$, als zugrunde liegende Menge $\mathcal{A} = \{f : [0, 2] \rightarrow \mathbb{R} : f(0) = 1, f(1) = 0, f(2) = 1\}$ und als Funktional die Bogenlänge $E(f) = \int_0^2 \sqrt{1 + f'^2(x)} dx$. Anschaulich ist es einleuchtend, dass $E(f) \geq 2\sqrt{2}$ für alle Funktionen aus $\mathcal{A}$ gilt, siehe auch Aufg. 2.12. Der offensichtliche Minimierer, nämlich die stückweise lineare Funktion mit Knick an der Stelle $x = 1$, gehört aber nicht zur Menge der zulässigen Funktionen, diese Funktion ist ja an der Stelle $x = 1$ nicht differenzierbar. Auch das Bogenlängenfunktional ist erst einmal für diese Funktion gar nicht definiert.

Dieses einfache Beispiel zeigt bereits, dass es für eine erfolgreiche Anwendung von Extremalprinzipien nötig sein wird, den Begriff der Differenzierbarkeit (tatsächlich auch den der Integrierbarkeit) zu erweitern, um sicherzustellen, dass beschränkte Funktionale immer einen Minimierer haben. Das führt unter anderem zu den Konzepten der schwachen oder distributionellen Ableitung und den sogenannten Sobolev-Räumen von Funktionen.

2.2.9 Reflexionen zum mathematischen Modellieren III: Ad-hoc-Modelle, Modelle nach First Principles und die „unreasonable effectiveness of mathematics in the natural science“

Wir kehren noch einmal zur Kabel-Aufgabe zurück. Offenbar gibt es einen großen Unterschied zwischen den Modellen 2 und 3. Das zweite Modell ist nach Augenschein des beobachteten Phänomens ausgewählt worden: Die einfachste mathematische Kurve, deren Gestalt dem Kabelverlauf ähnelt, ist die Parabel. Darüber hinaus gibt es keinen weiteren Grund, ausgerechnet eine Parabel zu wählen. Wir sprechen von einer Ad-hoc-Modellierung, die keinen weiteren Hintergrund benötigt.

Die Wahl der hyperbolischen Funktion erfolgt aufgrund eines Prinzips, welches in der realen Situation für den nicht vorgebildeten Betrachter in keiner Weise ersichtlich oder nahelegend ist: Modelliere zu der gegebenen Situation aufgrund physikalischer Überlegungen eine Energiefunktion und bestimme mit Hilfe recht ausgeklügelter Mathematik den Minimierer dieser Funktion. Dieser wird die Situation gut darstellen. So etwas nennen wir eine Modellierung nach *First-Principles*.

Ist es nun erstaunlich, dass aus diesem Prinzip nach Durchnudeln durch eine komplizierte mathematische Maschinerie ein mit der Realität so gut übereinstimmendes Ergebnis herauskommt? Oder allgemeiner: Ist es erstaunlich, dass zumindest die Physik so gut mit den Mitteln der Mathematik beschrieben werden kann? Ist die Mathematik nur ein großer Vorrat an Modellen, aus denen man geschickt eines auswählen kann, das ein bestimmtes natürliche Phänomen dann mehr oder weniger gut beschreibt, oder gehorcht die (physikalische) Natur tatsächlich mathematischen Gesetzmäßigkeiten, ist nach mathematischen Gesetzmäßigkeiten aufgebaut? In einem berühmten Aufsatz *The unreasonable effectiveness of mathematics in the natural science* [93] kommt der Physik-Nobelpreisträger Eugene Wigner zu dem Schluss:

“The miracle of the appropriateness of the language of mathematics for the formulation of the laws of physics is a wonderful gift which we neither understand nor deserve. We should be grateful for it and hope that it remain valid in future research and that it will extend, for better or worse, to our pleasure, even though perhaps also to our bafflement, to wide branches of learning.”

Nicht so verwunderlich findet Ortlieb die Übereinstimmung zwischen Vorhersage und empirischer Realität, siehe [69]. Er macht die spezifische Art der Naturerkenntnis dafür verantwortlich:

„Die Mathematik – das wußte bereits Kant – liegt nicht in der Natur, sondern in unserer spezifischen Erkenntnis der Natur. Erst recht gilt das für ein mathematisches Instrumentarium, dem mit dem Gewinn seiner Unabhängigkeit die Denknötwendigkeit verloren gegangen ist und das wir daher – in den Worten von Hertz – nach Aspekten der Zweckmäßigkeit auswählen, dem wir Details ‚nach Willkür hinzutun oder wegnehmen können‘. Daß ‚gewisse Übereinstimmungen

vorhanden sein (müssen) zwischen der Natur und unserem Geiste', wovon auch Hertz spricht, wird in der Physik dadurch gewährleistet, daß die Natur im Experiment an unseren Geist, also an die mathematischen Idealbedingungen angepaßt und die besagte Übereinstimmung damit erst hergestellt wird. Lassen sich dagegen die im Modell unterstellten Idealbedingungen nicht oder nur unzureichend herstellen, so bleiben die zu beobachtenden ‚Naturgesetze‘ letztlich mathematische Fiktionen, wie jeder wissen könnte, der einmal Modelle und Daten ‚gefittet‘ hat. Die Gesetzmäßigkeit steckt allein in der mathematischen Funktion des Modells, während die Abweichungen der Beobachtungsdaten davon durch externe ‚Störungen‘ erklärt werden, die sich der Modellierung entziehen. [...] Unter der Annahme, die Wirklichkeit folge mathematischen Gesetzen, versuchen wir diejenige mathematische Struktur und Gesetzmäßigkeit herauszufinden, die mit kontrollierten Beobachtungen am besten zusammenpasst. Offenbar funktioniert das in vielen Bereichen, nur folgt daraus eben nicht die Richtigkeit der zu Grunde liegenden Annahme. Umgekehrt wird es schlüssig: Durch die Wahl eines bestimmten Instrumentariums – das der exakten Wissenschaften – fokussieren wir und beschränken wir uns auf die Erkenntnis derjenigen Aspekte der Wirklichkeit, die sich mit diesem Instrumentarium erfassen lassen. Und es spricht nichts dafür, daß das schon die ganze Wirklichkeit wäre oder einmal werden könnte.“

Aufgaben

2.1 Der Bremsweg

Zeigen Sie, dass ein Auto, dessen Geschwindigkeit linear mit der Beschleunigungskonstante a zunimmt, vom Beschleunigen aus dem Stand bis zur Geschwindigkeit v die Strecke $s = \frac{v^2}{2a}$ zurücklegt.

2.2 Varianten des Modells „Anhalteweg“

- Variieren Sie das Modell aus Abschn. 2.1.2 so, dass sich die Autofahrerin nicht an ihrer direkten Vorgängerin, sondern an ihrer Vorvorgängerin orientiert. Berechnen Sie die optimale Geschwindigkeit und den zugehörigen Fluss.
- Implementieren Sie das Modell „Anhalteweg“ in einem TK oder DGS, sodass die Bremsbeschleunigung, die minimale Fahrzeuglänge, die durchschnittliche Fahrzeuglänge, die Länge der Schrecksekunde eingegeben werden können und die optimierte Geschwindigkeit und der zugehörige Fluss ausgegeben werden.

2.3 Minimierung von Wegzeiten bei unterschiedlichen Geschwindigkeiten

Horst-Hugo zeltet am Rhein-Herne-Kanal (ein stehendes Gewässer). Um kurz vor acht Uhr bemerkt er, dass sein Getränkevorrat sich dem Ende zuneigt. Auf der gegenüberliegenden Uferseite liegt ein Büdchen, das er so schnell wie möglich erreichen möchte. Da keine

Brücke in der Nähe ist, beschließt er zu schwimmen. Wie die meisten Menschen kann sich Horst-Hugo an Land schneller vorwärts bewegen als im Wasser.



- Modellieren Sie die Situation so allgemein wie möglich unter der Fragestellung, welche Stelle am gegenüberliegenden Ufer Horst-Hugo beim Schwimmen anvisieren muss, um das Büdchen möglichst schnell zu erreichen. Führen Sie eine passende Entdimensionalisierung durch.
- Implementieren Sie Ihre Modellierung in einem DGS oder TK so, dass die Parameterwerte eingegeben werden können und die anzuvisierende Stelle und die benötigte Zeit zum Büdchen ausgegeben werden.
- Erweitern Sie das Modell so, dass auch fließende Gewässer mit umfasst werden. Implementieren Sie auch dieses Modell in einem TK oder DGS.

2.4 Ein Markt mit zwei Anbietern

Als Schulaufgabe ist folgendes Modell zur Simulation der Marktsituation bei zwei Anbietern konzipiert: Für Anbieter 1 sei die Nachfragefunktion

$$f_1(x_1, x_2) = 70 - x_1 + \delta x_2$$

gegeben. Hierbei ist 70 die Grundnachfrage, x_1 der Preis bei Anbieter 1 und x_2 der Preis bei Anbieter 2. Analog gilt für die Nachfrage bei Anbieter 2:

$$f_2(x_1, x_2) = 70 - x_2 + \delta x_1$$

Hieraus ergeben sich die Gewinnfunktionen

$$g_1(x_1, x_2) = f_1(x_1, x_2) \cdot (x_1 - 20) \quad \text{und} \quad g_2(x_1, x_2) = f_2(x_1, x_2) \cdot (x_2 - 20),$$

wobei 20 die Herstellungskosten beider Anbieter darstellt. Lösen Sie zunächst folgende Aufgaben:

- a) Anbieter 1 unterliegt einer Preisbindung und ist verpflichtet, sein Produkt zu einem festen Preis von $\bar{x}_1$ Euro anzubieten. Bei welchem Preis erzielt Anbieter 2 den höchsten Gewinn?
- b) Bei welchen Preisen $(\bar{x}_1, \bar{x}_2)$ erzielen beide den maximalen Gewinn, wenn der jeweils andere Preis als fix betrachtet wird? (Man spricht dann von einem *Nash-Gleichgewicht*: Beide würden ihren Gewinn verringern, wenn sie ihren Preis aus diesem Gleichgewichtspunkt verändern würden.)
- c) Die beiden Anbieter treffen illegale Kartellabsprachen. Welche Möglichkeiten gibt es, den Gesamtgewinn zu maximieren?
- d) Welche Werte für δ sind in diesem Modell sinnvoll?

Die oben verwendeten Formeln ergeben, betrachtet man die Dimensionen, keinen Sinn. Wir ziehen daher die folgende Form vor:

$$f_1(x_1, x_2) = n - \alpha x_1 + \beta x_2 \quad \text{und} \quad g_1(x_1, x_2) = f_1(x_1, x_2)(x_1 - k)$$

mit beliebiger Grundnachfrage n und beliebigen Herstellungskosten k . Für Anbieter 2 gilt dies analog.

$[f] = A$ ist eine Anzahl und $[x_1] = [x_2] = G/A$ eine Geldmenge pro Anzahl.

- e) Welche Dimension haben die Parameter n, α, β, k bzw. der Gewinn g ?
- f) Bringen Sie die Funktion f_1 durch geschickte Entdimensionalisierung von f_1, x_1 und x_2 auf die Form

$$\varphi_1(y_1, y_2) = 1 - y_1 + \gamma y_2.$$

- g) Gibt es eine Möglichkeit, die Gewinnfunktion g_1 auf die Form

$$\gamma_1(y_1, y_2) = \varphi_1(y_1, y_2)(y_1 - \varepsilon)$$

zu bringen?

- h) Bearbeiten Sie die Aufgaben a), b) und c) in verallgemeinerter Fassung mit den in f) und g) gefundenen Formeln.

2.5 Die Taylor-Formel im Mehrdimensionalen

In (2.7) wurde die die Taylor-Formel für eine Funktion mehrerer Variablen benutzt. Diese soll hier aus der Taylor-Formel für Funktionen einer Variablen hergeleitet werden soll.

- a) Es sei $f : \mathbb{R}^n \rightarrow \mathbb{R}$ zweimal stetig differenzierbar, $x, h \in \mathbb{R}^n$. Zeigen Sie: Es gibt $\theta \in (0, 1)$, sodass

$$f(x + h) = f(x) + \sum_{i=1}^n \partial_i f(x) h_i + \frac{1}{2} \sum_{i,j=1}^n \partial_i \partial_j f(x + \theta h). \quad (2.19)$$

- b) Verallgemeinern Sie diese Formeln für Approximationen beliebiger Ordnungen und beweisen Sie diese.

2.6 Beweis des Fundamentallemmas der Variationsrechnung (stetige Version)

Beweisen Sie die folgende Version des *Fundamentallemmas der Variationsrechnung* durch einen Widerspruchsbeweis:

Sei $[a, b] \subset \mathbb{R}$ ein Intervall und $f : [a, b] \rightarrow \mathbb{R}$ stetig. Ferner gelte für alle stetigen Funktionen $\varphi : [a, b] \rightarrow \mathbb{R}$ mit $\varphi(a) = \varphi(b) = 0$, dass

$$\int_a^b f(x)\varphi(x) dx = 0,$$

dann ist f die Nullfunktion auf $[a, b]$.

2.7 Rechnungen zum Energiefunktional der hängenden Kette

- Rechnen Sie nach, dass ein Minimierer H des Energiefunktionals zur Lagrange-Dichte (2.11) die zugehörige Euler-Lagrange-Gleichung (2.12) löst. Im Wesentlichen ist das eine Übung der Anwendung der Kettenregel.
- Rechnen Sie nach, dass die Gl. (2.12) nicht durch parabelförmige Höhenverläufe gelöst wird.
- Rechnen Sie nach, dass ein Minimierer H von $E - ghH_0L$ die Gleichung (2.16) löst.

2.8 Höhenlinien und Gradienten

In Abschn. 2.2.5 wird behauptet, dass Höhenlinien und Gradienten einer Funktion stets senkrecht zueinander stehen. Die Aussage soll in dieser Aufgabe mathematisch präzisiert und bewiesen werden. Zunächst einmal muss geklärt werden, was in diesem Zusammenhang überhaupt mit einem Winkel gemeint ist. Gemeinhin sind lediglich Winkel zwischen Vektoren über das Skalarprodukt definiert.

- Sei $\gamma : (-1, 1) \rightarrow \mathbb{R}^2$ eine beliebige Kurve, die in einer Höhenlinie von $h : \mathbb{R}^2 \rightarrow \mathbb{R}$ verläuft. Zeigen Sie: Für alle $t \in (-1, 1)$ gilt

$$\dot{\gamma}(t) \cdot \nabla h(\gamma(t)) = 0.$$

- Inwiefern liefert dieser Sachverhalt eine mathematische Präzisierung der Aussage „Höhenlinien und Gradienten sind senkrecht zueinander“?

2.9 Beweis der Methode der Lagrange'schen Multiplikatoren im zweidimensionalen Fall

Beweisen Sie Satz 2.3, indem Sie den Beweis aus Abschn. 2.2.7 geeignet modifizieren.

2.10 Der Kosinushyperbolicus ist kein kritischer Punkt des Bogenlängenfunktional

Die Funktion $H(x) = k \cosh\left(\frac{x-\bar{x}}{k}\right) + H_0$ ist kein kritischer Punkt des Bogenlängenfunktional L !

- Machen Sie sich das anschaulich klar.
- Zeigen Sie das rigoros, indem Sie eine konkrete Testfunktion χ angeben und nachrechnen, dass $\delta L(H, \chi) \neq 0$ gilt.

Vorschlag: Wir gehen zunächst einmal von $H_1 = H_2$ aus. Betrachten Sie $\chi := H_1 - H$. Modifizieren Sie geeignet, um auf die Voraussetzung $H_1 = H_2$ zu verzichten.

2.11 Wie hängt ein elastisches Kabel?

In den betrachteten Modellen wurde vorausgesetzt, dass das Kabel ein unveränderliche Länge habe. Tatsächlich sind aber Kabel in der Realität elastisch. Unter Zugspannung verlängern sie sich ein bisschen, unter Druck kann man sie etwas zusammenpressen. Das Verlängern oder Verkürzen erfordert Energie. Aus der Schule ist das Hooke'sche Gesetz für elastische Federn bekannt. Demnach erfordert die Veränderung der Länge einer Feder eine Kraft F , die proportional zum Ausmaß der Verlängerung ist, $F = D(L - L_0)$ mit der „Ruhelänge“ L_0 . Die Proportionalitätskonstante D wird üblicherweise als Federhärte bezeichnet. Eine ausgelenkte Feder speichert die Energie $\frac{1}{2}D(L - L_0)^2$ (das ist das „Arbeitsintegral“ $\int_{L_0}^L D(l - L_0) dl$).

- Stellen Sie ein Energiefunktional $\tilde{E}$ für den Höhenverlauf eines elastischen Kabels mit Konstante D auf, indem Sie diesen Energieanteil in Ihrem Energiefunktional berücksichtigen und dafür auf die Nebenbedingung der konstanten Kabellänge verzichten.
- Versuchen Sie die Ableitung der Funktion $s \mapsto \tilde{E}(H + s\varphi)$ zu berechnen. Welche Probleme treten auf?

2.12 Die Kurve zwischen zwei Punkten mit minimaler Bogenlänge

Gegeben seien die Punkte $(0, 0)$ und $(\bar{x}, \bar{y})$ mit $\bar{x} > 0$. Bestimmen Sie mit Hilfe der Variationsrechnung diejenige Funktion

$$f : [0, \bar{x}] \longrightarrow \mathbb{R} \quad \text{mit} \quad f(0) = 0 \quad \text{und} \quad f(\bar{x}) = \bar{y},$$

welche die Bogenlänge minimiert. Denken Sie vorher eine Sekunde darüber nach, was wohl die Lösung sein muss.

2.13 Das Brachistochronenproblem

Ein klassisches Problem aus den Anfängen der Variationsrechnung besteht darin, zu zwei Punkten A und B , wobei A „höher“ liegt als B , diejenige Bahn zu finden, auf der sich eine Kugel nur unter dem Einfluss der Schwerkraft am schnellsten von A nach B bewegen würde.

„Nur unter dem Einfluss der Schwerkraft“ bedeutet hier, dass die Kugel nicht durch Reibung an der Bahn Energie verliert und langsamer wird. Eine experimentelle Realisierung findet sich z. B. in der Elementa in Mannheim, siehe Abb. 2.4. Die Reibungsfreiheit kann natürlich nur näherungsweise realisiert werden, in dem Sinne, dass die auftretenden Reibungskräfte im Vergleich zur Schwerkraft sehr klein sind.

In dieser Aufgabe sollen Sie das zu minimierende Zeitfunktional aufstellen: Zu einer gegebenen Bahn $y : [x_A, x_B] \rightarrow \mathbb{R}$ mit $y(x_A) = y_A$ und $y(x_B) = y_B$ muss die Zeit T berechnet werden, die eine Kugel benötigt, um die Bahn rein unter dem Einfluss der Schwerkraft zurückzulegen. Dafür müssen Sie das funktionale Geflecht der Größen x , y der Zeit t und der Geschwindigkeit v geschickt organisieren. Gegeben ist die Höhe $y = y(x)$. Bekannt sind die folgenden Beziehungen:

- 1) Für die Ortskoordinaten x , y und die Geschwindigkeit v als Funktionen der Zeit, $x(t)$, $y(t)$ und $v(t)$, gilt $\dot{x}^2(t) + \dot{y}^2(t) = v^2(t)$. Dabei steht (wie im gesamten Buch) der Punkt für Ableitungen nach der Zeit.
- 2) Die Bewegung erfolgt reibungsfrei. Damit bleibt die Summe aus potentieller und kinetischer Energie während der gesamten Bewegung erhalten: $\frac{1}{2}mv^2 + mgy = \text{const} = mgy_A$.
- 3) Für die benötigte Zeit T gilt $T = \int_0^T dt = \int_{x_A}^{x_B} t'(x) dx$.

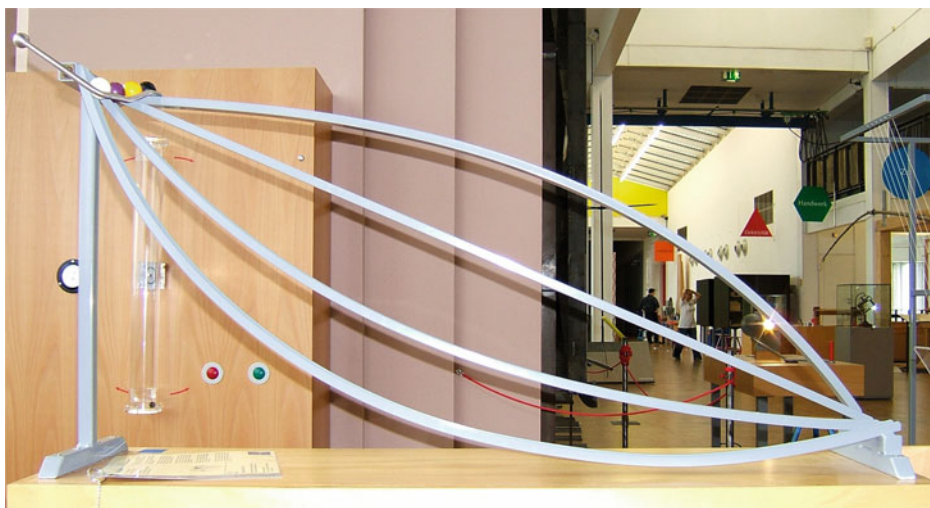


Abb. 2.4 Experimenteller Aufbau zur Demonstration der Laufzeiten unterschiedlicher Bahnkurven im Landesmuseum für Technik und Arbeit, Mannheim. (Foto: <https://commons.wikimedia.org/wiki/File:Brachistochronerutschbahn.jpg#file>)

Das Ziel ist also, das Integral über die Integrationsvariable der Zeit t durch ein Integral über die Integrationsvariable x in Termen von $y(x)$ und $y'(x)$ auszudrücken.

Berechnen Sie in einem zweiten Schritt die zugehörigen Euler-Lagrange-Gleichungen.

2.14 Bewegungen als Minimierer von Wirkungen

Die Physiker haben noch weitere Prinzipien, z. B. das *Prinzip von der minimalen Wirkung* (eigentlich müsste es Prinzip von der *stationären* Wirkung heißen). Es sagt aus, dass für die physikalisch realisierte Bewegung eines Körpers, die durch die Funktion $t \mapsto x(t)$ mathematisch dargestellt wird ($x(t)$ gibt den Ort des Körpers zur Zeit t an), das folgende Prinzip gilt: Das sogenannte Wirkungsfunktional

$$S(x) := \int_{t_1}^{t_2} \mathcal{L}(x(t), \dot{x}(t)) dt$$

ist für die realisierte Bewegung $x(t)$ minimal (eigentlich stationär). Dabei ist $\mathcal{L}$ die Differenz aus der kinetischen Energie T und der potentiellen Energie V des Körpers, $\dot{x}$ bezeichnet die Geschwindigkeit des Körpers, also die Zeitableitung.

In dieser Aufgabe sollen die aus der Schule bekannten Gleichungen für den senkrechten Wurf aus dem Prinzip der minimalen Wirkung hergeleitet werden.

Für die vertikale Bewegung im homogenen Schwerfeld der Erde ($x(t)$ gibt also die Höhe über dem Boden zur Zeit t an) gilt

$$T(x, \dot{x}) = \frac{1}{2} m \dot{x}^2 \quad \text{und} \quad V(x, \dot{x}) = mgx,$$

dabei ist m die Masse des Körpers und g die Erdbeschleunigung.

Bestimmen Sie mit Hilfe der Variationsrechnung die Bewegungsgleichung für einen Körper, der zur Zeit $t = 0$ mit der Geschwindigkeit v_0 senkrecht vom Boden in die Luft geschossen wird.



Zinsen, Strahlung und Kaninchen – Modellieren mit linearen Differenzengleichungen

3

Inhaltsverzeichnis

3.1 Guthaben und Zinsen	65
3.2 Der Radongehalt in der Luft eines unbelüfteten Kellerraums	66
3.3 Zusammenfassung, Systematisierung und graphische Analyse	68
3.4 Summen der n -ten Potenzen und Kaninchen – das Lösungsschema für lineare Gleichungen	70
Aufgaben	76

Das Modellierungswerkzeug in diesem Kapitel sind lineare Differenzengleichungen erster und zweiter Ordnung. Aufgrund der elementaren Mathematik sind die aufgeführten Beispiele zum großen Teil auch im Schulunterricht einsetzbar. Die Systematisierung in Abschn. 3.3 dient der Vorbereitung auf das folgende Kap. 4, in dem mit nichtlinearen Differenzengleichungen modelliert wird. Wir führen zudem in diesem mathematischen Kontext das „lineare Lösungsschema“ ein (Abschn. 3.4), das uns im Kontext der linearen Differentialgleichungen in Kap. 5 bis 8 immer wieder begegnen wird.

3.1 Guthaben und Zinsen

3.1.1 Guthaben

In diesem Kapitel betrachten wir ein elementares *normatives* Modell, nämlich Verzinsungsvorgänge. Im Gegensatz zu den meisten Modellierungen dieses Buchs, die sich mit Beschreibungen von Naturvorgängen befassen, also *deskriptiver* Natur sind, wird hier durch das mathematische Modell der Gegenstand erst erschaffen.

Das Guthaben auf dem Festgeldkonto einer Bank wird einmal im Jahr mit dem Zinssatz p verzinst. Wie groß ist das Guthaben nach n Jahren?

Modellierung

Wir bezeichnen und rechnen:

a_n : das Guthaben zu Beginn des n -ten Jahres

a_0 : das zum Startzeitpunkt eingezahlte Geld

$$a_1 = a_0 + pa_0 = (1 + p)a_0$$

$$a_2 = (1 + p)a_1$$

$$a_{n+1} = (1 + p)a_n \quad \text{für alle } n \in \mathbb{N}_0$$

Durch vollständige Induktion nach n zeigt man leicht $a_n = (1 + p)^n a_0$.

3.1.2 Festgeldkonto mit jährlicher Einlage

Zusätzlich zur Verzinsung wird jedes Jahr eine Einlage geleistet oder eine Auszahlung entnommen, die wir mit b bezeichnen. b kann also positive oder negative Werte annehmen. Wie hoch ist das Guthaben nach n Jahren?

Modellierung

Wir haben $a_{n+1} = (1 + p)a_n + b$ und versuchen aus dieser *Rekursionsformel* – man muss nacheinander jedes Folgeglied aus dem vorangegangenen berechnen – eine *explizite* Formel zu finden, mit deren Hilfe man z. B. das tausendste Folgeglied direkt berechnen kann, ohne die 999 Folgeglieder davor ausrechnen zu müssen. Durch explizites Berechnen der ersten paar Folgeglieder kommen wir auf die Vermutung

$$a_n = (1 + p)^n a_0 + b \cdot \sum_{j=0}^{n-1} (1 + p)^j,$$

die sich durch Induktion nach n beweisen lässt. Mit Hilfe der geometrischen Summenformel können wir diesen Term umformen zu

$$a_n = (1 + p)^n a_0 + b \cdot \frac{(1 + p)^n - 1}{p}. \quad (3.1)$$

Eine Anwendungen dieses Modells auf Immobilienkredite findet sich in Aufg. 3.1.

3.2 Der Radongehalt in der Luft eines unbelüfteten Kellerraums

Die oberen Gesteinsschichten der Erde enthalten in unterschiedlichen Konzentrationen das radioaktive Nuklid Radium 226. Dieses Nuklid zerfällt mit einer Halbwertszeit von $\tau_{\text{Ra226}} = 1600$ a in das radioaktive Nuklid Radon 222. Das bedeutet, dass von einer großen Anzahl von

Radium-Nukliden nach 1600 Jahren die Hälfte in Radon-Nuklide zerfallen ist. Das Radon-Nuklid liegt gasförmig vor und kann durch Spalten und Poren aufsteigen und sich z. B. in unbelüfteten Kellerräumen ansammeln. Es ist selber auch wieder radioaktiv und zerfällt mit einer Halbwertszeit von $\tau_{\text{Rn222}} = 3,825 \text{ d}$ in Polonium 218. Es stellt sich die Frage, ob sich Radon in gesundheitsschädlichen Mengen in solchen Kellerräumen ansammeln kann und ob sich der Radongehalt durch regelmäßiges Lüften wirksam reduzieren lässt.

Wir nehmen an, dass der Radongehalt im Kellerraum einmal pro Tag bestimmt wird.

3.2.1 Modellierung

Es sei a_n die Menge an Radon 222 am n -ten Tag. Radium hat im Gegensatz zu Radon eine sehr lange Halbwertszeit. Deshalb nehmen wir eine zeitlich konstante tägliche Zufuhr b von Radon-Nukliden durch den Zerfall von Radium-Nukliden an. Wir modellieren:

$$a_{n+1} = r \cdot a_n + b,$$

dabei gibt r den Anteil an Radon-Nukliden an, die an einem Tag nicht zerfallen sind. Wir müssen nun diesen *täglichen Zerfallsfaktor* r bestimmen. Ohne die Zufuhr von neuen Nukliden reduziert sich die Radon-222-Menge innerhalb von 3,825 d auf die Hälfte. Lassen wir diesen Zeitraum 1000-mal verstreichen, erhalten wir (ohne neue Zufuhr)

$$a_{n+3825} = \left(\frac{1}{2}\right)^{1000} a_n.$$

Der Zerfallsfaktor für 3825 Tage beträgt also $(1/2)^{1000}$ und der Zerfallsfaktor für einen Tag damit

$$r = \left(\frac{1}{2}\right)^{\frac{1000}{3825}} = 0.834.$$

Mit Hilfe der Formel (3.1) aus Abschn. 3.1.2 erhalten wir (mit $r = 1 + p$)

$$a_n = r^n \left(a_0 - \frac{b}{1-r} \right) + \frac{b}{1-r}.$$

Uns interessiert besonders, wie sich der Radongehalt entwickelt, wenn wir den Keller lange Zeit nicht lüften. Wir sind also am sogenannten *Langzeitverhalten* interessiert, das wir durch den Limes $\lim_{n \rightarrow \infty} a_n$ modellieren. Wegen $0 < r < 1$ gilt $\lim_{n \rightarrow \infty} r^n = 0$ und damit unabhängig vom Startwert a_0

$$\lim_{n \rightarrow \infty} a_n = \frac{b}{1-r} = 6,0334b.$$

3.2.2 Interpretation

Auf lange Sicht ist der Radongehalt in der Luft also circa sechsmal so groß wie die Menge an Radon, die täglich durch den Zerfall von Radium-Nukliden hinzukommt. Nach diesem Modell kann somit das tägliche Lüften des Kellerraums den Radon-Gehalt auf ungefähr eine Sechstel reduzieren.

Bei den beiden angeführten Modellierungen haben wir explizite Lösungsformeln für die Differenzengleichungen gefunden und diese dann benutzt. Das wird bei den später folgenden Modellierungen, die auf nichtlineare Differenzengleichungen führen werden, nicht mehr möglich sein. Um auch aus solchen Gleichungen auf das Verhalten der Lösungsfolgen der Gleichungen schließen zu können, werden wir zunächst ein paar Begriffe einführen und mathematisches Rüstzeug bereitstellen.

3.3 Zusammenfassung, Systematisierung und graphische Analyse

3.3.1 Definition, eindeutige Lösbarkeit und Darstellungsformel

1. Für eine Funktion $f : \mathbb{N}_0 \times \mathbb{R} \rightarrow \mathbb{R}$ heißt die Gleichung $a_{n+1} = f(n, a_n)$ eine Differenzengleichung (DZG). Falls die Funktion f nicht explizit von n abhängt, sprechen wir von einer *autonomen* DZG.
2. Durch Vorgabe eines Anfangswerts $\bar{a}_0$ ist jede Differenzengleichung eindeutig lösbar, d.h., es gibt genau eine Folge $(a_n)_{n \in \mathbb{N}_0}$ mit $a_0 = \bar{a}_0$ und $a_{n+1} = f(n, a_n)$ für alle $n \in \mathbb{N}_0$.
3. Mit $f(n, a) = f(a) = r \cdot a + b$ und $r, b \in \mathbb{R}$ heißt die DZG $a_{n+1} = f(a_n)$ eine *lineare DZG mit konstanten Koeffizienten*. Sie lässt sich explizit lösen mit der Lösungsformel

$$a_n = \left(a_0 + \frac{b}{1-r} \right) r^n + \frac{b}{1-r} \quad \text{für } r \neq 1 \quad (3.2)$$

$$a_n = a_0 + nb \quad \text{für } r = 1. \quad (3.3)$$

Der Beweis der eindeutigen Lösbarkeit erfolgt durch Induktion nach n , die Formel (3.2) aus (3.1) mit $r = 1 + p$ und (3.3) durch scharfes Hingucken (und Induktion).

3.3.2 Das Cobwebbing

Beim Cobwebbing handelt es sich um eine graphische Veranschaulichung der Lösung einer autonomen DZG $a_{n+1} = f(a_n)$. Dazu zeichnet man den Graphen von f und den Graphen der identischen Abbildung $\text{id} : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x$ in ein Koordinatensystem und verfolgt in diesem Diagramm den Verlauf von einzelnen Lösungsfolgen. Wir schauen uns das zunächst

für die lineare DZG mit $f(a) = r \cdot a + b$ an. Dabei unterscheiden wir die Fälle $1 < r$, $0 < r < 1$, $-1 < r < 0$ und $r < -1$. Die Grenzfälle -1 , 0 und 1 werden in Aufg. 3.2 behandelt. Die Stelle a^* ist durch die Fixpunktgleichung $a^* = r \cdot a^* + b$ bestimmt, es gilt also $a^* = b/(1 - r)$. Die Cobwebs in Abb. 3.1 verdeutlichen die folgenden Sachverhalte, die sich auch leicht anhand der Formel (3.2) verifizieren lassen:

- Für $1 < r$ haben wir
 - $\lim_{n \rightarrow \infty} a_n = +\infty$ und die Folge ist streng monoton wachsend, falls für den Startwert $a_0 > a^*$ gilt.
 - $\lim_{n \rightarrow \infty} a_n = -\infty$ und die Folge ist streng monoton fallend, falls für den Startwert $a_0 < a^*$ gilt.
- Für $0 < r < 1$ konvergieren alle Lösungsfolgen unabhängig vom Anfangsdatum gegen a^* , und zwar
 - streng monoton wachsend, falls $a_0 < a^*$ gilt,
 - streng monoton fallend, falls $a_0 > a^*$ gilt.

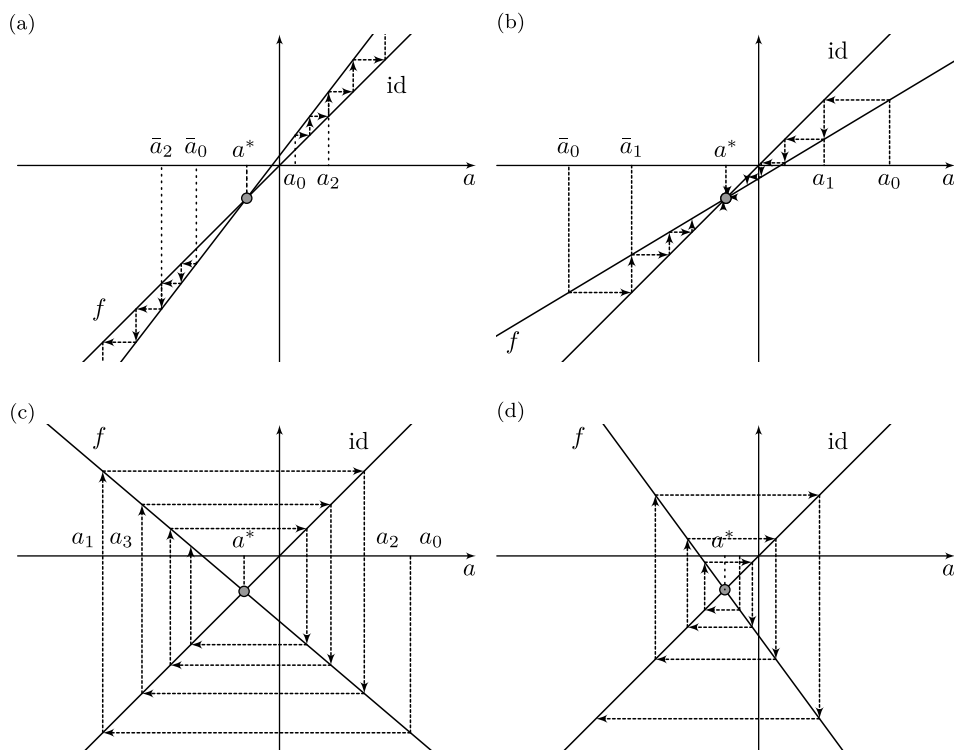


Abb. 3.1 Cobweb der linearen DZG für $1 < r$ (a) und $0 < r < 1$ (b), $-1 < r < 0$ (c) und $r < -1$ (d)

- Für $-1 < r < 0$ konvergieren alle Lösungsfolgen unabhängig vom Anfangsdatum gegen a^* , und zwar – mit Ausnahme der konstanten Folge (a^*) – oszillierend, d. h., es ist abwechselnd ein Folgenwert größer und ein Folgenwert kleiner als a^* .
- Für $r < -1$ und einen Startwert $a_0 \neq a^*$ gilt $\lim_{n \rightarrow \infty} |a_n| = \infty$ und die Folge oszilliert um a^* .

Diese Beobachtungen veranlassen einige Definitionen:

Definition 3.1

Gegeben sei eine DZG $a_{n+1} = f(a_n)$.

- Die Lösungen der Gleichungen $a^* = f(a^*)$ heißen *stationäre Punkte* der Differenzengleichung.
- Im Fall $f(a) = r \cdot a + b$ heißt der (in den aufgeführten Fällen eindeutige) stationäre Punkt a^*
 - monoton instabil, falls $r > 1$,
 - monoton asymptotisch stabil, falls $0 < r < 1$,
 - oszillierend asymptotisch stabil, falls $-1 < r < 0$,
 - oszillierend instabil, falls $r < -1$.



3.4 Summen der n -ten Potenzen und Kaninchen – das Lösungsschema für lineare Gleichungen

Wir haben die Lösungsformel (3.1) durch scharfes Hingucken erraten und anschließend durch Induktion nach n bewiesen. In diesem Kapitel soll ein anderer Lösungsweg beschrieben werden, der von der linearen Struktur der Differenzengleichung Gebrauch macht und dessen Struktur uns in vielen anderen Situationen, die lineare Gleichungen hervorbringen, wiederbegegnen wird. Sie kennen ihn bereits von der Lösung linearer Gleichungssysteme her.

3.4.1 Das lineare Lösungsschema

Dazu betrachten wir allgemein eine *inhomogene lineare Differenzengleichung mit konstanten Koeffizienten*

$$x_{n+1} = rx_n + b \quad \text{für alle } n \in \mathbb{N}_0, \quad (3.4)$$

mit festen reellen Koeffizienten $r, b \in \mathbb{R}$. Eine reelle Folge $(x_n)_{n \in \mathbb{N}_0}$ ist genau dann eine Lösung von (3.4), wenn diese Gleichung für alle Folgeglieder dieser Folge gilt. Wir nehmen

jetzt an, wir hätten zwei Lösungsfolgen (y_n) und (z_n) von (3.4). Durch einfaches Nachrechnen sehen wir, dass die Differenzenfolge $(w_n) = (y_n - z_n)$ dann die *homogene lineare Differenzengleichung*

$$x_{n+1} = r x_n \quad \text{für alle } n \in \mathbb{N}_0 \quad (3.5)$$

löst. Haben wir dagegen zwei Lösungsfolgen (w_n) und (v_n) der homogenen Gl. (3.5), so ist auch jede Linearkombination $\alpha(w_n) + \beta(v_n)$ mit $\alpha, \beta \in \mathbb{R}$ eine Lösung von (3.5). Dabei definieren wir die Multiplikation mit den Skalaren α und β komponentenweise. Die Lösungen der homogenen Gl. (3.5) bilden also einen Untervektorraum im $\mathbb{R}$ -Vektorraum der Folgen mit reellen Komponenten. Da ganz allgemein jede Lösung einer Differenzengleichung $x_{n+1} = f(x_n)$ durch die Angabe ihres Anfangswerts x_0 eindeutig bestimmt ist, ist der Lösungsraum von (3.5) eindimensional und wird z.B. von der Lösung zum Anfangswert $x_0 = 1$ aufgespannt. Wir müssen also nur eine Lösungsfolge der homogenen Gleichung finden (die nicht die Nullfolge ist), alle anderen Lösungsfolgen von (3.5) sind dann lineare Vielfache. Wir tun so, als würden wir die Lösung noch nicht kennen, und versuchen es mit dem Ansatz $x_n = x^n$ mit einem geeignet zu bestimmenden festen $x \in \mathbb{R}$. Dieser Ansatz funktioniert genau dann, wenn die folgende Gleichung gilt:

$$x_{n+1} = x^{n+1} \stackrel{!}{=} r \cdot x_n = r \cdot x^n \quad \text{für alle } n \in \mathbb{N}_0.$$

Auflösen nach x liefert $x = 0$ oder $x = r$. Damit haben wir den ersten Schritt gemacht, nämlich alle Lösungen der homogenen Gleichung zu berechnen:

$$x_n^{\text{hom}} = c r^n \quad \text{mit } c \in \mathbb{R}.$$

In einem zweiten Schritt versuchen wir irgendeine Lösung der inhomogenen Gl. (3.4) zu finden. Auch hier helfen häufig geschickt gewählte Ansätze. Wir bemerken zunächst, dass wir die inhomogene Gl. (3.4) allgemeiner formulieren könnten, nämlich in der Form

$$x_{n+1} = r x_n + b_n \quad \text{für alle } n \in \mathbb{N}_0, \quad (3.6)$$

mit einer gegebenen Folge (b_n) . Im betrachteten Fall ist die Folge (b_n) einfach eine konstante Folge. Die Grundidee, die häufig zum Ziel führt, ist, die partikuläre Lösung *ähnlich* zur vorliegenden *Inhomogenität* anzusetzen. Im vorliegenden Fall ist die Inhomogenität eine konstante Folge, also versuchen wir es bei der partikulären Lösung auch mit einer konstanten Folge, machen also den Ansatz $x_n^{\text{part}} = k$ für alle $n \in \mathbb{N}$ mit einem geeignet zu wählenden $k \in \mathbb{R}$. Wir setzen diesen Ansatz in die Gl. (3.4) ein und sehen, dass er genau dann funktioniert, wenn die folgende Gleichung gilt:

$$k = x_{n+1}^{\text{part}} \stackrel{!}{=} r \cdot x_n^{\text{part}} + b = r \cdot k + b \quad \text{für alle } n \in \mathbb{N},$$

also genau dann, wenn $k = \frac{b}{1-r}$. Wir bemerken, dass dieser Ansatz im Fall $r = 1$ nicht funktioniert, und können das darauf zurückführen, dass die konstante Folge dann bereits

eine Lösung der homogenen Gleichung ist, also keine Lösung der inhomogenen Gleichung sein kann, wenn $b \neq 0$ ist.

Mit der speziellen partikulären Lösung und den Lösungen der homogenen Gleichung haben wir nun für den Fall $r \neq 1$ alle Lösungsfolgen der inhomogenen Gl. (3.6) gefunden:

$$x_n^{\text{all}} = x_n^{\text{hom}} + x_n^{\text{part}} = c \cdot r^n + \frac{b}{1-r} \quad \text{mit } c \in \mathbb{R}.$$

Geben wir nun einen bestimmten Anfangswert x_0 vor, können wir c geeignet wählen:

$$x_0 = cr^0 + \frac{b}{1-r}, \quad \text{also } c = x_0 - \frac{b}{1-r}.$$

Insgesamt erhalten wir also die Lösung

$$x_n = \left(x_0 - \frac{b}{1-r} \right) r^n + \frac{b}{1-r}.$$

Diese Lösung hatten wir freilich auch schon in (3.1) mit $r = 1 + p$ gefunden. Jetzt haben wir aber ein Schema entwickelt, das mutatis mutandis in vielen weiteren Situationen einsetzbar ist.

Wir fassen das *lineare Lösungsschema* zusammen:

Das lineare Lösungsschema

Gegeben sei die affin-lineare Gleichung

$$x_{n+1} = rx_n + b_n \quad \text{für alle } n \in \mathbb{N}_0$$

mit einer vorgegebenen Zahl $r \in \mathbb{R}$ und einer vorgegebenen Folge (b_n) .

I Berechne *alle* Lösungen der homogenen Gleichung

$$x_{n+1}^{\text{hom}} = rx_n^{\text{hom}} \quad \text{für alle } n \in \mathbb{N}_0.$$

II Finde *irgendeine* Lösung der inhomogenen Gleichung

$$x_{n+1}^{\text{part}} = rx_n^{\text{part}} + b_n \quad \text{für alle } n \in \mathbb{N}_0.$$

Häufig helfen hier Ähnlichkeitsansätze.

III Die allgemeine Lösung setzt sich aus den homogenen und einer partikulären Lösung zusammen:

$$(x_n^{\text{all}}) = (x_n^{\text{hom}}) + (x_n^{\text{part}})$$

IV Falls ein Anfangswert vorgegeben ist, spezifiziere für diesen Anfangswert.

Wir betrachten zwei weitere Anwendungen des linearen Lösungsschemas: Zum einen sollen Formeln für die Summe von Potenzen hergeleitet werden, wie sie bei der elementaren Integration von Potenzfunktionen benötigt werden, zum anderen betrachten wir Differenzgleichungen zweiter Ordnung mit dem prominenten Beispiel der Fibonacci-Folge.

3.4.2 Eine Anwendung: die Summe der ersten n -Kubikzahlen

Im Rahmen der elementaren Integration von Quadratfunktionen als Grenzwert von Riemann-Summen wird eine Formel für die Summe der ersten n -Quadratzahlen benötigt, die in der Schule häufig vom Himmel fällt und eventuell per Induktion bestätigt wird. Entsprechend benötigt man für die Integration der kubischen Funktionen die Summe der ersten n -Kubikzahlen und allgemein für die Integration der k -ten Potenz eine Summenfolge für die k -ten Potenzen. Für die niedrigen Potenzen gibt es einige anschaulich geometrische Herleitungen. Allgemeine Potenzsummen lassen sich algebraisch rekursiv mit Hilfe des binomischen Lehrsatzes bestimmen, siehe Aufg. 3.3. Dieses Verfahren hat den Nachteil, dass es rekursiv ist. Zur Bestimmung der Summenformel der 10-ten Potenzen werden Summenformeln für alle niedrigeren Potenzen benötigt.

Im Gegensatz dazu lässt sich mit dem linearen Lösungsschema für Differenzgleichungen die Summenformel für eine beliebige Potenz direkt ansteuern. Das soll am Beispiel der Summe über die ersten n Kuben vorgeführt werden.

Gesucht ist eine Formel für die Summe

$$s_n = \sum_{k=1}^n k^3.$$

Die Folge (s_n) löst die Differenzgleichung

$$s_{n+1} = s_n + (n+1)^3,$$

eine inhomogene lineare Differenzgleichung mit der Inhomogenität $(b_n) = (n+1)^3$, einem Polynom dritter Ordnung in n . Wir suchen die Lösung dieser Differenzgleichung mit dem Anfangswert $s_0 = 0$. Wir versuchen eine Lösung nach dem linearen Lösungsschema zu finden.

- I Die homogene Gleichung $x_{n+1} = x_n$ hat als Lösungen gerade die konstanten Folgen $x_n^{\text{hom}} = c$ für alle $n \in \mathbb{N}_0$ mit einem $c \in \mathbb{R}$.
- II Nach dem Ähnlichkeitsprinzip bietet sich für die partikuläre Lösung (x_n^{part}) ein Polynom dritter Ordnung an: $x_n^{\text{part}} = a_0 + a_1 n + a_2 n^2 + a_3 n^3$. Dieser Ansatz (die interessierte Leserin probiere es aus) funktioniert jedoch nicht. Das liegt daran, dass der konstante Anteil a_0 der Lösungsfolge keinen Beitrag zur Lösung der inhomogenen Gleichung liefern kann, weil die konstanten Folgen gerade die homogene Gleichung

lösen. Der Ausweg besteht darin, ein Polynom vierter Ordnung anzusetzen. Auf den konstanten Term können wir verzichten, weil er wie gesagt keinen Beitrag zur Lösung des inhomogenen Problems liefern kann. Der Ansatz lautet also

$$x_n^{\text{part}} = a_1 n + a_2 n^2 + a_3 n^3 + a_4 n^4.$$

Wir setzen den Ansatz in die Differenzengleichung ein:

$$\begin{aligned} a_1(n+1) + a_2(n+1)^2 + a_3(n+1)^3 + a_4(n+1)^4 \\ = a_1 n + a_2 n^2 + a_3 n^3 + a_4 n^4 + (n+1)^3 \end{aligned}$$

für alle $n \in \mathbb{N}_0$. Wir haben auf der linken und rechten Seite jeweils ein Polynom vierter Ordnung in n stehen, das an abzählbar vielen Stellen, nämlich gerade den natürlichen Zahlen, wertgleich sein soll. Nach dem Fundamentalsatz der Algebra müssen dann die beiden Polynome identisch sein. Wir können also die Koeffizienten in den einzelnen Ordnungen von n vergleichen. Wir starten von oben: In der vierten Ordnung entsteht die triviale Gleichung $1 = 1$, die keine Informationen liefert. In der dritten Ordnung erhalten wir

$$a_3 + 4a_4 = a_3 + 1, \quad \text{also} \quad a_4 = \frac{1}{4}.$$

In der zweiten Ordnung lesen wir die folgende Gleichung ab:

$$a_2 + 3a_3 + 6a_4 = a_2 + 3, \quad \text{also} \quad a_3 = \frac{1}{2}.$$

Für die erste Ordnung berechnet man

$$a_1 + 2a_2 + 3a_3 + 4a_4 = a_1 + 3, \quad \text{also} \quad a_1 = \frac{1}{4},$$

und in der nullten Ordnung ergibt sich schließlich

$$a_1 + a_2 + a_3 + a_4 = 1, \quad \text{also} \quad a_1 = 0.$$

Die partikuläre Lösung hat somit die Form

$$x_n^{\text{part}} = \frac{1}{4}n^2 + \frac{1}{2}n^3 + \frac{1}{4}n^4 = \frac{n^2 \cdot (n+1)^2}{4}.$$

III Die allgemeine Lösung lautet also $x_n^{\text{all}} = c + \frac{n^2 \cdot (n+1)^2}{4}$.

IV Der Anfangswert $s_0 = 0$ wird gerade für $c = 0$ realisiert.

Insgesamt haben wir die Summenformel

$$\sum_{k=1}^n k^3 = \frac{n^2 \cdot (n+1)^2}{4}$$

hergeleitet. Anhand dieses Beispiels lässt sich erkennen, dass für die Herleitung der Summe über die k -ten Potenzen ein lineares $k \times k$ -Gleichungssystem zu lösen ist, das aber sofort in oberer Dreiecksform auftritt, also leicht zu lösen ist. In Aufg. 3.3 soll dieses Verfahren auf die Summe der Quadratzahlen angewendet werden.

3.4.3 Eine zweite Anwendung: die Fibonacci-Folge

In einer der frühesten bekannten mathematischen Modellierungen biologischer Phänomene beschreibt Leonardo Fibonacci 1202 in seinem Buch *Liber abaci* die Populationsdynamik einer Kaninchenpopulation durch eine Differenzengleichung zweiter Ordnung. Die Modellierungsüberlegungen sehen folgendermaßen aus: Ein Kaninchenpaar benötigt ein Jahr, um geschlechtsreif zu werden. Anschließend sorgt es pro Jahr für ein neues Pärchen, welches wiederum ein Jahr benötigt, um geschlechtsreif zu werden. Ansonsten haben die Kaninchen die bemerkenswerte Eigenschaft, unsterblich zu sein. Zu Anfang ist ein junges Paar vorhanden, das noch nicht geschlechtsreif ist. Damit wird die Entwicklung durch die Differenzengleichung

$$x_{n+1} = x_n + x_{n-1} \quad \text{für alle } n \in \mathbb{N}. \quad (3.7)$$

zusammen mit den Anfangswerten $x_1 = 1$ und $x_2 = 1$ beschrieben. Die Folge beginnt also mit den Zahlen 1; 1; 2; 3; 5; 8; 13; ... Damit eindeutig eine Folge durch diese Gleichung bestimmt ist, müssen nun zwei Anfangswerte x_0 und x_1 vorgegeben werden. Fibonacci wollte die Anzahl von Kaninchenpaaren beschreiben und hat dazu $x_1 = 1$ und $x_2 = 1$ gewählt. Diese Folge erhalten wir auch, wenn wir $x_0 = 0$ und $x_1 = 1$ vorschreiben.

Auch wenn die Modellierungsüberlegungen in Bezug auf Kaninchenpopulationen wohl recht fragwürdig sind, finden sich vor allem in der Pflanzenwelt durchaus Phänomene, die sich durch Fibonacci-Zahlen beschreiben lassen, siehe [65]. Aber auch inner-mathematisch hat sich die Fibonacci-Folge als sehr interessant erwiesen, so dass sie sich sogar eine eigene Fachzeitschrift, den *Fibonacci Quarterly*, verdient hat. Wir wollen hier eine explizite Darstellung der Folge ermitteln, sowie die Beziehung zum *Goldenen Schnitt* in Aufg. 3.5 klären. Auch bei der Differenzengleichung von Fibonacci handelt es sich um eine lineare Gleichung, sogar um eine homogene Gleichung: Haben wir zwei Lösungsfolgen von (3.7), dann löst auch jede Linearkombination der beiden Folgen die Gleichung. Da wir zwei Anfangswerte für eine eindeutig bestimmte Lösung vorgeben müssen, ist der Lösungsraum zweidimensional. Wir starten mit demselben Ansatz wie im vorherigen Kapitel, $x_n = x^n$, und setzen diesen Ansatz in die Differenzengleichung ein:

$$x^{n+1} = x_{n+1} \stackrel{!}{=} x_n + x_{n-1} = x^n + x^{n-1} \quad \text{für alle } n \in \mathbb{N},$$

lösen nach x auf und erhalten die Lösungen $x = 0$ oder $x^2 - x - 1 = 0$. Die zweite Gleichung hat die beiden Lösungen

$$\bar{x}_1 = \frac{1 + \sqrt{5}}{2} \quad \text{und} \quad \bar{x}_2 = \frac{1 - \sqrt{5}}{2}.$$

Das ergibt zwei linear unabhängige Lösungsfolgen $(x_{1,n}) = (\bar{x}_1^n)$ und $(x_{2,n}) = (\bar{x}_2^n)$ und wir können alle Lösungsfolgen der Fibonacci-Gleichung als Linearkombination dieser beiden Folgen darstellen:

$$x_n^{\text{hom}} = c_1 \bar{x}_1^n + c_2 \bar{x}_2^n \quad \text{mit} \quad c_1, c_2 \in \mathbb{R}.$$

Eine partikuläre Lösung müssen wir nicht bestimmen, weil die Gleichung bereits homogen ist. Wir bestimmen jetzt c_1, c_2 so, dass die Anfangswerte $x_0 = 0$ und $x_1 = 1$ erfüllt sind. Diese beiden Forderungen führen auf das lineare Gleichungssystem für c_1 und c_2 :

$$0 = c_1 + c_2 \quad 1 = c_1 \bar{x}_1 + c_2 \bar{x}_2$$

mit der Lösung $c_{1/2} = \pm 1/\sqrt{5}$. Insgesamt erhalten wir also die Lösung

$$x_n = \frac{1}{\sqrt{5}} \left(\left(\frac{1 + \sqrt{5}}{2} \right)^n - \left(\frac{1 - \sqrt{5}}{2} \right)^n \right).$$

Aufgaben

3.1 Immobilienkredite

Zur Finanzierung einer Immobilie wird ein Immobilienkredit der Höhe K_0 mit einer bestimmten Laufzeit L und einem festen jährlichen Zinssatz p aufgenommen. In der Regel wird ein fester Tilgungssatz t verabredet. Die Abzahlung des Kredits verläuft dann folgendermaßen: Der Kreditnehmer zahlt jeden Monat $(p + t)/12$ des aufgenommenen Kredits K_0 an die Bank zurück. Davon streicht die Bank den Anteil $p/12$ der in diesem Monat noch vorhandenen Restschuld als Zinsen ein, der Rest wird zur Tilgung der Schuld verwendet. Es sei a_n die Restschuld nach n Monaten.

- Stellen Sie eine Differenzengleichung für a_n auf und geben Sie eine explizite Lösungsformel an.
- Welche Restschuld ist bei gegebenen Werten von K_0 , p und t nach zehn Jahren noch übrig?
- Nach wie vielen Monaten ist die Schuld bei gegebenem p und t vollständig abbezahlt?
- Wie hoch muss die Tilgungsrate t bei gegebenem Zinssatz p gewählt werden, damit der Kredit in zehn Jahren zurückbezahlt ist?

3.2 Cobwebs für die Sonderfälle

Zeichnen Sie Cobwebs für die lineare DZG $a_{n+1} = r a_n + b$ für die Fälle $r = -1$, $r = 0$ und $r = 1$.

3.3 Die Formel für die Summe der ersten n Quadratzahlen

Die Formel

$$\sum_{k=1}^n k^2 = \frac{n \cdot (2n + 1) \cdot (n + 1)}{6}$$

soll auf drei Weisen hergeleitet werden.

a) Gehen Sie von der Teleskopsumme

$$(n + 1)^3 - 1 = \sum_{k=1}^n (k + 1)^3 - k^3$$

aus und nutzen Sie die bekannten Formeln

$$\sum_{k=1}^n k^0 = n \quad \text{und} \quad D_n := \sum_{k=1}^n k^1 = \frac{n \cdot (n + 1)}{2}.$$

b) Nutzen Sie das lineare Lösungsschema.

c) Lesen Sie die Formel jeweils aus den beiden geometrischen Konfigurationen in Abb. 3.2 heraus.

3.4 Eine Lösungsformel für lineare Differenzgleichungen mit variablen Koeffizienten

In dieser Aufgabe soll auch für die lineare Differenzgleichung mit variablen Koeffizienten eine explizite Lösungsformel hergeleitet werden.

Dazu seien (α_n) und (β_n) zwei vorgegebene reelle Folgen. Weiter sei (a_n) eine Lösung der Differenzgleichung

$$a_{n+1} = \alpha_n \cdot a_n + \beta_n \quad \text{für alle } n \in \mathbb{N}_0.$$

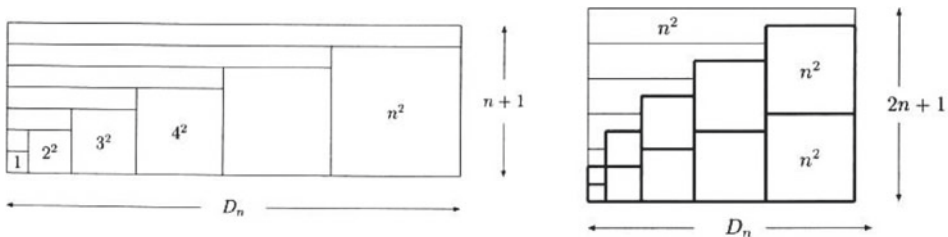


Abb. 3.2 Geometrische Konfigurationen zur Herleitung der Summe über die Quadratzahlen. Abbildungen aus [79] (©(2016)Springer Spektrum)

- a) Es sei $\alpha \in \mathbb{R}$ und es gelte $\alpha_n = \alpha$ für alle $n \in \mathbb{N}_0$. Beweisen Sie, dass $(a_n)_{n \in \mathbb{N}_0}$ das explizite Bildungsgesetz

$$a_n = \alpha^n \cdot a_0 + \sum_{k=0}^{n-1} \alpha^k \cdot \beta_{n-1-k} \quad \text{für } n \in \mathbb{N}_0$$

erfüllt.

- b) Bestimmen Sie im allgemeinen Fall, wenn (α_n) eine beliebig vorgegebene reelle Folge ist, das explizite Bildungsgesetz für die Folge (a_n) und beweisen Sie es.

3.5 Weiteres zur Fibonacci-Folge

Wir betrachten die Fibonacci-Gleichung

$$x_{n+1} = x_n + x_{n-1} \quad \text{für alle } n \in \mathbb{N}.$$

- a) Sei (x_n) eine Lösungsfolge von (3.7) (zu irgendwelchen Anfangswerten). Zeigen Sie, dass die Folge der Quotienten $q_n = x_{n+1}/x_n$ konvergiert, und bestimmen Sie ihren Grenzwert. Inwiefern hängt dieser Grenzwert von den Anfangswerten x_0 und x_1 ab?
- b) Die Essener Kaninchen sind fruchtbarer als die Kaninchen von Fibonacci. Jedes geschlechtsreife Paar wirft jede Woche statt einem Paar zwei Paare neugeborener Kaninchen. Auch diese Kaninchen brauchen wieder eine Woche, um geschlechtsreif zu werden, und sind unsterblich. Anfangs sei wieder ein neugeborenes Paar vorhanden. Finden Sie eine Modellierung und bestimmen Sie eine explizite Lösungsformel.
- c) Bestimmen Sie für $r, s, b \in \mathbb{R}$ die Lösungen der inhomogenen Differenzengleichung

$$a_{n+1} = ra_n + sa_{n-1} + b. \quad (3.8)$$

3.6 Iterative Preisbestimmung

In Aufg. 2.4 wird ein Nachfrage-Preis-Modell für zwei Anbieter diskutiert, das enddimensionalisiert die folgende Form hat:

$$\begin{aligned} \varphi_1(y_1, y_2) &= 1 - y_1 + \gamma y_2, & \gamma_1(y_1, y_2) &= \varphi_1(y_1, y_2)(y_1 - \varepsilon) \\ \varphi_2(y_1, y_2) &= 1 - y_2 + \gamma y_1, & \gamma_2(y_1, y_2) &= \varphi_2(y_1, y_2)(y_2 - \varepsilon). \end{aligned}$$

Dabei sind φ_1 die Nachfrage bei Anbieter 1, y_1 der von Anbieter 1 verlangte Preis, γ_1 der von Anbieter 1 erzielte Gewinn und Entsprechendes für Anbieter 2.

Die Anbieter bestimmen nun nach der folgenden Strategie ihren Preis:

- Anbieter 1 startet mit einem Preis.
- Anbieter 2 justiert jetzt seinen Preis so, dass er bei dem gewählten Preis von Anbieter 1 den maximalen Gewinn erwirtschaftet.

- Jetzt justiert Anbieter 1 seinen Preis neu, so dass er nun den maximalen Gewinn bei dem von Anbieter 2 gewählten Preis erwirtschaftet.
 - Jetzt ändert wieder Anbieter 2 seinen Preis, und so fort.
- a) Modellieren Sie die Preisentwicklung mit Hilfe einer Differenzengleichung.
- b) Wie entwickeln sich die Preise auf lange Sicht? Was hat das Ganze mit dem Nash-Gleichgewicht zu tun?



Insekten zwischen Stabilität und Chaos – Modellieren mit nichtlinearen Differenzengleichungen

4

Inhaltsverzeichnis

4.1 Die logistische Differenzengleichung – ein Prototyp für den Übergang von einem stabilen zu einem chaotischen Verhalten	81
4.2 Das Prinzip der Linearisierung	86
4.3 Zurück zur logistischen Differenzengleichung	89
4.4 Insektenpopulation mit nichtüberlappender Generationenfolge	93
Aufgaben	103

4.1 Die logistische Differenzengleichung – ein Prototyp für den Übergang von einem stabilen zu einem chaotischen Verhalten

Wir wollen die zeitliche Entwicklung einer Population gewisser Lebewesen in einem gewissen Gebiet mit Hilfe von Differenzengleichungen beschreiben, z. B. Hirsche in einem Waldgebiet, Insekten auf einer Wirtspflanze oder Bakterien in einer Petrischale. In festen Zeitabständen Δt wird eine Zählung/Messung des Bestandes vorgenommen, z. B. einmal im Jahr vor der Brunft, nach der Geburt oder dem Schlüpfen, ... Der Bestand der n -ten Zählung wird mit a_n bezeichnet. Die Dimension von der Größe a_n kann dabei unterschiedlich sein:

- eine Anzahl, z. B. bei Großtieren,
- Biomasse, z. B. bei Algen,
- die besiedelte Fläche, z. B. bei Bakterien,
- ...

4.1.1 Das lineare Modell

Im n -ten Zeitschritt ändert sich die Populationsgröße im betrachteten Gebiet wie folgt:

$$a_{n+1} = a_n + \text{Geburten} - \text{Todesfälle}$$

im betrachteten Zeitintervall. Wir wollen von den folgenden Modellierungsannahmen ausgehen: Die Anzahlen der Geburten- und Todesfälle im betrachteten Zeitintervall sind proportional zur Größe der am Anfang des Zeitintervalls vorhandenen Population a_n , wobei die Proportionalitätskonstanten durch äußere, zeitlich konstante Bedingungen bestimmt sind. Wir führen also die Geburtskonstante $b \geq 0$ und die Sterbewahrscheinlichkeit $0 \leq d \leq 1$ für das Zeitintervall Δt ein und erhalten das Modell

$$a_{n+1} = a_n + ba_n - da_n = (1 + b - d)a_n. \quad (4.1)$$

Die Lösungen dieser Gleichung kennen wir bereits: Eine positive Anfangspopulation a_0 wird entweder aussterben, nämlich dann, wenn $1 + b - d < 1$ ist, oder über alle Grenzen wachsen, nämlich dann, wenn $1 + b - d > 1$ ist, die Geburtskonstante also größer als die Sterbewahrscheinlichkeit ist. Die Konstante $r := 1 + b - d$ scheint wichtig zu sein und wird auch als die *Pro-Kopf-Reproduktionsrate* bezeichnet, bezogen auf den Zeitschritt Δt .

4.1.2 Das logistische Modell

Während nach dem linearen Modell die Population entweder ausstirbt oder über alle Grenzen wächst, beobachtet man in der Realität häufig, dass die Bestände sich bei einer gewissen Größe stabilisieren und nur geringfügig um diese Größe herum in einer zufällig anmutenden Art und Weise variieren. Das lineare Modell muss also verändert werden, um das angesprochene Verhalten darstellen zu können. Im linearen Modell ergibt sich die Änderung des Bestandes durch $a_{n+1} = ra_n$. Wir modellieren jetzt r abhängig von der Populationsgröße $r = r(a_n)$, und zwar so, dass $r(a)$ positiv ist, wenn $a < K$, und negativ, wenn $a > K$ ist. Dabei ist K eine bestimmte charakteristische Populationsgröße. (In späteren Modellen wird mit K die Kapazität des betrachteten Lebensraums bezeichnet. Hier spielt K jedoch nicht diese Rolle.) Das mathematisch einfachste Modell mit diesen Eigenschaften ist wiederum linear, also

$$r(a) = r_0 - \frac{r_0}{K}a.$$

Zur Vereinfachung der Notation bezeichnen wir die Konstante r_0 wieder mit r :

$$a_{n+1} = \left(r - \frac{r}{K}a_n\right)a_n = ra_n \left(1 - \frac{a_n}{K}\right),$$

oder anders ausgedrückt:

$$a_{n+1} = f(a_n) \quad \text{mit} \quad f(a) = ra \left(1 - \frac{a}{K}\right). \quad (4.2)$$

Diese Gleichung nennen wir die *logistische Differenzengleichung*. Wir berechnen zunächst die stationären Punkte der Gl. (4.2), also diejenigen Populationsgrößen, die zeitlich konstant

bleiben. Die Gleichung

$$a = ra \left(1 - \frac{a}{K}\right)$$

hat die beiden Lösungen $a_0^* = 0$ und $a_1^* = \left(1 - \frac{1}{r}\right) K$.

Wir wollen nun zunächst das Modell entdimensionalisieren, um einen Parameter loszuwerden. Genau wie im ersten Kapitel folgen wir dem Schema

- dimensionslose Variable einführen

$$x_n := \frac{a_n}{\bar{a}} \quad \text{also} \quad a_n = \bar{a} x_n,$$

- die dimensionslose Gleichung ausrechnen

$$x_{n+1} = rx_n \left(1 - \frac{\bar{a}}{K} x_n\right),$$

- und die Skala wählen: $\bar{a} = K$.

Unser entdimensionalisiertes Modell

$$x_{n+1} = rx_n(1 - x_n) \tag{4.3}$$

hat die stationären Punkte $x_0^* = 0$ und $x_1^* = 1 - \frac{1}{r}$.

Eine kurze Bemerkung: Unter reinen Modellierungsgesichtspunkten ist die logistische Differenzgleichung durchaus fragwürdig (warum?). Ihre eigentliche Funktion hier liegt in der Generierung wichtiger Phänomene an einem mathematisch möglichst einfachen Beispiel. Im Gegensatz dazu wird die logistische Differentialgleichung in Kap. 5 auch bei der Modellierung von Realsituationen von großer Wichtigkeit sein.

4.1.3 Von stabilen Zuständen in das Chaos

Bevor wir weitergehende analytische Betrachtungen vornehmen, wollen wir das Verhalten der Lösungen von Gl. (4.3) numerisch untersuchen. Dazu benutzen wir ein Tabellenkalkulationsprogramm, implementieren die Folge (x_n) und variieren den Parameter r zwischen $0 < r \leq 4$, den Startwert x_0 zwischen $0 < x_0 < 1$ und lassen uns die Graphen jeweils plotten. Einige ausgewählte Graphen mit dem Startwert $x_0 = 0,5$ sind in Abb. 4.1 dargestellt. Die unscheinbare DZG (4.3) weist also eine sehr reichhaltige Dynamik auf. Wir fassen die Befunde zusammen und nutzen die Bezeichnung $x_0^* = 0$ und $x_1^* = 1 - \frac{1}{r}$. Für den Anfangswert x_0 gilt stets $0 < x_0 < 1$.

1. Für $0 < r < 1$ konvergiert die Lösungsfolge für alle Startwerte streng monoton gegen x_0^* .

2. Für $1 < r < 2$ gilt für alle Startwerte $\lim_{n \rightarrow \infty} x_n = 1 - \frac{1}{r}$. Die Folge ist streng monoton, und zwar fallend, falls $x_0 > 1 - \frac{1}{r}$, bzw. steigend, falls $x_0 < 1 - \frac{1}{r}$. Falls $x_0 = x^*$, ist die Lösungsfolge konstant.
3. Für $2 < r < 3$ konvergiert die Folge für alle Startwerte gegen x^* , und zwar oszillierend: Eine der beiden Teilfolgen (x_{2n}) und (x_{2n+1}) konvergiert monoton fallend und die andere monoton steigend.
4. Für $3 < r \lesssim 3,27$ gibt es zwei Häufungspunkte x_2^* und x_3^* mit $x_0^* < x_2^* < x_1^* < x_3^* < 1$ und die Lösungsfolge springt zwischen ihnen hin und her. Es gilt

$$\lim_{n \rightarrow \infty} x_{2n} = x_{2/3}^* \quad \text{und} \quad \lim_{n \rightarrow \infty} x_{2n+1} = x_{3/2}^*.$$

Die Konvergenz beider Teilfolgen ist monoton. Welche Teilfolge gegen welchen Häufungspunkt konvergiert, hängt vom Anfangsdatum ab.

5. Für $3,27 \lesssim r \lesssim 3,45$ konvergieren die Teilfolgen oszillierend gegen die beiden Häufungspunkte.
6. Für $3,45 \lesssim r \lesssim 3,55$ gibt es vier Häufungspunkte x_4^*, x_5^*, x_6^* und x_7^* mit

$$x_0^* < x_4^* < x_2^* < x_5^* < x_1^* < x_6^* < x_3^* < x_7^* < 1.$$

Jeder der beiden Häufungspunkte aus 4. hat sich gewissermaßen in zwei neue Häufungspunkte aufgespalten.

7. Für $3,55 \lesssim r \lesssim 3,55$ gibt es acht Häufungspunkte.
8. Für $3,56 \lesssim r < 4$ scheint das Verhalten der Lösungsfolge chaotisch zu sein. Chaotisches Verhalten ist insbesondere dadurch charakterisiert, dass Lösungen zu Anfangswerten, die einen beliebig kleinen, aber positiven Abstand voneinander haben, sich schnell weit auseinanderentwickeln können. Dieses Verhalten ist in Abb. 4.1 im letzten Graphen illustriert, in dem die Lösungsfolgen zu den Anfangswerten $x_0 = 0,5$ und $\tilde{x}_0 = 0,5001$ geplottet sind.

Im nächsten Kapitel wird dieses überraschende Verhalten bei Variation des Parameters r zumindest ansatzweise geklärt. Ferner werden wir dann mit der dort vorgestellten Methode *berechnen* können, wann sich das Verhalten ändert, und es nicht nur staunend zur Kenntnis nehmen. Zentral ist dabei das Prinzip der Linearisierung: *Wenn dir ein System zu schwierig erscheint, dann linearisiere es um die Stelle, die dich gerade interessiert.*

Das Phänomen des *deterministischen Chaos* spielt in vielen Situationen eine Rolle. Der Begriff *deterministisch* bedeutet, dass die Entwicklung des Systems durch den Anfangszustand eindeutig festgelegt ist (im Gegensatz zu den stochastischen Systemen). Der Begriff *Chaos* umfasst das obige Phänomen: Egal wie nah zwei Anfangszustände zusammenliegen, kann ihre Entwicklung beliebig schnell auseinandergehen. Das hat weit reichende Auswirkungen auf die Vorhersagekraft solcher Modelle: In der Praxis lassen sich Anfangszustände in der Regel nur mit einer endlichen Genauigkeit messen. Deshalb lassen sich keine sicheren Prognosen über die Entwicklung des Systems treffen, weil sich im Rahmen der Messge-

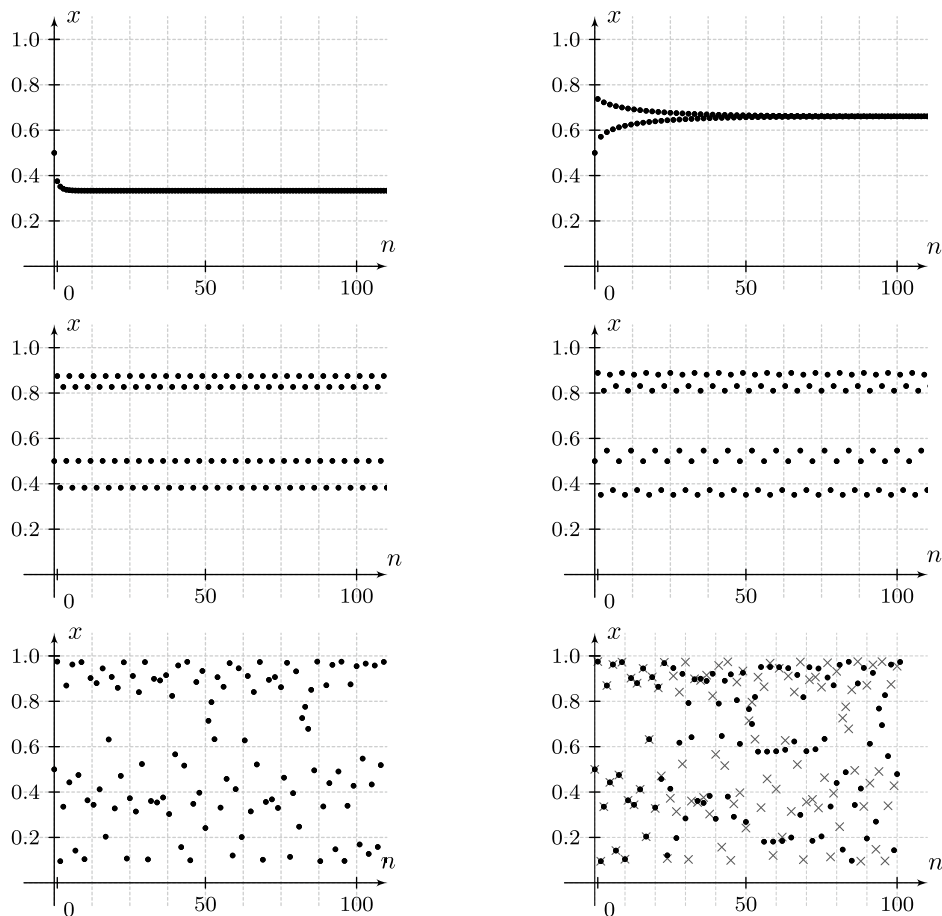


Abb. 4.1 Die logistische Folge $x_{n+1} = rx_n(1 - x_n)$ mit Anfangswert $x_0 = 0,5$ und $r = 1,5, 2,95, 3,55, 3,555$ und $3,9$. In der letzten Abbildung ist zusätzlich die Folge mit dem Anfangswert $\tilde{x}_0 = 0,5001$ mit Kreuzen geplottet

nauigkeit viele Anfangszustände befinden, die sich vollkommen unterschiedlich entwickeln können. Ein prominentes Beispiel für ein System, das chaotisches Verhalten zeigt, ist das Wetter. Hier gilt der Spruch vom Flügelschlag eines Schmetterlings (kleine Abweichung in den Anfangsbedingungen), der einen Orkan auslösen kann (großer Unterschied in der Entwicklung).

4.2 Das Prinzip der Linearisierung

Um die reichhaltige Dynamik der logistischen Differenzengleichung in den Griff zu bekommen, führen wir einige Bezeichnungen ein, die helfen sollen, das unterschiedliche Verhalten der Lösungsfolgen zu beschreiben. Eine wichtige Rolle scheinen wieder die stationären Punkte zu spielen.

Definition 4.1

Gegeben sei eine Differenzengleichung

$$x_{n+1} = f(x_n) \quad \text{für alle } n \in \mathbb{N}. \quad (4.4)$$

- i) Ein Punkt x^* heißt *stationärer Punkt*, wenn $f(x^*) = x^*$ gilt.
- ii) Ein stationärer Punkt x^* von (4.4) heißt *asymptotisch stabil*, wenn es eine Umgebung U um x^* gibt, sodass für jede Lösungsfolge $(x_n)_{n \in \mathbb{N}}$ von (4.4) gilt: Enthält U ein Folgenglied der Folge (x_n) , dann konvergiert (x_n) gegen x^* . Wir nennen einen asymptotisch stabilen stationären Punkt
 - *monoton*, wenn jede Lösungsfolge von (4.4), die gegen x^* konvergiert, ab irgendeinem Index monoton ist, und
 - *oszillierend*, wenn jede Lösungsfolge von (4.4), die gegen x^* konvergiert, ab irgendeinem Index um x^* oszilliert, also abwechselnd größer und kleiner als x^* ist.
- iii) Ein stationärer Punkt x^* von (4.4) heißt *stabil*, wenn es für jede Umgebung U von x^* eine weitere Umgebung V von x^* gibt, sodass für jede Lösungsfolge (x_n) von (4.4) das Folgende gilt: Aus $x_{\hat{n}} \in V$ folgt $x_n \in U$ für alle $n \geq \hat{n}$.
- iv) Ein stationärer Punkt x^* heißt *instabil*, wenn er weder asymptotisch stabil noch stabil ist. ♦

Bemerkung 1

- i) Wir können die beiden Stabilitätskonzepte sprachlich so fassen: Ein stationärer Punkt ist *asymptotisch stabil*, wenn jede Folge, die dem Punkt hinreichend nah kommt, gegen den Punkt konvergiert. Ein stationärer Punkt heißt *stabil*, wenn jede Folge, die dem Punkt hinreichend nah kommt, in der Nähe des Punktes bleibt.
- ii) Die lineare Differenzengleichung $f(x) = rx + b$ hat für $r \neq 1$ den einzigen stationären Punkt $x^* = \frac{b}{1-r}$, und dieser ist
 - *instabil* für $|r| > 1$,
 - *stabil* für $-1 \leq r < 1$,
 - *asymptotisch stabil* für $-1 < r < 1$,
 - *oszillierend asymptotisch stabil* für $-1 < r < 0$,
 - *monoton asymptotisch stabil* für $0 < r < 1$.

Für $r = -1$ ist der Punkt x^* stabil, aber nicht asymptotisch stabil. Für $r = 1$ und $b = 0$ ist jeder Punkt $x^* \in \mathbb{R}$ stabil, aber nicht asymptotisch stabil, denn alle Lösungsfolgen sind ab dem zweiten Folgenglied konstant. Für $r = 1$ und $b \neq 0$ gibt es keine stationären Punkte.

Es gibt auch für nichtlineare Differenzengleichungen einfache Stabilitätskriterien: In vielen Fällen vererbt sich das Stabilitätsverhalten der Linearisierung um einen stationären Punkt auf das Stabilitätsverhalten der nichtlinearen Gleichung. Das wird im folgenden Satz formuliert und anschließend bewiesen.

Satz 4.1 (Satz von der linearisierten Stabilität)

Es sei $f : \mathbb{R} \rightarrow \mathbb{R}$ stetig differenzierbar und x^* ein stationärer Punkt der Differenzengleichung

$$x_{n+1} = f(x_n) \quad \text{für alle } n \in \mathbb{N}.$$

Dann gelten die folgenden Stabilitätskriterien:

i) Für $|f'(x^*)| < 1$ ist x^* asymptotisch stabil. Weiterhin gilt:

- Für $-1 < f'(x^*) < 0$ ist x^* oszillierend.
- Für $0 < f'(x^*) < 1$ ist x^* monoton.

ii) Für $|f'(x^*)| > 1$ ist x^* instabil.

Bevor wir den Beweis angehen, soll zunächst einmal die zugrunde liegende fundamentale Idee dargestellt werden.

Das Linearisierungsrezept

Ist ein Problem zu kompliziert, um es direkt anzugehen, ersetze es lokal an der Stelle, die von Interesse ist, durch die Linearisierung des Problems.

Dazu schreiben wir die Lösungsfolge (x_n) als stationären Punkt zuzüglich der Auslenkung aus dem stationären Punkt:

$$x_n = x^* + r_n \quad \text{bzw.} \quad r_n = x_n - x^*.$$

Die Folge der Auslenkungen (r_n) löst dann die Differenzengleichung:

$$r_{n+1} = x_{n+1} - x^* = f(x_n) - x^* = f(x^* + r_n) - x^*$$

Wir nutzen jetzt zum einen, dass x^* ein stationärer Punkt ist, und nehmen zum anderen an, dass die Auslenkungen aus dem stationären Punkt klein sind, also $|r_n| \ll 1$. Die Linearisierung des Problems besteht nun darin, die Funktion f durch ihre lineare Approximation an der Stelle x^* zu ersetzen:

$$r_{n+1} \approx f(x^*) + f'(x^*)r_n - x^* = f'(x^*)r_n \quad (4.5)$$

Sie löst also, solange sie nah an dem stationären Punkt x^* verbleibt, approximativ eine lineare Differenzengleichung. Der Linearisierungssatz 4.1 sagt nun gerade aus, dass, bis auf die Grenzfälle $f'(x^*) = \pm 1$, das nichtlineare Problem die Stabilität vom linearen Problem erbt.

Beweis 4.1 Der Beweis folgt aus dem Mittelwertsatz der Differentialrechnung. Es gilt

$$x_{n+1} - x^* = f(x_n) - f(x^*) = f'(\xi_n)(x_n - x^*) \quad (4.6)$$

mit einer Stelle ξ_n zwischen x_n und x^* .

- i) Falls $|f'(x^*)| < 1$, gibt es aufgrund der Stetigkeit von f' ein $0 < q < 1$ und eine ganze Umgebung U um x^* , sodass

$$|f'(x)| \leq q \quad \text{für alle } x \in U.$$

Haben wir nun $x_n \in U$ für ein $n \in \mathbb{N}$ und es gelte, sagen wir, $x_{\hat{n}} - x^* > 0$, folgt mit (4.6)

$$|x_{\hat{n}+1} - x^*| \leq q |x_{\hat{n}} - x^*|$$

und durch Induktion nach n

$$|x_{\hat{n}+n} - x^*| \leq q^n |x_{\hat{n}} - x^*| \quad \text{für alle } n \in \mathbb{N}.$$

Die Folge (x_n) konvergiert also gegen x^* . Falls $0 < f'(x^*) < 1$, kann die Umgebung U zusätzlich so gewählt werden, dass

$$0 < f'(x) < 1 \quad \text{für alle } x \in U.$$

Ebenfalls mit (4.6) folgt, dass sich dann das Vorzeichen der Differenz $x_{\hat{n}+n} - x^*$ auf das Vorzeichen der folgenden Differenz $x_{\hat{n}+n+1} - x^*$ vererbt. Die Folgen konvergieren also schließlich monoton. Mit der entsprechenden Argumentation ist klar, dass im Fall $-1 < f'(x^*) < 0$ die Folge schließlich um den stationären Punkt oszilliert.

- ii) Wir nehmen nun an, dass $|f'(x^*)| > 1$. Wie eben finden wir ein $q \in \mathbb{R}$ und eine beschränkte Umgebung U um x^* , sodass

$$1 < q \leq |f'(x)| \quad \text{für alle } x \in U.$$

Sei nun $x_{\hat{n}} \in U$. Für alle anschließenden Folgeglieder, die sich noch in U befinden gilt die Abschätzung

$$|x_{\hat{n}+n} - x^*| \geq q^n |x_{\hat{n}} - x^*|.$$

Im Fall $x_{\hat{n}} \neq x^*$ muss die Folge also die Umgebung U nach endlich vielen Schritten verlassen, egal wie nah $x_{\hat{n}}$ auch an x^* gelegen haben mag. Damit ist der Punkt x^* instabil. Das Monotonie- bzw. Oszillationsverhalten folgt wie im stabilen Fall aus dem Vorzeichen von $f'(x^*)$.

4.3 Zurück zur logistischen Differenzengleichung

Das neu erworbene Rüstzeug soll jetzt auf die logistische Differenzengleichung angewendet werden.

4.3.1 Die Stabilität der stationären und 2-periodischen Punkte und das Bifurkationsdiagramm

Zur Bestimmung der Stabilität der beiden stationären Punkte $x_0^* = 0$ und $x_1^* = 1 - \frac{1}{r}$ berechnen wir die Ableitung der Funktion $f(x) = rx(1 - x)$ und setzen die beiden Punkte ein. Wir erhalten:

$$f'(x) = r(1 - 2x)$$

$$f'(x_0^*) = r$$

$$f'(x_1^*) = 2 - r$$

Mit Satz 4.1 schließen wir:

- i) x_0^* ist monoton stabil für $0 < r < 1$ und monoton instabil für $1 < r$.
- ii) x_1^* ist
 - monoton instabil für $0 < r < 1$,
 - monoton stabil für $1 < r < 2$,
 - oszillierend stabil für $2 < r < 3$ und
 - oszillierend instabil für $3 < r$.

Aufgrund der numerischen Befunde erwarten wir für $3 < r$ die Häufungspunkte x_2^* und x_3^* , die jeweils Grenzwert der geradzahlgigen bzw. der ungeradzahlgigen Teilfolgen von (x_n) sind. Diese beiden Punkte sollten also stationäre Punkte der Differenzenfolge $x_{n+1} = f(f(x_n))$ sein und somit die Gleichung $x^* = f(f(x^*))$ lösen. Solche Punkte, die keine stationären Punkte von f selber sind, nennen wir die *2-periodischen* Punkte von f . Entsprechend definieren wir auch die *n-periodischen* Punkte. Um die 2-periodischen Punkte zu bestimmen, muss die Gleichung

$$x = f(f(x)) = r^2(1-x)x(1-r(1-x)x)$$

gelöst werden. Wir überlegen erst einmal, ob wir dazu in der Lage sind. Es handelt sich um eine Gleichung vierter Ordnung, zu der es eine algebraische Lösungsformel gibt. Da wir ja aber schon die beiden Lösungen x_0^* und x_1^* kennen, könnten wir auch die beiden Linearfaktoren $x - x_0^*$ und $x - x_1^*$ durch Polynomdivision herausdividieren und müssten nur noch eine quadratische Gleichung lösen, was auch händisch geht. Nachdem wir uns davon überzeugt haben, dass wir es nicht müssten, verwenden wir trotzdem ein Computer-Algebra-System und erhalten die 2-periodischen Lösungen

$$x_2^* = \frac{1}{2r} \left(r + 1 + \sqrt{r^2 - 3r - 2} \right) \quad \text{und} \quad x_3^* = \frac{1}{2r} \left(r + 1 - \sqrt{r^2 - 3r - 2} \right).$$

Man erkennt, dass diese beiden Lösungen gerade bei $r = 3$ reell werden. Um die Stabilität der 2-periodischen Lösungen zu bestimmen, müssen wir die Ableitung von $f \circ f$ bilden und die beiden Punkte einsetzen. Das sieht unangenehm aus, glücklicherweise gibt es aber einen Trick:

Zunächst gilt $(f(f(x)))' = f'(f(x)) \cdot f'(x)$. Für die beiden 2-periodischen Lösungen gilt $f(x_2^*) = x_3^*$ und $f(x_3^*) = x_2^*$. Das erkennt man folgendermaßen: Es gilt

$$f(f(f(x_2^*))) = f(f(f(x_2^*))) = f(x_2^*).$$

Also ist auch $f(x_2^*)$ eine Lösung von $x = f(f(x))$ und somit in der Menge $\{x_0^*, x_1^*, x_2^*, x_3^*\}$ enthalten. Damit muss $f(x_2^*) = x_3^*$ gelten, weil x_2^* ansonsten bereits ein stationärer Punkt von f wäre. Wir erhalten

$$(f \circ f)'(x_2^*) = f'(x_2^*) \cdot f'(x_3^*) = 4 + 2r - r^2 = (f \circ f)'(x_3^*).$$

Die beiden Punkte wechseln also ihr Stabilitätsverhalten bei

- $r = 1 + \sqrt{5} = 3.236$ von monoton asymptotisch stabil zu oszillierend asymptotisch stabil und bei
- $r = 1 + \sqrt{6} = 3.449$ von oszillierend asymptotisch stabil zu oszillierend instabil.

Aufgrund der numerischen Ergebnisse erwarten wir für $r > 1 + \sqrt{6}$ das Einsetzen von 4-periodischen Lösungen, deren Stabilität mit wachsendem r von monoton über oszillierend asymptotisch stabil zu oszillierend instabil wechselt. Können wir diese 4-periodischen Lösungen berechnen? $f \circ f \circ f$ ist ein Polynom vom Grad acht. Von der Gleichung $(f \circ f \circ f)(x) = x$ kennen wir bereits vier Lösungen. Durch Polynomdivision lässt sich der Grad der zu lösenden Gleichung auf vier verringern und dafür gibt es eine algebraische Lösungsformel. Die zu lösende Gleichung $(f \circ f \circ f \circ f)(x) - x = 0$ ist dann vom Grad 16, acht Lösungen sind bekannt, das Problem ist also auf eine Gleichung achten Grades reduzierbar. Dafür gibt es aber keine algebraischen Lösungsformeln mehr. Die

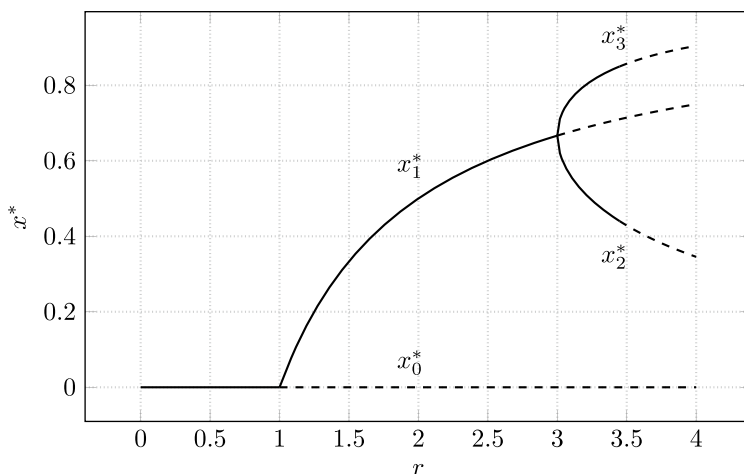


Abb. 4.2 Bifurkationsdiagramm der logistischen Differenzengleichungen mit den stationären und den 2-periodischen Punkten. Durchgezogene Kurven stehen für asymptotisch stabile Punkte, gestrichelte Kurven für instabile Punkte

Galois-Theorie sagt ja, dass es zu Gleichungen vom Grad größer als vier keine algebraischen Lösungsformeln geben kann. Wir sind also an dieser Stelle mit unserem algebraischen Latein am Ende. Immerhin können wir das Verhalten bis $r = 1 + \sqrt{6}$ restlos aufklären. Die erzielten Ergebnisse lassen sich übersichtlich in einem sogenannten *Bifurkationsdiagramm* festhalten. In einem Bifurkationsdiagramm werden die stationären Punkte und in diesem Fall auch die periodischen Punkte gegen den sogenannten *Bifurkationsparameter*, in diesem Fall r , geplottet. Die stationären und 2-periodischen Punkte sind in Abb. 4.2 dargestellt. Ein weitergehendes Bifurkationsdiagramm der logistischen DZG lässt sich mit numerischen Mitteln erstellen, siehe Abb. 4.3. Es weist eine interessante selbstähnliche Struktur auf, die im Rahmen der Theorie diskreter dynamischer Systeme untersucht wird, siehe z. B. [78]. In Aufg. 4.5 soll so ein Bifurkationsdiagramm in Ansätzen mit einem TK erstellt werden.

4.3.2 Der Satz von Sarkovskii und die Feigenbaum-Konstante

In der logistischen Differenzengleichung erscheinen mit zunehmenden Bifurkationsparametern in dieser Reihenfolge die 2-, die 4- und die 8-periodischen Punkte, bevor die Gleichung ein chaotisch anmutendes Verhalten aufweist. Es stellt sich die Frage, ob das eine spezielle Eigenschaft der logistischen Differenzengleichung ist oder eine allgemeine Eigenschaft von Differenzengleichungen. Eine überraschende Antwort auf diese Frage gibt der Satz von Sarkovskii, den der ukrainische Mathematiker Oleksandr Sarkovskii im Jahr 1964 bewies, siehe [90]. In der Kurzversion sagt er das Folgende: Wenn $f : \mathbb{R} \rightarrow \mathbb{R}$ eine stetige Funktion ist und die Differenzengleichung $x_{n+1} = f(x_n)$ einen 3-periodischen Punkt besitzt, dann

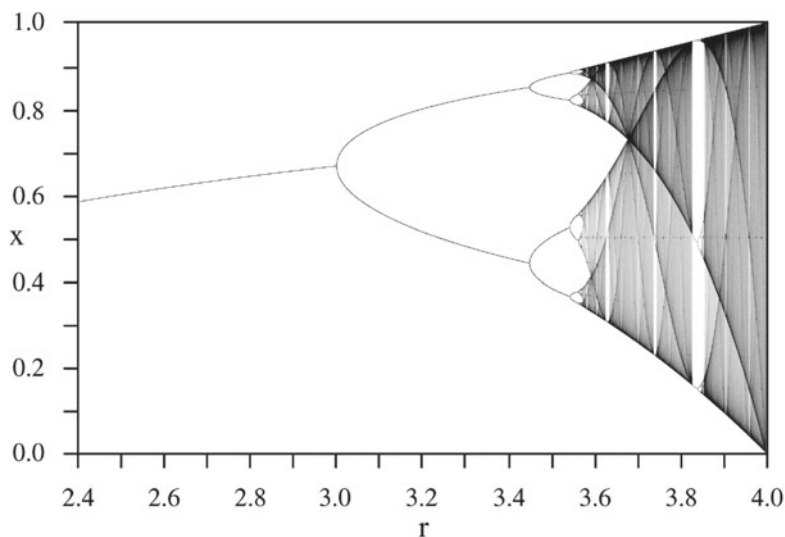


Abb. 4.3 Numerisch erstelltes Bifurkationsdiagramm der logistischen Differenzengleichung. (Quelle: gemeinfrei, <https://commons.wikimedia.org/w/index.php?curid=323398>)

besitzt die Differenzengleichung für alle $n \in \mathbb{N}_0$ auch einen n -periodischen Punkt. Noch kürzer und einprägsamer: Periode 3 bedeutet Chaos.

Genauer gilt das Folgende: Auf der Menge der natürlichen Zahlen definieren wir die sogenannte Sarkovskii-Ordnung $\lesssim$. Diese beginnt mit den ungeraden Zahlen größer eins in aufsteigender Reihenfolge:

$$3 \lesssim 5 \lesssim 7 \lesssim \dots \quad (4.7)$$

Es schließen sich die zweifachen der ungeraden Zahlen in aufsteigender Reihenfolge an

$$2 \cdot 3 \lesssim 2 \cdot 5 \lesssim 2 \cdot 7 \dots,$$

darauf die vierfachen, die achtfachen usw. der ungeraden Zahlen größer eins. Zum Schluss kommen die Zweierpotenzen in abfallender Reihung:

$$\dots \lesssim 2^4 \lesssim 2^3 \lesssim 2^2 \lesssim 2^1 \lesssim 2^0 = 1$$

Diese Relation definiert eine totale Ordnung auf $\mathbb{N}_0$, d.h., sie ist reflexiv, antisymmetrisch, transitiv und total. Der Satz von Sarkovskii sagt im Detail: Gilt $a \lesssim b$ und hat f einen a -periodischen Punkt, dann hat f auch einen b -periodischen Punkt. Periodische Punkte implizieren also in jedem Fall auch stationäre Punkte. Gibt es einen periodischen Punkt, dessen Periode keine Potenz von Zwei ist, gibt es unendlich viele periodische Punkte und unendlich viele verschiedene Perioden. Interessanterweise ist der Beweis des Satzes zwar recht lang, aber sehr elementar. Hauptwerkzeuge sind der Zwischenwertsatz für stetige Funktionen, der

stets für die Existenz von periodischen Punkten sorgt, sowie einfache Teilbarkeitstheorie. Eine ausführliche Darstellung des Beweises findet sich in [1].

Nicht nur die Reihenfolge, in der periodische Lösungen auftreten, ist universal, sondern auch das Muster der Periodenverdoppelung weist eine universelle Eigenschaft auf, die von Mitchell J. Feigenbaum entdeckt wurde [26]. Dazu sei (r_n) die Folge der Bifurkationsparameter, welche gerade die Parameterwerte enthält, zu denen eine Periodenverdoppelung eintritt. Für die betrachtete logistische Differenzgleichung haben wir $r_1 = 3$ und $r_2 = 1 + \sqrt{6}$ berechnet. Allgemein gilt, dass die Folge der Quotienten von aufeinanderfolgenden Verdoppelungsparametern gegen die sogenannte Feigenbaum-Konstante δ konvergiert,

$$\lim_{n \rightarrow \infty} \frac{r_{n-1} - r_n}{r_n - r_{n+1}} = \delta,$$

wobei eine numerische Approximation durch

$$\delta \approx 4,66920$$

angegeben wird. Ein numerischer Wert für diese Konstanten wurde erstmals von S. Großmann und dem Essener Physiker Stefan Thomae publiziert [36]. Bis heute scheint die Frage offen zu sein, ob diese Konstante algebraisch oder transzendent ist.

4.4 Insektenpopulation mit nichtüberlappender Generationenfolge

Differenzgleichungen haben von ihrer Natur her einen diskontinuierlichen Charakter und eignen sich daher insbesondere zur Beschreibung von Vorgängen, die ebenfalls einen diskontinuierlichen Charakter haben. Interessante Beispiele finden sich bei der Modellierung von Populationsgrößen bestimmter Insektenarten mit *nichtüberlappender Generationenfolge*. Wir stellen zunächst eine Familie von Grundmodellen auf und diskutieren anschließend eine Bekämpfungsstrategie.

4.4.1 Grundmodelle für Insektenpopulationen mit nichtüberlappender Generationenfolge

Namensgebend für solche Arten ist es, dass es eindeutig identifizierbare aufeinander folgende Generationen gibt mit der Eigenschaft, dass die vorhergehende Generation ausgestorben ist, bevor die Mitglieder der nachkommenden Generation aus den Eiern schlüpfen. Ganz grob gesprochen besteht also ein Lebenszyklus aus dem Schlüpfen, dem Überleben bis zur Paarung, der Eiablage und dem Sterben. Pro Generation soll einmal die Größe der Population erhoben werden. Wir einigen uns auf den Zeitpunkt direkt nach dem Schlüpfen. Der Wert a_n gibt damit die Anzahl der Insekten der betrachteten Population in der n -ten Generation

an, die aus den Eiern schlüpfen. Ziel ist ein Modell, welches die Entwicklung dieses Wertes über viele Generationen hinweg beschreiben soll. Wir sehen davon ab, dass Insekten sich im Allgemeinen geschlechtlich fortpflanzen und paaren müssen, sondern wollen lediglich zwei Effekte berücksichtigen: die Wahrscheinlichkeit für ein Insekt, vom Schlüpfen bis zur Eiablage zu überleben, und die durchschnittliche Anzahl von Nachkommen pro Insekt, das Eier ablegt, siehe Abb. 4.4. Wir beginnen mit dem zweiten Effekt und nennen r die durchschnittliche Anzahl von Nachkommen pro Insekt, die Pro-Kopf-Reproduktionsrate PKR, oder englisch *Per-capita Reproduction Rate*. Damit haben wir das Grundmodell

$$a_{n+1} = a_n \cdot s \cdot r, \quad (4.8)$$

wobei s die Wahrscheinlichkeit für ein Insekt angibt, vom Schlüpfen bis zur Eiablage zu überleben. Es erscheint plausibel, diese Wahrscheinlichkeit als abhängig von der Anzahl der mitschlüpfenden Insekten zu modellieren ($s = s(a_n)$), weil die Insekten häufig in einem Konkurrenzverhältnis um Futterpflanzen, Verstecke und Ähnliches stehen.

Es gibt eine Vielzahl von Modellen für die Überlebenswahrscheinlichkeit $s(a)$. Wir starten mit dem Modell der *exakten Kompensation* von Hassell.

4.4.2 Das Hassell-Modell mit exakter Kompensation

Wir stellen die Forderungen auf, dass

- 1.) die Überlebenswahrscheinlichkeit mit wachsender Populationsgröße abnimmt und gegen null geht, s also eine monoton fallende Funktion von a ist mit $\lim_{a \rightarrow \infty} s(a) = 0$,
- 2.) für verschwindend geringe Populationen das Überleben so gut wie sicher ist, also $s(0) = 1$,
- 3.) die Überlebenswahrscheinlichkeit sich für große Populationen antiproportional zur Populationsgröße verhält, man spricht von *exakter Kompensation*.

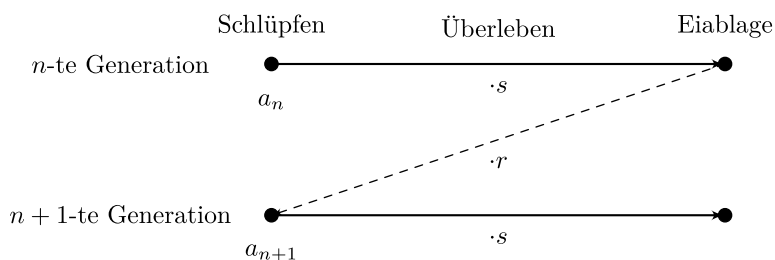


Abb. 4.4 Grundmodell für Insektenpopulationen mit nichtüberlappender Generationenfolge

Auf den ersten Blick scheinen 2.) und 3.) unvereinbar, allerdings soll 3.) nur für hinreichend große Populationen gültig sein. Eine geeignete Mathematisierung besteht aus der Forderung

$$\lim_{a \rightarrow \infty} a \cdot s(a) = c$$

mit einer passenden Konstante c . Wir müssen nun einen passenden Funktionsterm finden, der die Forderungen 1.) bis 3.) erfüllt und möglichst einfach ist. Die vielleicht einfachste Lösung dafür ist es, eine passende Hyperbel zu wählen:

$$s(a) = \frac{1}{1 + \frac{a}{c}} \quad (4.9)$$

Die geneigte Leserin rechne nach, dass die Forderungen tatsächlich erfüllt sind. Die Gl. (4.8) und (4.9) zusammen ergeben das *Hassell-Modell* [48]. Wir rechnen nun nach, welche Prognosen dieses Modell über die langfristige Entwicklung der Insektenpopulation in Abhängigkeit von den beiden in das Modell eingehenden Parametern r und c trifft. Wir werden dazu das Modell entdimensionalisieren sowie die stationären Punkte und ihren Charakter bestimmen.

Entdimensionalisierung Mit $a_n = \bar{a}x_n$ erhalten wir die dimensionslose Gleichung

$$x_{n+1} = \frac{rx_n}{1 + \frac{\bar{a}}{c}x_n} = \frac{rx_n}{1 + x_n}$$

mit der naheliegenden Skalenwahl $\bar{a} = c$.

Stationäre Punkte, Stabilität und Langzeitverhalten Die Gleichung $x = f(x) = \frac{rx}{1+x}$ hat die beiden Lösungen $x_0^* = 0$ und $x_1^* = r - 1$. Zur Bestimmung der Stabilität dieser beiden Punkte berechnen wir die Ableitung der Funktion f , $f'(x) = \frac{r}{(1+x)^2}$ und setzen die beiden stationären Punkte ein: $f'(x_0^*) = f'(0) = r$ und $f'(x_1^*) = f'(r-1) = 1/r$. Aus dem Satz über die linearisierte Stabilität, Satz 4.1, können wir nun folgern, dass es im Wesentlichen die beiden Fälle $0 < r < 1$ und $1 < r$ zu unterscheiden gilt.

Im ersten Fall, $0 < r < 1$, ist x_0^* ein monoton asymptotisch stabiler stationärer Punkt und x_1^* ist negativ (spielt also im Kontext keine Rolle) und instabil. Ein Cobweb gibt Aufschluss über das globale Verhalten des Systems in diesem Fall, siehe Abb. 4.5: Egal mit welchem Anfangsdatum wir starten, die Lösungsfolge konvergiert monoton gegen 0. Im Kontext bedeutet dies, dass die Population aussterben wird. Eine Pro-Kopf-Reproduktionsrate PKR kleiner eins führt also zu keiner überlebensfähigen Population. Es sollte also gar keine Insektenpopulationen geben, deren PKR so klein ist.

Im zweiten Fall, $1 < r$, ist x_0^* instabil und x_1^* monoton asymptotisch stabil. Wir erwarten also, dass sich die Populationsgröße auf lange Sicht unabhängig vom Anfangsdatum im stationären Punkt x_1^* stabilisieren wird. Das können wir durch ein Cobweb bestätigen. Für

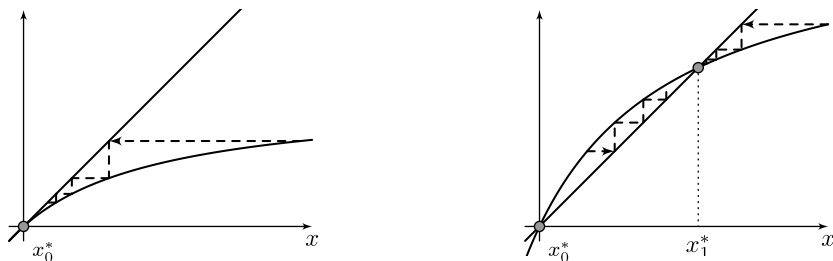


Abb. 4.5 Cobweb für das Hassell-Modell mit exakter Kompensation für $r = 0,9$ und $r = 2,1$

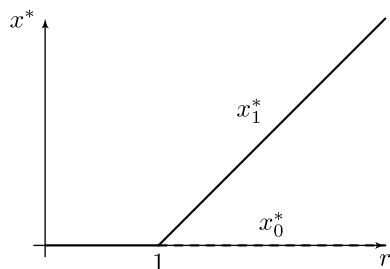
das qualitative Verhalten der Population spielt also die PKR die Hauptrolle, in die auf lange Sicht zu erwartende Populationsgröße $a_1^* = \bar{a}x_1^* = c(r - 1)$ geht auch der Parameter c ein.

Die Gesamtdynamik des Hassell-Modells fassen wir wieder in einem Bifurkationsdiagramm zusammen, vgl. Abb. 4.6: ein sehr überschaubares Bild, verglichen mit der reichhaltigen Dynamik der logistischen Differenzgleichung.

4.4.3 Modelle mit Über- und Unterkompensation

Das Hassell-Modell aus (4.9) lässt nur ein sehr reguläres Verhalten für Insektenpopulationen zu, tatsächlich sind sie schon ausgestorben, $r < 1$, oder die Population hat über Generationen hinweg eine ähnliche Größe, die lediglich den Schwankungen unterworfen ist, welche die sich ändernde Umwelt ihnen aufprägt. Tatsächlich gibt es aber Insektenarten mit nichtüberlappender Generationenfolge, deren Größe sich periodisch zu ändern scheint bzw. – zumindest in Laborversuchen – sogar chaotisch zu sein scheint. Berühmt ist in dieser Hinsicht das Experiment von Nicholson an *blowflies* unter unterschiedlichen Fütterungsbedingungen, siehe [67]. Es besteht also die Notwendigkeit das Hassell-Modell so zu verallgemeinern, dass es periodisches und chaotisches Verhalten zulässt. Mathematisch gesprochen müssen wir einen zusätzlichen Parameter einführen, so dass ein stationärer Punkt sich von einem stabilen stationären Zustand zu einem oszillierend instabilen Zustand

Abb. 4.6 Bifurkationsdiagramm für das Modell mit exakter Kompensation



verändern kann, wenn an diesem Parameter gedreht wird. Ebenfalls von Hassell wurde die folgende Erweiterung studiert:

$$s(a_n) = \frac{1}{\left(1 + \frac{a_n}{c}\right)^b} \quad (4.10)$$

mit einer charakteristischen Anzahl an Individuen c und einer reellen Zahl $b > 0$. Die zugehörige entdimensionalisierte Version lautet

$$x_{n+1} = \frac{r x_n}{(1 + x_n)^b}. \quad (4.11)$$

Man spricht von Modellen mit *Unterkompensation*, falls $0 < b < 1$ ist, von Modellen mit *exakter Kompensation*, falls $b = 1$ ist und von Modellen mit *Überkompensation*, falls $b > 1$ ist. Das Modell mit exakter Kompensation kennen wir bereits. Der Ausdruck Überkompensation rührt daher, dass die zunehmende Konkurrenz unter den Insekten eine Vergrößerung der Population überkompensieren würde, eine Vergrößerung der Population in der n -ten Generation also zu einer Verkleinerung der Population in der $n + 1$ -ten Generation führen kann.

Wir lösen die Gl. (4.11) nach x auf und erhalten:

$$\begin{aligned} x^* = \frac{r x^*}{(1 + x)^b} &\iff x^* = 0 \quad \vee \quad (1 + x^*)^b = r \\ &\iff x^* = 0 \quad \vee \quad x^* = r^{1/b} - 1 \end{aligned}$$

Somit haben wir die stationären Punkte

$$x_0^* = 0, \quad x_1^* = r^{1/b} - 1.$$

Um die Stabilität zu klären, berechnen wir zunächst $f'(x)$:

$$f'(x) = r \left[\frac{1}{(1 + x)^b} - \frac{bx}{(1 + x)^{b+1}} \right] \quad (4.12)$$

Wir setzen x_0^* in die Ableitung ein und erhalten

$$f'(0) = r.$$

Damit ist x_0^* monoton stabil für $0 < r < 1$ und monoton instabil für $1 < r$. Der Parameter b spielt also bei diesem stationären Punkt keine Rolle.

Nun zu x_1^* :

$$\begin{aligned} f'(x_1^*) &= r \left[\frac{1}{r^{\frac{b}{b}} - 1} - \frac{b \cdot (r^{1/b} - 1)}{r^{\frac{b+1}{b}} - 1} \right] \\ &= 1 - b \cdot r^{1+1/b-1-1/b} + b \cdot r^{1-1-1/b} \\ &= 1 - b + \frac{b}{r^{1/b}} \end{aligned}$$

Damit ist x_1^*

- monoton instabil, falls $1 < 1 - b + \frac{b}{r^{1/b}}$,
- monoton asymptotisch stabil, falls $0 < 1 - b + \frac{b}{r^{1/b}} < 1$,
- oszillatorisch asymptotisch stabil, falls $-1 < 1 - b + \frac{b}{r^{1/b}} < 0$,
- und oszillatorisch instabil, falls $1 - b + \frac{b}{r^{1/b}} < -1$.

Um einen Überblick zu erhalten, welches Verhalten für welche Parameter auftritt, stellen wir r als Funktion von b dar, wenn das Stabilitätsverhalten

- von monoton instabil auf monoton stabil,
- von monoton stabil auf oszillierend stabil,
- von oszillierend stabil auf oszillierend instabil

wechselt. Für den ersten Übergang (monoton instabil zu monoton stabil) müssen wir die Gleichung

$$1 - b + \frac{b}{r^{1/b}} = 1 \quad (4.13)$$

nach r auflösen und erhalten $r = 1$. Aus (4.13) sieht man, dass $1 - b + \frac{b}{r^{1/b}} < 1$, falls $r > 1$ ist, und $1 - b + \frac{b}{r^{1/b}} > 1$, falls $r < 1$ ist. Damit ist x_1^* monoton instabil, falls $r < 1$, unabhängig von b .

Nun untersuchen wir den zweiten Übergang (monoton nach oszillatorisch stabil): Dazu müssen wir die Gleichung

$$1 - b + \frac{b}{r^{1/b}} = 0 \quad (4.14)$$

nach r auflösen und erhalten nach ein paar Umformungen

$$\left(1 + \frac{1}{b-1}\right)^b = r.$$

Diese Gleichung ist nur für $b > 1$ lösbar. Für den letzten Übergang (oszillatorisch stabil nach oszillatorisch instabil) lösen wir

$$1 - b + \frac{b}{r^{1/b}} = -1 \quad (4.15)$$

nach r auf und erhalten

$$\left(1 + \frac{2}{b-2}\right)^b = r.$$

Diese Gleichung ist nur für $b > 2$ lösbar. Das bedeutet für $r > 1$: Bei exakter und Unterkompensation ist x_1^* stabil, unabhängig von r , die Populationsgröße befindet sich also (zumindest nach etwas Warterei) nahezu im stationären Punkt. Periodisches oder gar chaotisches Ver-

halten kann nur im überkompensatorischen Bereich auftreten. In Abb. 4.7 sind die Graphen zu

$$\ln(r) = b \cdot \ln\left(1 + \frac{1}{b-1}\right)$$

$$\ln(r) = b \cdot \ln\left(1 + \frac{2}{b-2}\right)$$

geplottet.

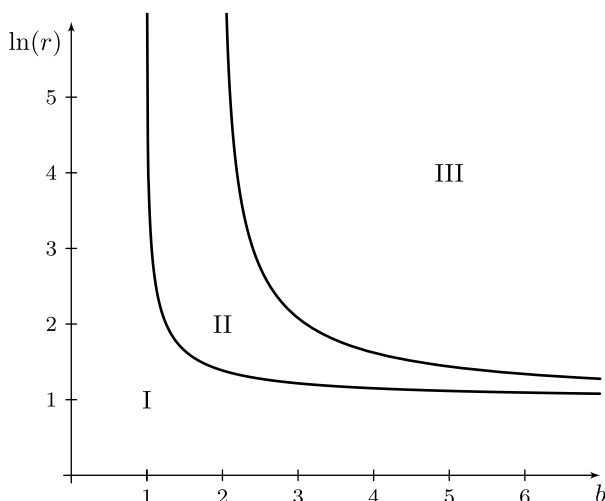
Im Bereich I verhält sich die Population monoton stabil, im Bereich II oszillierend stabil und im Bereich III oszillierend instabil. Trotz einiger spektakulärer Beispiele wie z. B. dem *Colorado beetle* und den angesprochenen *blowflies* fristen die meisten realen Insektenpopulationen ihr Dasein im Bereich I.

4.4.4 Eine Bekämpfungsstrategie: sterile Insekten

In vielen Situation, vor allem in der Land- oder Forstwirtschaft, werden Insektenpopulationen vom Menschen als unerwünscht wahrgenommen und sollen bekämpft oder zumindest in ihrer Größe kontrolliert werden. Eine eingesetzte Möglichkeit zur Bekämpfung besteht im Aussetzen von sterilen Insekten der gleichen Art. Auf diese Weise soll der Fortpflanzungserfolg vermindert und dadurch die Populationsgröße gesenkt werden; man spricht von der *Sterile Insect Release Method* (SIRM).

Wir wollen nun versuchen, ein mathematisches Modell aufzustellen, das diese Bekämpfungsstrategie beschreibt, um aus der mathematischen Modellierung Schlussfolgerungen

Abb. 4.7 Parameterraum für das verallgemeinerte Hassell-Modell (4.11). Bereich I: monoton stabil, Bereich II: oszillierend stabil, Bereich III: oszillierend instabil



für den Einsatz dieser Methode zu ziehen. Dazu nehmen wir an, dass in jeder Generation K sterile Insekten ausgesetzt werden und die Insekten sich zufällig paaren. Für die Wahrscheinlichkeit, dass ein fertiles Insekt einen fertilen Partner findet, setzen wir die Laplace-Wahrscheinlichkeit, also den Quotienten aus der Anzahl der günstigen Möglichkeiten und der Anzahl aller Möglichkeiten an: $p = a_n/(a_n + K)$. Genau genommen müssten wir a_n durch $a_n - 1$ ersetzen, da das Insekt sich ja nicht mit sich selbst paaren kann. Wenn die Population aber hinreichend groß ist, sollte das keinen relevanten Unterschied machen. Ferner haben wir hier die Annahme getroffen, dass zum Zeitpunkt der Paarung noch alle geschlüpften Insekten vorhanden sind. Ansonsten hätten wir a_n durch eine reduzierte Anzahl ersetzen müssen. In Aufg. 4.10 soll untersucht werden, welche Auswirkungen das hätte. Insgesamt erhalten wir aufbauend auf dem Modell der exakten Kompensation das Modell

$$a_{n+1} = r a_n \cdot \frac{a_n}{a_n + K} \cdot \frac{1}{1 + \frac{a_n}{c}} \quad \text{für alle } n \in \mathbb{N}.$$

In einem ersten Schritt entdimensionalisieren wir auch dieses Modell mit der alten Skala $\bar{a} = c$ und erhalten mit $k = K/c$

$$x_{n+1} = f(x_n) = r \cdot x_n \cdot \frac{x_n}{x_n + k} \cdot \frac{1}{1 + x_n} \quad \text{für alle } n \in \mathbb{N}. \quad (4.16)$$

Die Fragestellung lautet nun, wie groß k gewählt werden sollte, damit dieses Modell keinen positiven stationären Punkt mehr besitzt. Dabei nehmen wir die PKR r als fix an und setzen $r > 1$ voraus. Andernfalls gäbe es gar nicht die Notwendigkeit, diese Population zu bekämpfen. Wir lösen zunächst einmal die Fixpunktgleichung $x = f(x)$ und erhalten

$$x = f(x) \Leftrightarrow x = 0 \quad \vee \quad x^2 + (k + 1 - r)x + k = 0 \quad (4.17)$$

mit den Lösungen

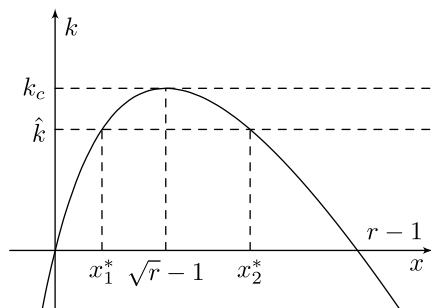
$$x_0^* = 0 \quad \text{und} \quad x_{1/2}^* = -\frac{k + 1 - r}{2} \pm \frac{1}{2} \sqrt{(k + 1 - r)^2 - 4k}. \quad (4.18)$$

Wir könnten jetzt diese Lösungen in ihren Abhängigkeiten von k und r analysieren, gehen aber anders vor: Anstatt zu vorgegebenem k die stationären Punkte $x_{1/2}(k)$ zu berechnen, fragen wir, wie zu vorgegebenem $0 < x < \infty$ das k bestimmt werden muss, damit das gewählte x zu einem stationären Punkt wird. Anders gesagt: Anstatt die zweite Gleichung in (4.17) nach x aufzulösen, lösen wir sie nach k auf und erhalten

$$k = k(x) = \frac{rx}{1 + x} - x.$$

Eine Kurvendiskussion der Funktion $x \mapsto k(x)$ über dem Intervall $[0, \infty)$ zeigt, dass k streng monoton wachsend über dem Intervall $[0, \sqrt{r} - 1]$ ist und dort alle Werte aus dem Intervall $[0, k_c]$ mit $k_c = (\sqrt{r} - 1)^2$ annimmt, außerdem auf dem Intervall $[\sqrt{r} - 1, \infty)$ streng monoton fallend ist und eine Nullstelle in $r - 1$ hat (siehe Abb. 4.8). Das heißt: Für

Abb. 4.8 Graph von k : Für $\hat{k} > k_c$ gibt es keine positiven stationären Punkte, für $0 < \hat{k} < k_c$ gibt es zwei stationäre Punkte $x_1^* = x_1^*(\hat{k})$ und $x_2^* = x_2^*(\hat{k})$ mit $0 < x_1^* < \sqrt{r} - 1 < x_2^* < r - 1$



$k > k_c$ gibt es keine Lösung $x \geq 0$ der Gleichung $x^2 + (k+1-r)x + k = 0$, es gibt also nur den einzigen nicht-negativen stationären Punkt $x_0^* = 0$. Für $k = k_c$ gibt es genau einen positiven stationären Punkt $x^* = \sqrt{r} - 1$, der sich für $0 < k < k_c$ in die beiden stationären Punkte x_1^* und x_2^* mit $0 < x_1^* < \sqrt{r} - 1 < x_2^* < r - 1$ aufspaltet.

Wie sieht es nun mit der Stabilität dieser Punkte aus? Wir könnten natürlich versuchen, die explizite Darstellung der stationären Punkte aus (4.18) in die Ableitung der Funktion f aus (4.16) einzusetzen, und dann mit Hilfe des Satzes über die linearisierte Stabilität auf die Stabilität der Punkte schließen. Stattdessen können wir die Stabilität aber auch aufgrund einer globalen Analyse bestimmen, siehe Abb. 4.9.

1. Fall, $k > k_c$: Da es keinen positiven stationären Punkt gibt und $\lim_{x \rightarrow \infty} f(x) = r$ gilt, haben wir $f(x) < x$ für alle $0 < x$. Damit ist jede Lösungsfolge von (4.16) streng monoton fallend und nach unten durch $x_0^* = 0$ beschränkt. Also konvergiert jede Lösungsfolge und der Grenzwert muss ein stationärer Punkt von (4.16) sein, somit konvergiert also jede Lösungsfolge gegen x_0^* .

Der Grenzfall, $k = k_c$: Es gibt den einen positiven stationären Punkt $x^* = \sqrt{r} - 1$. Wegen $f(0) = 0$, $f'(0) = 0$ und $\lim_{x \rightarrow \infty} f(x) = r$ gilt $f(x) < x$ für alle $x > 0$ und $x \neq x^*$. Damit sind alle Lösungsfolgen von (4.16) mit Anfangsdatum $x_0 > x^*$ streng monoton

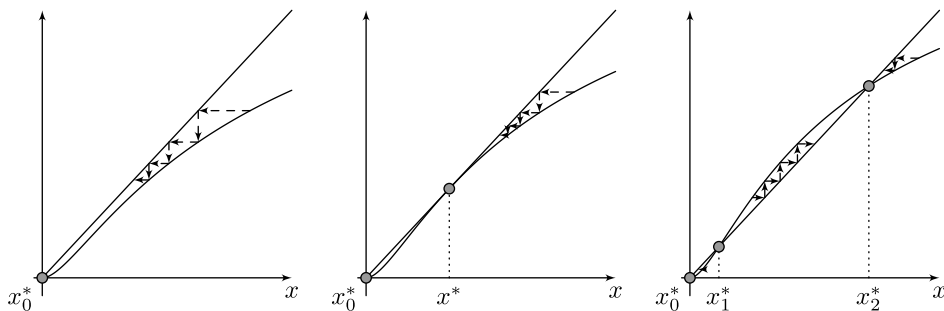


Abb. 4.9 Cobweb für (4.16) mit $r = 9$, $k = 5$, $k = 4$ und $k = 3$

fallend und konvergieren gegen x^* ; alle Lösungsfolgen mit Anfangsdaten $0 < x_0 < x^*$ sind ebenfalls streng monoton fallend und konvergieren gegen x_0^* .

2. Fall, $k > k_c$: Wir haben die beiden positiven stationären Punkte $x_1^* < \sqrt{r} - 1 < x_2^*$. Um zu entscheiden, ob der Graph von f den Graphen der Identität an den beiden Stellen lediglich berührt oder schneidet, schätzen wir ab:

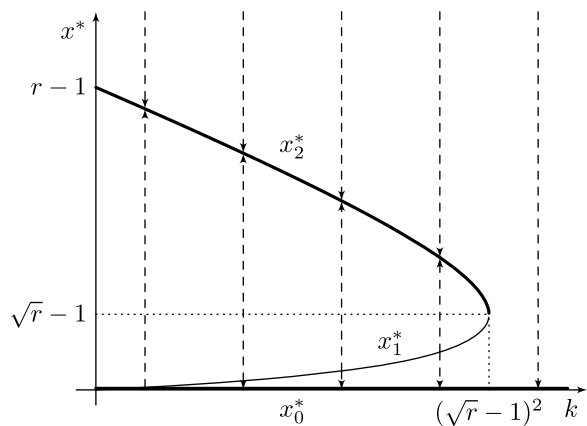
$$\begin{aligned} f(x^*) &= f(\sqrt{r} - 1) = r \cdot \frac{(\sqrt{r} - 1)^2}{\sqrt{r} \cdot (\sqrt{r} + k - 1)} = (\sqrt{r} - 1) \cdot \frac{\sqrt{r}(\sqrt{r} - 1)}{\sqrt{r} + k - 1} \\ &> (\sqrt{r} - 1) \cdot \frac{\sqrt{r}(\sqrt{r} - 1)}{\sqrt{r} + (\sqrt{r} - 1)^2 - 1} = (\sqrt{r} - 1) \cdot \frac{r - \sqrt{r}}{r - \sqrt{r}} \\ &= (\sqrt{r} - 1). \end{aligned}$$

Hierbei haben wir in der Abschätzung von der ersten zur zweiten Zeile benutzt, dass $k > k_c = (\sqrt{r} - 1)$ gilt. Also muss der Graph von f den Graphen der Identität zweimal schneiden und zwar bei x_1^* von unten nach oben und an der Stelle x_2^* von oben nach unten. Damit gilt $f(x) > x$ für $x_1^* < x < x_2^*$ und $f(x) < x$, falls $0 < x < x_1^*$ oder $x_2^* < x$. Das Verhalten der Lösungsfolgen unterscheidet sich also je nach Anfangsdatum: Für Anfangsdaten $0 < x_0 < x_1^*$ konvergieren die Lösungsfolgen monoton fallend gegen x_0^* , für Anfangsdaten $x_1^* < x_0 < x_2^*$ konvergieren die Lösungsfolgen monoton wachsend gegen x_2^* und für Anfangsdaten $x_2^* < x_0$ konvergieren die Lösungsfolgen streng monoton fallend gegen x_2^* .

Es ist übersichtlich, das Lösungsverhalten in Form eines Bifurkationsdiagramms festzuhalten: Dazu tragen wir die stationären Punkte mitsamt ihrem „Einzugsbereich“ gegen den Bifurkationsparameter k auf (siehe Abb. 4.10).

Welche Aussagen liefert die mathematische Analyse des Problems nun für die Problemstellung? In der Ausgangssituation sind keine sterilen Insekten vorhanden, die Insektenpopulation, auch „das System“ genannt, befindet sich (gemessen in Vielfachen der

Abb. 4.10 Die drei stationären Punkte x_0^* , x_1^* und x_2^* bilden jeweils Zweige im Bifurkationsdiagramm. Die stabilen x_0^* - und x_2^* -Zweige sind gefettet dargestellt, der instabile x_1^* -Zweig dünn. Der Einzugsbereich der stabilen stationären Punkte ist gestrichelt angedeutet



gewählten Maßeinheit/Skala $\bar{a} = c$) in dem stabilen Zustand $x_2^* = x_2^*(k = 0) = r - 1$. Nun werden zur Bekämpfung der Insektenpopulation über mehrere Generationen sterile Insekten hinzugefügt. Wird nun die Anzahl K der Insekten so gewählt, dass $k = K/c < k_c$ gilt, wird sich das System über mehrere Generationen der Größe $x_2^*(k)$ annähern; es verbleibt folglich auf dem x_2^* -Ast, die Population wird also mehr oder weniger verkleinert, sie stirbt jedoch nicht aus. Erst wenn der Kontrollparameter k über den Wert k_c vergrößert wird, sodass der stabile Zustand x_2^* nicht mehr existiert, wird der Bestand an Insekten zusammenbrechen und die Population in dem betrachteten Gebiet sogar auf lange Sicht verschwinden; das Ziel ist erreicht.

Was passiert jetzt, wenn sich der Landwirt zurücklehnt und, weil ja keine Schädlingspopulation mehr vorhanden ist und die Herstellung steriler Insekten Geld kostet, das Aussetzen steriler Insekten einstellt? Das System befindet sich dann im stationären Punkt $x_0^* = 0$ beim Kontrollparameter $k = 0$. Für diesen Parameter ist der stationäre Punkt aber instabil, eine noch so kleine Einwanderung von Insekten von außen sorgt dann dafür, dass die Population sich in den stabilen stationären Punkt $x_2^*(0) = r - 1$ entwickelt und der Landwirt sich wieder in der ursprünglichen Situation befindet.

Würde er jedoch weiterhin eine deutlich kleinere Anzahl $\hat{k} < k_c$ steriler Insekten aussetzen, so ließe sich die Insektenpopulation wirkungsvoll kontrollieren: Der stationäre Punkt $x_0^* = x_0^*(\hat{k}) = 0$ ist nämlich jetzt stabil, eine Einwanderung bis zur Größe $x_1^*(\hat{k})$ läge im Einflussbereich des stationären Punkts x_0^* (siehe Abb. 4.10) und die Population würde wieder aussterben.

Wie sieht nun eine sinnvolle Bekämpfungsstrategie aus? In einem ersten Schritt muss eine große Anzahl steriler Insekten ausgesetzt werden, bis die Population auf praktisch null zusammensmilzt. Anschließend sollte prophylaktisch und langfristig eine geringe Anzahl steriler Insekten bereitgestellt werden, um die Einwanderung aus benachbarten Gebieten zu kontrollieren.

Hier tritt also ein bemerkenswerter Effekt auf, der uns noch an mehreren Stellen in diesem Buch begegnen wird: Die Antwort des Systems auf den Kontrollparameter $0 < k < k_c$ kann unterschiedlich sein, x_0^* oder x_2^* , und hängt von der Vorgeschichte ab, man spricht von *Hysterese*.

Aufgaben

4.1 Einzugsbereiche von stationären Punkten

Bei der linearen DZG $a_{n+1} = ra_n + b$ mit $r \neq 1$ ist der stationäre Punkt $a^* = \frac{b}{1-r}$ bekanntlich genau dann asymptotisch stabil, wenn $|r| < 1$ gilt. In diesem Fall konvergiert für jeden Startwert $\bar{a}_0$ die Lösungsfolge (a_n) der Gleichung

$$a_0 = \bar{a}_0, \quad a_{n+1} = ra_n + b, \quad \forall n \in \mathbb{N}_0$$

gegen den stationären Punkt a^* .

Das ist bei Lösungsfolgen der nichtlinearen Gleichungen

$$a_0 = \bar{a}_0, \quad a_{n+1} = f(a_n), \quad \forall n \in \mathbb{N}_0 \quad (4.19)$$

nicht mehr unbedingt der Fall. Untersuchen Sie systematisch anhand von (gezeichneten) Beispielen, wie der *Einzugsbereich* eines asymptotisch stabilen Punkts a^* von f von der Gestalt des Graphen von f abhängt. Versuchen Sie allgemeingültige Sätze zu entdecken und diese zu beweisen. Hierbei ist der Einzugsbereich eines stationären Punkts a^* von f definiert als die Menge aller Startwerte $\bar{a}_0$, für die die Lösungsfolge von (4.19) gegen a^* konvergiert.

4.2 Stabil und asymptotisch stabil

Häufig wird in der Definition von asymptotisch stabilen Punkten zusätzlich gefordert, dass der Punkt stabil sein muss, siehe z. B. [52]. Folgt das nicht schon aus Definition 4.1?

Beweisen oder widerlegen Sie: Es sei $f : \mathbb{R} \rightarrow \mathbb{R}$ eine beliebig oft stetig differenzierbare Funktion und x^* ein stationärer Punkt der Differenzengleichung

$$x_{n+1} = f(x_n) \quad \text{für alle } n \in \mathbb{N}_0.$$

Ist x^* asymptotisch stabil, dann ist x^* auch stabil.

4.3 Die Ausnahmefälle im Satz über die linearisierte Stabilität

- a) Klären Sie das Verhalten von Lösungen der linearen Differenzengleichung

$$x_{n+1} = rx_n + b$$

für $r \in \{-1; 0; 1\}$. Wie sieht es mit stationären Punkten aus? Wie ist es um die Stabilität dieser stationären Punkte bestellt?

- b) Es sei $f : \mathbb{R} \rightarrow \mathbb{R}$ stetig differenzierbar. Zeigen Sie durch Beispiele (durchaus auch in graphischer Darstellung), dass sich für stationäre Punkte x^* von f mit $f'(x^*) \in \{-1; 1\}$ kein allgemeiner Stabilitätssatz formulieren, bzw. im Fall $f'(x^*) = 0$ keine allgemeine Aussage über das Monotonieverhalten von Lösungsfolgen in einer Umgebung von x^* treffen lässt.
- c) Sei nun $f : \mathbb{R} \rightarrow \mathbb{R}$ hinreichend oft stetig differenzierbar, x^* ein stationärer Punkt von f mit $f'(x^*) = 1$. Finden Sie Bedingungen an f , unter denen x^* asymptotisch stabil, stabil, bzw. instabil ist. „Beweisen“ Sie Ihre Behauptung „anschaulich“.

4.4 Möglich oder nicht?

Zeichnen Sie, falls möglich, den Graphen einer stetigen Funktion f , so dass die resultierende Differenzengleichung

$$x_{n+1} = f(x_n)$$

die folgenden Eigenschaften hat, bzw. erläutern Sie, warum es so eine Funktion nicht geben kann.

- Die Folge (x_n) hat genau zwei stationäre Punkte $x_1^* = 0$ und $x_2^* > 0$, und für jedes Anfangsdatum $x_0 > 0$ konvergiert sie gegen x_2^* .
- Die Folge (x_n) hat genau drei stationäre Punkte $0 = x_1^* < x_2^* < x_3^*$, so dass x_1^* monoton asymptotisch stabil ist, x_2^* monoton instabil ist und x_3^* oszillierend instabil ist.
- Die Folge (x_n) hat genau drei stationäre Punkte $0 = x_1^* < x_2^* < x_3^*$, so dass alle drei Punkte asymptotisch stabil sind.
- Die Folge (x_n) hat genau drei stationäre Punkte $0 = x_1^* < x_2^* < x_3^*$, so dass alle drei Punkte instabil sind.

4.5 Weiteres zur logistischen Differenzgleichung

Wir betrachten die logistische Differenzgleichung

$$x_{n+1} = rx_n(1 - x_n).$$

- Nach den analytischen Ergebnissen aus Abschn. 4.3 existieren ab $r = 3$ Perioden der Ordnung 2, die ab $r = 1 + \sqrt{6}$ instabil werden. Bestimmen Sie numerisch mit einem TK möglichst genau das r , ab dem es
 - Perioden der Ordnung 8
 - Perioden der Ordnung 6
 - Perioden der Ordnung 3
 gibt und beschreiben Sie Ihr Vorgehen.
- Im Internet finden sich unter dem Schlagwort „Bifurkationsdiagramm der logistischen Differenzgleichung“ solche Abbildungen wie Abb. 4.3. Unter der in der Abbildungsunterschrift angegebenen Internetadresse wird erklärt, wie die Grafik entstanden ist.

„The horizontal axis is the r parameter, the vertical axis is the x variable. The image was created by forming a 1601×1001 array representing increments of 0.001 in r and x . A starting value of $x = 0.25$ was used, and the map was iterated 1000 times in order to stabilize the values of x . 100,000 x -values were then calculated for each value of r and for each x value, the corresponding (x, r) pixel in the image was incremented by one. All values in a column (corresponding to a particular value of r) were then multiplied by the number of non-zero pixels in that column, in order to even out the intensities. Values above 250,000 were set to 250,000, and then the entire image was normalized to 0-255. Finally, pixels for values of r below 3.57 were darkened to increase visibility.“

- Erläutern Sie, wie diese Graphik produziert wurde.
- Erklären Sie, warum es vor dem Hintergrund des Satzes von Sarkovskii erstaunlich ist, dass im Bereich $r \approx 3,82$ anscheinend nur drei Perioden der Länge 3 vorhanden sind. Wie erklären Sie sich das vor dem Hintergrund der Entstehung der Graphik?
- Erstellen Sie selbst mit einem Programm Ihrer Wahl so eine Graphik.

4.6 Periodische Punkte implizieren stationäre Punkte

In dieser Aufgabe soll eine erste Teilaussage aus dem Satz von Sarkovskii bewiesen werden: Periodische Punkte implizieren stationäre Punkte. Wir wiederholen die grundlegende Definition:

Definition 4.2

Für eine stetige Funktion $f : \mathbb{R} \rightarrow \mathbb{R}$ und $n \in \mathbb{N}_0$ bezeichnen wir mit $f^{[n]}$ die n -fach Iterierte von f , die durch $f^{[0]}(x) = x$ und $f^{[n+1]}(x) = f(f^{[n]}(x))$ für alle $x \in \mathbb{R}$ definiert ist.

Eine Punkt $x^* \in \mathbb{R}$ heißt n -periodischer Punkt von f , wenn $f^{[n]}(x^*) = x^*$ und $f^{[j]}(x^*) \neq x^*$ für alle $j = 1, \dots, n-1$. $\blacklozenge$

Beweisen Sie die folgenden beiden Sätze:

Satz 1 (Fixpunktsatz)

Es sei $f : \mathbb{R} \rightarrow \mathbb{R}$ stetig und $J \subset \mathbb{R}$ ein kompaktes Intervall mit $J \subset f(J)$. Dann gibt es ein $x^* \in J$ mit $f(x^*) = x^*$.

Satz 2 (Periodische Punkte implizieren stationäre Punkte.)

Es sei $f : \mathbb{R} \rightarrow \mathbb{R}$ stetig und es gebe einen n -periodischen Punkt von f für ein $n \geq 1$. Dann hat f auch einen 1-periodischen Punkt, also einen Fixpunkt.

4.7 Das Gompertz-Modell

Ein weiteres Modell zum Populationswachstum ist das sogenannte Gompertz-Modell. Es wird durch die bereits entdimensionalisierte Differenzengleichung

$$x_{n+1} = f(x_n) \quad \text{mit} \quad f(x) = rx \ln\left(\frac{1}{x}\right) \quad \text{für} \quad x > 0 \quad (4.20)$$

beschrieben.

- Zeigen Sie, dass sich die Funktion f durch 0 stetig in $x = 0$ fortsetzen lässt, und nutzen Sie im weiteren Verlauf diese Fortsetzung der Funktion f , die dann auch für $x = 0$ definiert ist.
- Bestimmen Sie die stationären Punkte des Modells sowie die Stabilität des positiven stationären Punktes.
- Skizzieren Sie das Cobweb für $r = 1/2$. Diskutieren Sie dazu den Graphen von f .
- Vergleichen Sie das Gompertz-Modell (4.20) qualitativ mit dem logistischen Modell:

$$x_{n+1} = rx_n(1 - x_n) \quad (4.21)$$

4.8 Ein exponentielles Modell

Eine Population wird (bereits entdimensionalisiert) durch die Differenzengleichung

$$x_{n+1} = x_n e^{\alpha - x_n} = f(x_n) \quad (4.22)$$

mit $\alpha > 0$ beschrieben.

- Berechnen Sie die stationären Punkte und ihre Stabilität in Abhängigkeit von α .
- Diskutieren Sie fachmännisch die Funktion f und zeichnen Sie Cobwebs mit sinnvollen Fallunterscheidungen für die Werte von α . Zeichnen Sie ein Bifurkationsdiagramm bezüglich des Bifurkationsparameters α .
- Das Modell (4.22) soll auch im oszillierend instabilen Bereich stabilisiert werden, indem in jedem Zeitschritt, nachdem alle Geburten und Todesfälle stattgefunden haben, ein Anteil $0 < p < 1$ der Population beseitigt wird, also nur noch der Anteil $q = 1 - p$ übrig bleibt. Das führt auf das Modell

$$x_{n+1} = q x_n e^{\alpha - q x_n} . \quad (4.23)$$

Für welche Werte von q führt das auf asymptotisch stabile positive stationäre Punkte?

- Wir kehren zurück zur Gl. (4.22): Weisen Sie nach (z.B. durch eine geeignete Kurvendiskussion), dass es für $\alpha > 2$ zwei-periodische Lösungen gibt. (Sie müssen nur die Existenz nachweisen, nicht die Lösungen berechnen.)

4.9 Der Pazifische Lachs – eine Population mit zweijährigem Lebenszyklus

Der Pazifische Lachs weist den folgenden zweijährigen Lebensrhythmus auf. Nach dem Schlüpfen in Süßwasserbächen wandern die Junglachse in den Ozean, wachsen dort heran, werden geschlechtsreif und wandern anschließend in ihre Geburtsbäche zurück, wo sie ablaichen und anschließend sterben. Die Lachspopulation besteht also aus zwei Teilpopulationen: der Population der jungen, noch nicht geschlechtsreifen Tiere und der Teilpopulation der erwachsenen geschlechtsreifen Tiere. Wir bezeichnen die Populationsgrößen der Jung- und der Alttiere im n -ten Jahr mit j_n bzw. a_n . Die jährlichen Übergänge von Jung- zu Alttieren bzw. (über Eiablage und Schlüpfen) von Alt- zu Jungtieren modellieren wir mit Hilfe zweier Funktionen f und g als

$$a_{n+1} = f(j_n) \quad \text{und} \quad j_{n+1} = g(a_n).$$

- Begründen Sie, warum die folgenden Modellannahmen für f und g sinnvoll sind:
 - f und g sind monoton wachsende Funktionen mit $f(0) = 0$ und $g(0) = 0$.
 - $f(j) < j$ für alle $j > 0$ und $\lim_{a \rightarrow \infty} g(a) = g_0$.
- Bestimmen Sie Gleichungen, welche die Fixpunkte (j^*, a^*) dieses Systems erfüllen müssen, und interpretieren Sie diese geometrisch in einem j - a -Diagramm.

- c) Experimentieren Sie numerisch und geometrisch und finden Sie Kriterien, ausgedrückt in den Ableitungen $g'(a^*)$ und $(f^{-1})'(a^*)$, ob ein stationärer Punkt (j^*, a^*) asymptotisch stabil oder instabil ist (Sie müssen diese nicht unbedingt beweisen).
- d) Zeigen Sie mit geometrischen Überlegungen, dass im Fall mehrerer stationärer Punkte diese abwechselnd stabil und instabil sind, wobei die Reihenfolge in der naheliegenden Weise definiert ist.
- e) Gegeben sei eine Situation mit drei stationären Punkten $P_0^* = (0, 0)$, $P_1^* = (j_1^*, a_1^*)$ und $P_2^* = (j_2^*, a_2^*)$. Beschreiben Sie die unterschiedlichen möglichen Phänomene für das Langzeitverhalten $n \rightarrow \infty$.

4.10 Unterschiedliche Geschlechter und weitere Bekämpfungsmethoden

- a) Stellen Sie ein Modell für Insektenpopulationen mit nichtüberlappender Generationenfolge auf, welches berücksichtigt, dass es männliche und weibliche Insekten gibt. Untersuchen Sie die Lösungen Ihres Modells so weit wie möglich.
- b) Wir betrachten das *Sterile Insect Release Modell* aus Abschn. 4.4.4. Es erscheint merkwürdig, dass bei der Laplace-Wahrscheinlichkeit für die Paarung mit einem fertilen Partner, $a_n/(a_n + K)$, die Populationsgröße beim Schlüpfen zugrunde gelegt wird. Stellen Sie ein Modell auf, das berücksichtigt, dass die Insekten erst einmal bis zur Paarung überleben müssen, und untersuchen Sie die Lösungen Ihres Modells so weit wie möglich.
- c) Um die Insekten zu bekämpfen, wird nun eine *Absammelstrategie* verfolgt: Jedes Jahr wird zu einem bestimmten Zeitpunkt ein Anteil k der Population abgesammelt. Legen Sie einen Zeitpunkt im Lebenszyklus der Insekten fest, zu dem abgesammelt werden soll, und stellen Sie ein geeignetes Modell auf.
Berechnen Sie, wie groß k mindestens gewählt werden muss, damit die Insektenpopulation gemäß diesem Modell ausstirbt.
Vergleichen Sie Ihre Strategie mit der Strategie sterile Insekten auszusetzen, die in Abschn. 4.4.4 behandelt wurde.



Populationsmodelle und Befischung – Modellieren mit Differentialgleichungen

5

Inhaltsverzeichnis

5.1	Modellieren mit Differentialgleichungen	109
5.2	Qualitatives Lösen von autonomen Differentialgleichungen	115
5.3	Populationsmodelle mit Bejagung.....	117
5.4	Der Tannenwickler und die Balsamtanne – erster Teil	124
5.5	Eine einfache Lösungstheorie für Differentialgleichungen – die Methode von der Trennung der Variablen	132
5.6	Ein paar ausgewählte Lösungsverfahren	140
5.7	Ein einfaches Verfahren zur numerischen Lösung von Anfangswertproblemen	144
5.8	Modellierungen in Abituraufgaben – Blutalkohol	145
	Aufgaben	155

5.1 Modellieren mit Differentialgleichungen

Im zweiten Kapitel haben wir Prozesse untersucht, bei denen es sinnvoll war anzunehmen, dass sich die zu modellierende Größe in bestimmten Abständen sprunghaft änderte (Zinsen und Insektenpopulationen mit nicht-überlappender Generationenfolge); dafür waren Differenzenfolgen besonders geeignet. Wir haben gesehen, dass diese Modelle sowohl stabile als auch chaotisch verlaufende Phänomene beschreiben können. In den folgenden Kapiteln wollen wir uns nun mit Phänomenen beschäftigen, die eine zeitkontinuierliche Beschreibung sinnvoll erscheinen lassen: Die Veränderungen vollziehen sich dabei weniger sprunghaft, sondern eher „gleichmäßig“ in der Zeit und wir werden versuchen, sie durch glatte, das heißt stetig differenzierbare, Funktionen der Zeit zu modellieren.

Wir werden es vielfach mit der Beschreibung von Populationsgrößen bestimmter Lebewesen zu tun haben. Dabei stellt sich die Frage, inwiefern sich denn hier glatte Vorgänge vollziehen: Populationen bestehen ja häufig aus Individuen, und Änderungen der Anzahl der Individuen können ja nicht gleichmäßig vonstattengehen, sondern die Anzahl muss sich

mindestens um eins ändern. Wir betrachten jedoch große Populationen, in denen die Anzahlen im Bereich von Tausenden und Millionen und größer liegen, sodass ein kontinuierlicher Ansatz sinnvoll erscheint, der uns nämlich ermöglicht, die innere Struktur und unser Wissen über stetige und differenzierbare Funktionen zu nutzen. In einem ersten Abschnitt werden wir erläutern, welche Überlegungen zu mathematisch analysierbaren Modellen, nämlich zu Differentialgleichungen führen. Wir wählen dazu das recht zugängliche Problem der Modellierung von Bevölkerungsgrößen.

5.1.1 Das Modell des exponentiellen Wachstums

Wir stellen uns nun die Aufgabe, die zeitliche Entwicklung einer bestimmten Population in einem bestimmten Gebiet mathematisch zu modellieren. Dabei könnte es z. B. um einen Fischbestand in einem bestimmten Bereich eines Ozeans gehen, um Bakterien in einer Petrischale oder auch um die Bevölkerungsentwicklung in Frankreich. Im Gegensatz zu Kap. 4 gehen wir jetzt aber davon aus, dass sich die interessierende Größe durch eine stetig differenzierbare Funktion $P(t)$ beschreiben lässt, wobei $P(t)$ die Größe der Population zur Zeit t angibt. Die Variable t hat dabei immer die Dimension der Zeit. Die abhängige Variable P kann unterschiedliche Dimensionen haben: Häufig wird man Anzahlen von Individuen zählen wollen, die Dimension ist dann also eine Anzahl. Bei Bakterienpopulationen wird aber häufig auch die besiedelte Fläche gemessen, Fischpopulationen (und Fischfänge) misst man in der Regel nach Masse.

Wir wollen uns nun überlegen, durch welche fundamentalen Mechanismen sich die Populationsgröße von einem Zeitpunkt $\hat{t}$ bis zu einem späteren Zeitpunkt t ändern kann. Wir unterscheiden vier Mechanismen, die so fundamental sind, dass sich schwerlich weitere finden lassen: Es können Individuen neu geboren werden oder durch Zellteilung entstehen, es können Individuen absterben, es können Individuen von außen in das betrachtete Gebiet zuwandern und es können Individuen aus dem betrachteten Gebiet abwandern. Wir erhalten damit die Gleichung

$$P(t) - P(\hat{t}) = B(t) - D(t) + I(t) - E(t), \quad (5.1)$$

wobei $B(t)$, $D(t)$, $I(t)$ und $E(t)$ die Größe der im Zeitintervall von $\hat{t}$ bis t neugeborenen, gestorbenen, immigrierten und emigrierten Population angibt.

Wir gehen nun davon aus, dass unsere Glattheitsannahme auch für die Teilprozesse Geborenwerden, Sterben, Einwandern und Auswandern gültig ist. Durch Differentiation von (5.1) erhalten wir die Gleichung

$$\dot{P}(t) = b(t) - d(t) + i(t) - e(t) \quad (5.2)$$

mit den (zeitabhängigen) Geburten-, Sterbe-, Immigrations- und Emigrationsratenfunktionen b , d , i und e . Die Modellbildung besteht nun darin, die Abhängigkeit dieser vier Ratenfunktionen von der Populationsgrößenfunktion P und weiteren Einflüssen zu bestimmen;

die anschließende mathematische Arbeit besteht darin, zu untersuchen, ob die entstehenden Gleichungen Lösungen besitzen und wie diese aussehen.

Wir beginnen mit den denkbar einfachsten Modellen: Dafür blenden wir die Migrationseinflüsse aus: $e(t) = i(t) = 0$ (es würde auch $i(t) - e(t) = 0$ reichen) für alle Zeiten t . Wir gehen davon aus, dass sich die äußeren Umstände, in denen die Population existiert, nicht ändern und die Geburten- und Sterberaten zur Zeit t nur von der Populationsgröße zur Zeit t abhängen. Das bedeutet, dass wir insbesondere von jeglicher Altersstruktur der Bevölkerung absehen.

In dem einfachsten denkbaren Modell besteht ein proportionaler Zusammenhang, $b \sim P$ und $d \sim P$: Ist die Population doppelt so groß, werden auch doppelt so viele geboren und doppelt so viele werden sterben. Wir erhalten also mit zwei positiven Konstanten α und β die Gleichung

$$\dot{P}(t) = \alpha P(t) - \beta P(t) = (\alpha - \beta)P(t) = rP(t). \quad (5.3)$$

Welche Dimensionen haben die hier auftretenden Größen und Konstanten? Wir bezeichnen die Dimension von P mit A wie Anzahl. Zählen wir aber nicht die Individuen der Population, so könnte A jetzt auch für Masse oder Fläche stehen. Die Dimension von $\dot{P}$ ist dann Anzahl pro Zeit, also A/T (siehe auch Abschn. 1.3.3). Damit muss auch die rechte Seite diese Dimension haben. Es folgt, dass die Konstanten α , β und r die Dimension einer inversen Zeit haben, $1/T$. Hat die Größe P die Einheit Anzahl, so bedeutet z.B. eine Konstante $\alpha = 0,001 \text{ d}^{-1}$, dass im Durchschnitt jedes Individuum pro Tag 0,001 Nachkommen zur Welt bringt, oder anders gesagt, dass 1000 Individuen im Schnitt einen Nachkommen pro Tag zur Population hinzufügen. Gl. (5.3) ist eine Gleichung zwischen der Funktion P und ihrer Ableitung $\dot{P}$, wir sprechen von einer Differentialgleichung. Lösungen dieser einfachen Differentialgleichung sind aus der Schule bekannt, es sind die Exponentialfunktionen

$$P(t) = P_0 e^{r(t-t_0)}, \quad (5.4)$$

wobei P_0 gerade die Populationsgröße zur Zeit t_0 ist, $P(t_0) = P_0$. Die Differentialgleichung (5.3) heißt deshalb auch die Differentialgleichung des *exponentiellen Wachstums*. Falls $r < 0$ ist, beschreibt sie allerdings einen Zerfallsvorgang. Aus Modellierungsgründen ist es nun wichtig, dass es keine weitere Funktion $\tilde{P}$ gibt, die die Differentialgleichung (5.3) löst und zur Zeit t_0 den Wert P_0 hat, sonst wüsste man ja nicht, welche der verschiedenen Funktionen man zur Beschreibung des Vorgangs heranziehen sollte. Eine einfache Rechnung zeigt nun, dass das sogenannte *Anfangswertproblem* (AWP)

$$\dot{P}(t) = rP(t) \quad \text{und} \quad P(t_0) = P_0 \quad (5.5)$$

nur die eine Lösung (5.4) hat: Angenommen $\tilde{P}$ sei eine zweite Lösung von (5.5). Falls $P_0 \neq 0$, betrachten wir den Quotienten

$$Q(t) = \frac{\tilde{P}(t)}{P(t)}.$$

Es gilt $Q(t_0) = 1$, sowie

$$Q' = \frac{\tilde{P}'P - \tilde{P}P'}{P^2} = r \frac{\tilde{P}P - \tilde{P}P}{P^2} = 0$$

und damit $\tilde{P} = P$.

Sei nun $P_0 = 0$ (und damit $P(t_0) = 0$), $\tilde{P}$ ebenfalls eine Lösung von (5.5) und $\tilde{P}(\hat{t}) = \hat{P}_0 \neq 0$ für irgendein $\hat{t} \neq 0$. Dann gilt aber nach dem eben Gesagten $\tilde{P}(t) = \hat{P}_0 e^{t-\hat{t}}$ und damit insbesondere $\tilde{P}(0) = \hat{P}_0 e^{-\hat{t}} \neq 0$ im Widerspruch zur Annahme.

Die Lösungen der Differentialgleichung (5.3) wachsen entweder exponentiell gegen $\pm\infty$, falls $r > 0$, die Geburtenrate also die Sterberate überwiegt, oder fallen exponentiell gegen 0 für $r < 0$. Damit ist klar, dass durch die Differentialgleichung des exponentiellen Wachstums für $r > 0$ die zeitliche Entwicklung einer bestimmten Population nur über einen begrenzten Zeitraum beschrieben werden kann, da die Population ja sonst beliebig groß werden würde, was offensichtlich nicht sein kann.

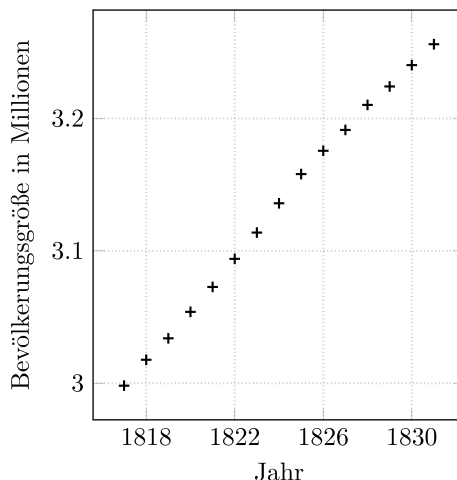
5.1.2 Das logistische Modell

Während der englische Ökonom Thomas Malthus (1766–1834) davon ausging, dass sich im Bevölkerungswachstum von Staaten wie dem Vereinigten Königreich regelmäßige Abschnitte mit exponentiellem (er spricht von geometrischem) Wachstum mit unregelmäßigen Abschnitten katastrophaler Zusammenbrüche infolge von Hungersnöten und Seuchen abwechseln würden, siehe [61], kam der belgische Mathematiker Pierre-François Verhulst (1804–1849) zu anderen Schlüssen, die wir hier kurz darstellen wollen. Er untersuchte die Bevölkerungsentwicklung mehrerer europäischer Staaten, u. a. Frankreich und Belgien, und fand in den Daten Phänomene, die sich mit der Malthus'schen Theorie nur unzureichend erklären ließen. In Abb. 5.1 sind die von Verhulst erhobenen Daten zu Frankreich dargestellt. Offensichtlich handelt es sich um einen Zeitraum ohne katastrophale Zusammenbrüche, andererseits ist aber auch deutlich erkennbar, dass es sich nicht um ein exponentielles Wachstum handeln kann, weil die Wachstumsgeschwindigkeit zunächst zuzunehmen scheint, aber ca. ab dem Jahr 1824 eher abnehmend wirkt. Verhulst schlägt nun in seiner Schrift *Notice sur la loi que la population poursuit dans son accroissement* [88] aus dem Jahr 1838 vor, die zugrunde liegende Differentialgleichung $\dot{P} = rP$ durch einen negativen Summanden $-\varphi(P)$ zu modifizieren, und schlägt als *mathematisch einfachste* Möglichkeit einen quadratischen Term in P vor:

$$\dot{P} = rP - nP^2 \tag{5.6}$$

Die resultierende Differentialgleichung nennt er *logistische Differentialgleichung*. Verhulst berechnet die Lösungen dieser Differentialgleichung (wir lernen in Abschn. 5.5, wie das geht) und zeigt, dass die erhobenen Daten durch Lösungen dieser Differentialgleichungen

Abb. 5.1 Die Daten von Verhulst zur Bevölkerungsentwicklung in Frankreich



gut beschrieben werden können. Damit hat sich die logistische Differentialgleichung als ein Grundmodell der Beschreibung von Wachstumsvorgängen von Populationsgrößen etabliert.

Bevor wir zum Lösen der Differentialgleichung kommen, soll noch ein Blick auf die Modellierung geworfen werden. Verhulsts Argument für den Term $-nP^2$ war rein mathematisch: Ein quadratischer Term ist nach den linearen Termen der einfachste. Lässt sich der quadratische Term auch inhaltlich begründen? Kann man Mechanismen annehmen, die innerhalb einer Population wirken und einen quadratischen Term plausibel erscheinen lassen?

Der quadratische Term als Konkurrenzeffekt

Die Abweichung vom exponentiellen Wachstum lässt sich durch Konkurrenz unter den Mitgliedern der Population erklären: Sie konkurrieren um Ressourcen wie Platz und Nahrung, im Falle von menschlichen Populationen vielleicht auch um Arbeitsplätze und Ähnliches. Konkurrenz taucht im exponentiellen Modell nicht auf, die Parameter α und β ändern sich nicht mit wachsender Populationsgröße. Wir geben zwei Überlegungen an, die beide zu derselben Differentialgleichung führen.

Beim ersten Zugang überlegt man sich, wann Konkurrenz entsteht. Diese entsteht, wenn mehrere Individuen der Population in irgendeiner Form aufeinandertreffen. Wenn man keinen trifft, kann man auch mit keinem konkurrieren. Wie lässt sich nun die Häufigkeit solcher potentiell Konkurrenz erzeugenden Zusammentreffen mathematisch modellieren? Dazu stellen wir uns sehr vereinfacht vor, die Mitglieder der Population würden sich zufällig in ihrem Populationsgebiet bewegen. Aus der Sicht eines einzelnen Individuums der Population ist dann die Wahrscheinlichkeit, in einem bestimmten Zeitraum auf ein anderes Mitglied der Population zu treffen, proportional zur Größe der Population: Eine doppelt so große Population bedeutet doppelt so viele Zusammentreffen mit anderen für ein einzel-

nes Individuum. Wir gehen vereinfacht davon aus, dass diese Anzahl für jedes Individuum dieselbe ist. Um die Gesamtzahl der Zusammenstöße von zwei Individuen in der gesamten Population zu ermitteln, müssen wir nun die Zahl der Zusammenstöße eines Individuums mit der Gesamtzahl der Individuen multiplizieren und durch zwei teilen, weil jeder Zusammenstoß doppelt gezählt wurde. Wir erhalten also, dass die Anzahl der Zweierzusammenstöße proportional zum Quadrat der Populationsgröße ist.

Da diese Zusammentreffen sich aufgrund der Konkurrenzeffekte überwiegend negativ auf die Populationsgröße auswirken, bekommt dieser Term in der Differentialgleichung ein negatives Vorzeichen. Welche inhaltliche Bedeutung hat nun der Proportionalitätsfaktor n ? Er hat die Dimension $1/(TA)$. Ein Wert von z. B. $n = \frac{0,1}{a} \cdot \frac{1}{1000}$ bedeutet, dass von zehn Individuen pro Jahr und pro Tausend konkurrierender Individuen eines aufgrund von Konkurrenzeffekten sterben würde. Hätten wir z. B. eine Population, die durch einen Deus ex Machina von außen durch passendes Hinzufügen oder Wegnehmen von Individuen künstlich auf einer Größe von konstant 500 Exemplaren gehalten würde, müssten wir davon ausgehen, dass in einem Jahr $\frac{0,1}{1000} \cdot 500^2 = 25$ Exemplare aufgrund von Konkurrenzeffekten sterben würden. Der Parameter n steuert also die Intensität der Konkurrenz.

Der quadratische Term als variabler Pro-Kopf-Reproduktionskoeffizient

Eine andere Überlegung führt zum gleichen Ergebnis. Wir starten hier wieder mit dem exponentiellen Modell $\dot{P} = rP$, nehmen aber an, dass $r = \alpha - \beta$ nicht konstant ist, sondern wiederum von der Populationsgröße abhängt, also $r = r(P)$. Insbesondere ist es plausibel anzunehmen, dass für sehr kleine Populationsgrößen $r = r_0$ positiv ist und ab einer bestimmten Populationsgröße K , die von dem betrachteten Gebiet abhängt, negativ wird. Das einfachste mathematische Modell für so ein Verhalten von r ist linear:

$$r(P) = r_0 \cdot \left(1 - \frac{P}{K}\right)$$

Insgesamt erhalten wir also das Modell

$$\dot{P} = r(P) \cdot P = r_0 \left(1 - \frac{P}{K}\right) \cdot P. \quad (5.7)$$

Die Parameter r_0 und K haben die Bedeutung einer Pro-Kopf-Reproduktionsrate für sehr kleine Populationen (Dimension $1/T$) bzw. einer Kapazität des Lebensraums (Dimension A). Dabei ist die Kapazität gerade die Populationsgröße, die der Lebensraum auf lange Sicht tragen kann. Wir werden gleich sehen, dass Lösungen dieser Differentialgleichung mit positiven Anfangswerten auf lange Sicht gegen K konvergieren.

Natürlich handelt es sich bei den Differentialgleichungen (5.6) und (5.7) um dieselbe Differentialgleichung und die Parameter stehen in der Relation $r = r_0$ und $n = r_0/K$.

Für die logistische Differentialgleichung dürften in der Regel die Lösungen nicht aus der Schule bekannt sein. Es ist vielleicht erst einmal sogar zweifelhaft, ob diese Differentialgleichung für alle Anfangswerte Lösungen hat und ob diese eindeutig sind. Wir werden

im weiteren Verlauf zweischrittig vorgehen: Wir werden in Abschn. 5.2 bis 5.4 voraussetzen, dass Anfangswertprobleme zu den dort betrachteten Differentialgleichungen stets eine eindeutige Lösung haben, und wollen mit elementaren Überlegungen die wichtigsten qualitativen Eigenschaften dieser Lösungen wie Monotonie und das sogenannte Langzeitverhalten ermitteln, ohne Lösungen explizit, d. h. in Form von Formeln, zu bestimmen. Mit diesen einfachen mathematischen Mitteln werden wir dann in den Abschn. 5.3 und 5.4 auch komplexere Modellierungen vornehmen und Rückschlüsse aus den mathematischen Modellen auf die beschriebene Realsituation ziehen.

In Abschn. 5.5 werden wir den mit Schulmitteln durchführbaren Existenz- und Eindeutigkeitssatz mit Hilfe der Trennung der Variablen vorstellen, der in einigen Situationen auch ein Rezept angibt, wie explizite Lösungen von Differentialgleichungen berechnet werden können. Mit Hilfe dieser expliziten Lösungsmethoden werden wir uns in Abschn. 5.8 einigen Kontexten aus Abituraufgaben zuwenden.

5.2 Qualitatives Lösen von autonomen Differentialgleichungen

In diesem Abschnitt orientieren wir uns an der Darstellung in [6]. Der folgende Sachverhalt liegt allen anschließenden Überlegungen zugrunde: Wenn die *Existenz und Eindeutigkeit* der Lösungen des Anfangswertproblems

$$\dot{x} = f(x) \quad \text{und} \quad x(t_0) = x_0 \tag{5.8}$$

für eine stetig differenzierbare Funktion f (der Einfachheit halber nehmen wir als Definitionsbereich die gesamten reellen Zahlen an) akzeptiert wird, lässt sich die Lösungsgesamtheit der Differentialgleichung allein aus der Kenntnis des Graphen von f mit einem Schlag qualitativ skizzieren. Zur Illustration wählen wir eine kubische Funktion (siehe Abb. 5.2). Zunächst bemerken wir, dass sich unterschiedliche Lösungsgraphen nicht schneiden dürfen, da sonst die Eindeutigkeit des Anfangswertproblems für den Schnittpunkt verletzt ist. Eine besondere Rolle spielen die Nullstellen x_0^* , x_1^* und x_2^* von f . Wählen wir sie als Anfangswert, so führen sie zu konstanten Lösungen. Lösungen mit Anfangswerten zwischen x_0^* und x_1^* müssen also zwischen x_0^* und x_1^* bleiben. Sie sind streng monoton fallend, da f in diesem Bereich negativ ist, und konvergieren für $t \rightarrow -\infty$ gegen x_1^* und für $t \rightarrow \infty$ gegen x_0^* . Außerdem gehen alle Lösungen zwischen x_0^* und x_1^* durch Verschiebung (also durch Zeittranslation) auseinander hervor. Entsprechend verhält es sich mit Lösungen zu Anfangsdaten zwischen x_1^* und x_2^* . Lösungen zu Anfangsdaten größer als x_2^* konvergieren monoton fallend gegen x_2^* und Lösungen zu Anfangsdaten kleiner als x_0^* konvergieren monoton wachsend gegen x_0^* . In Abschn. 5.5 werden wir uns sowohl um eine fachmathematische als auch um eine präformale Begründung dieser Aussagen kümmern.

Allgemein können wir also sagen: Nicht-konstante Lösungen von (5.8) entwickeln sich streng monoton von einer Nullstelle von f zur nächsten bzw. von $\pm\infty$ zu der größten/kleinsten Nullstelle von f bzw. von der größten/kleinsten Nullstelle von f nach $\pm\infty$.

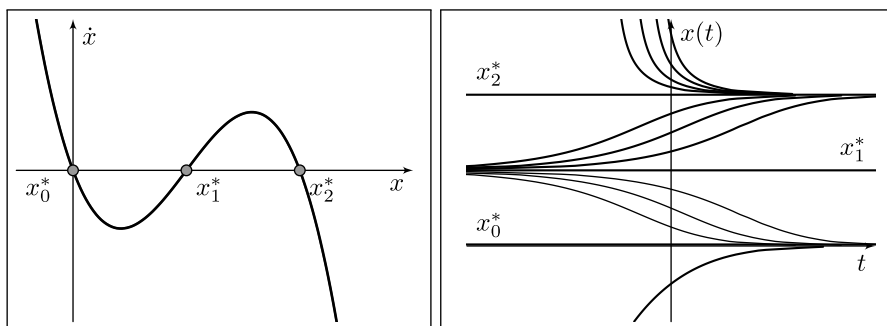
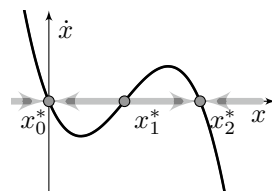


Abb. 5.2 Graph von f und Graphen einiger Lösungen von $\dot{x} = f(x)$

Ob eine Nullstelle von f , also ein stationärer Punkt, „anziehend“ oder „abstoßend“ ist, lässt sich am Vorzeichenwechsel von f erkennen: Hat die Funktion an der Nullstelle einen Vorzeichenwechsel von $+$ nach $-$, so ist sie anziehend, hat sie einen Vorzeichenwechsel von $-$ nach $+$, so ist sie abstoßend (siehe Abb. 5.3). Nullstellen ohne Vorzeichenwechsel ziehen von der einen Seite an und stoßen von der anderen Seite ab. Fachsprachlich handelt es sich um asymptotisch stabile bzw. um instabile stationäre Punkte.

In Anwendungssituationen ist häufig vor allem dieses Langzeitverhalten der Lösungen von (5.8) interessant. In welchem Zustand befindet sich die modellierte Größe, wenn sie sich eine Zeit lang ungestört entwickeln konnte? Diese Frage lässt sich durch einen Blick auf den Graphen von f beantworten (siehe Abb. 5.3): In unserem Beispiel konvergieren Lösungen zu Anfangswerten größer als x_1^* gegen x_2^* , Lösungen zu Anfangsdaten kleiner als x_1^* konvergieren gegen x_0^* . Die Größe befindet sich also, je nachdem, wo sie gestartet ist, nach einiger Zeit entweder im Zustand x_2^* oder im Zustand x_0^* . Was passiert nun mit dem abstoßenden Punkt x_1^* ? Mathematisch gesehen gibt es die konstante Lösung $x(t) = x_1^*$ für alle t . In Anwendungssituationen kann sich eine Größe allerdings nicht lange in dem abstoßenden stationären Punkt x_1^* aufhalten, da sie irgendwann durch irgendeinen kleinen Einfluss, der nicht mitmodelliert ist, aus dem instabilen Gleichgewicht gebracht wird und sich dann je nach Richtung der Störung entweder in den Punkt x_0^* oder in den Punkt x_2^* entwickelt.

Abb. 5.3 Graph von f mit Langzeitverhalten der Lösungen der Differentialgleichung (5.8)



5.3 Populationsmodelle mit Bejagung

5.3.1 Das Lösungsverhalten der logistischen Differentialgleichung

Wir wenden nun die Methode des qualitativen LöSENS auf die logistische Differentialgleichung

$$\dot{P} = f(P) = r \left(1 - \frac{P}{K}\right) P \quad (5.9)$$

an. Die Funktion hat die beiden Nullstellen $P_0^* = 0$ und $P_1^* = K$. Der Vorzeichenwechsel bei P_0^* lautet $-/+$, damit ist P_0^* instabil; der Vorzeichenwechsel bei P_1^* lautet $+/-$, damit ist P_1^* stabil. Lösungen mit Anfangswerten $0 < P_0 < K$ wachsen monoton und konvergieren von unten gegen die Kapazität, Lösungen zu Anfangswerten $K < P_0$ fallen monoton und konvergieren von oben gegen die Kapazität. Die Lösung hat ihre größte Wachstumsgeschwindigkeit bei der Größe $K/2$ erreicht und wächst dort mit einer Rate $rK/4$ (siehe Abb. 5.4). Wenn sich die Population also lange genug ungestört entwickeln kann, erwarten wir eine Populationsgröße K . Wir sagen auch: Wir erwarten das System im Zustand K .

5.3.2 Fangstrategien: Fischfang mit konstanter Fangrate und mit konstantem Aufwand

Wir stellen uns einen großen See mit stabilen äußeren Umständen vor und wollen den Fischbestand einer bestimmten Fischart modellieren. Im logistischen Wachstumsmodell gilt (5.9). Wir modellieren, welche Auswirkungen bestimmte Fangstrategien auf den Fischbestand haben, und vergleichen zwei Fangstrategien miteinander.

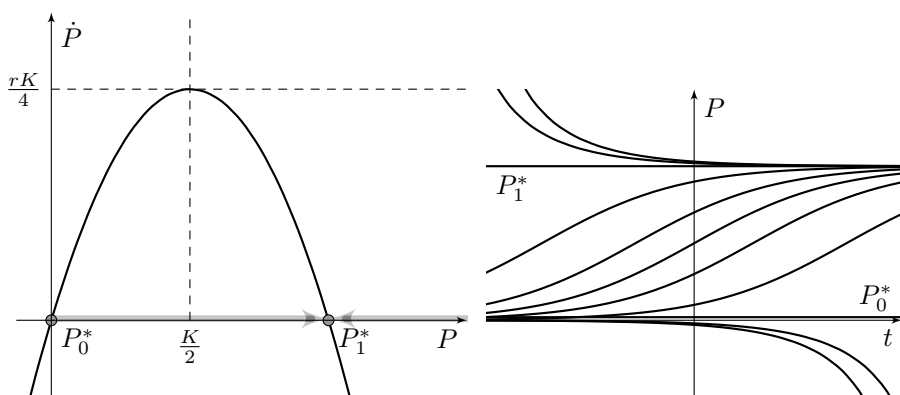


Abb. 5.4 Das qualitative Lösungsverhalten der logistischen Differentialgleichung und einige ausgewählte Lösungen

Fangstrategie konstante Fangrate

Nach der ersten Fangstrategie soll der Fischbestand mit einer konstanten Fangrate befischt werden, das heißt, pro Zeitraum soll immer dieselbe Menge an Fisch dem See entnommen werden. Das erinnert an die internationalen Fischereiübereinkommen, in denen auch feste Fangquoten vereinbart werden. Wir bezeichnen diese Menge pro Zeit mit F . Dadurch ändert sich unser Modell zu

$$\dot{P} = rP \left(1 - \frac{P}{K}\right) - F =: f_F(P). \quad (5.10)$$

Wir wollen nun untersuchen, wie groß diese Fangquote werden darf, ohne dass der Fischbestand ausstirbt. Man spricht von der maximalen nachhaltigen Fangrate oder dem *Maximum Sustainable Yield* (MSY). Dazu modifizieren wir unsere qualitative Auswertungsmethode noch ein wenig: Anstatt den Graphen der Gesamtfunktion f_F zu plotten, nutzen wir die additive Struktur dieser Funktion und plotten den Graphen der Funktion f sowie den Graphen der konstanten Funktion F in ein Koordinatensystem. Die stationären Punkte sind jetzt die Schnittstellen der beiden Graphen und ein stationärer Punkt ist stabil, wenn der Graph von f links der Schnittstelle über und rechts der Schnittstelle unterhalb des Graphen von F verläuft, bzw. instabil, wenn die Situation umgekehrt ist. Solange $F < \frac{rK}{4}$ behalten wir zwei Fixstellen $P_0^*(F)$ und $P_1^*(F)$, an denen die Ableitung $\dot{P}$ verschwindet, (siehe Abb. 5.5). Nach dem Obigen gilt: $P_0^*(F)$ ist abstoßend und $P_1^*(F)$ ist anziehend. Auf lange Sicht wird die Population gegen $P_1^*(F)$ streben, wenn sie von einem Wert größer als $P_0^*(F)$ startet. Ist sie jedoch zu Beginn kleiner als $P_0^*(F)$, stirbt sie in einem endlichen Zeitraum aus (dabei ist zu beachten, dass das Modell nur für $P \geq 0$ sinnvoll ist).

Vergößern wir nun F , laufen die beiden stationären Punkte auf der P -Achse zusammen. Für $F = \frac{rK}{4}$ gibt es genau einen stationären Punkt $a = \frac{K}{2}$. Für $F > \frac{rK}{4}$ gibt es keine Fixstelle mehr und die Ableitung ist stets negativ, sodass der Bestand stets abnimmt und in endlicher Zeit die Null erreicht. Was bedeutet das nun? Fischen wir mit einer Rate $F > \frac{rK}{4}$, wird der Fischbestand in endlicher Zeit aussterben. Fischen wir mit einer Rate $F < \frac{rK}{4}$,

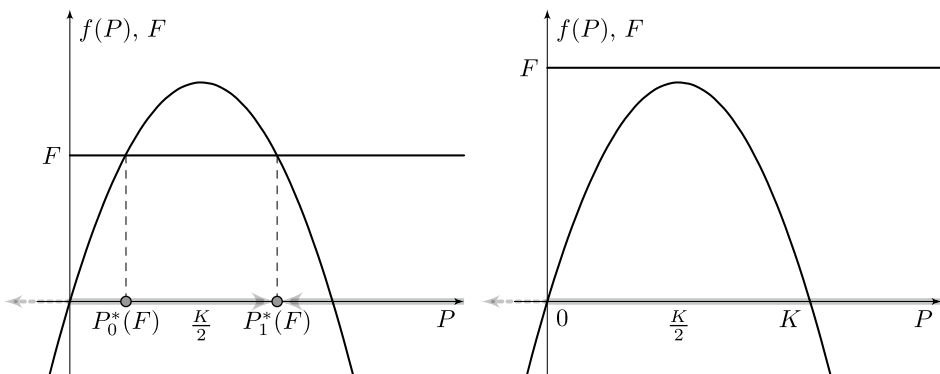


Abb. 5.5 Fischfang bei kleiner und bei großer konstanter Fangrate

stabilisiert sich der Bestand bei der Fixstelle $P_1^*(F)$. Je näher wir jedoch mit der Rate F an die kritische Grenze $\frac{rK}{4}$ heranrücken, desto kritischer wird die Situation des Fischbestandes: Der Bestand ist immer auch weiteren kleineren Schwankungen unterworfen, die Ursachen haben, welche in diesem Modell nicht berücksichtigt sind. Ist F nun sehr nah an $\frac{rK}{4}$, so kann der Bestand durch solch eine Schwankung unter den Wert $P_0^*(F)$ gedrückt werden, was dann zum Aussterben führt.

Wir können die Stabilität des Bestandes auch durch die Zeit charakterisieren, die der Bestand braucht, um sich von einer kleinen Störung zu erholen: Dazu schreiben wir $P(t) = P_1^*(F) + S(t)$ und betrachten $S(t)$ als Störung zur Zeit t . Die Störung S löst die Differentialgleichung

$$\dot{S} = \dot{P} = f_F(P) = f_F(P_1^*(F) + S).$$

Wir werden jetzt wieder das mächtige Instrument der Linearisierung nutzen. Wir gehen von einer kleinen Störung aus, also $|S/P_1^*(F)| \ll 1$. In diesem Fall ist es sinnvoll, die Funktion f_F durch die Linearisierung um $P_1^*(F)$ zu ersetzen:

$$\dot{S} \approx f_F(P_1^*(F)) + f'_F(P_1^*(F))(P_1^*(F) + S - P_1^*(F)) = f'_F(P_1^*(F))S.$$

Die Differentialgleichung lässt sich also, solange sich die Lösung in einer kleinen Umgebung um $P_1^*(F)$ befindet, durch eine lineare Differentialgleichung annähern. Es lässt sich zeigen, dass dann auch die Lösungen nah zusammen liegen. Einen Beweis und quantitative Abschätzungen zwischen der Lösung der vollen Differentialgleichung und der linearen Approximation finden sich z. B. in [92]. Akzeptieren wir diesen Sachverhalt, erhalten wir

$$S(t) \approx S_0 e^{f'_F(P_1^*(F))t}$$

(die Differentialgleichung des exponentiellen Wachstums können wir ja schon explizit lösen), dabei ist S_0 die anfängliche Störung zur Zeit $t = 0$. Ein Maß für die Stabilität des Bestandes ist die so genannte Rekonvaleszenzzeit t_R , die die Population benötigt, um die Größe der Störung auf den halben Wert zu reduzieren. Einsetzen liefert nun „in linearer Näherung“ $t_R = -\frac{\ln 2}{f'_F(P_1^*(F))}$. Andererseits gilt $f'_F(P_1^*(F)) = r \left(1 - \frac{2P_1^*(F)}{K}\right)$. Nähern wir uns nun mit F von unten der kritischen Quote $\frac{rK}{4}$, so konvergiert $P_1^*(F)$ von oben gegen $K/2$ und die Rekonvaleszenzzeit konvergiert gegen ∞ .

Insgesamt handelt es sich also um eine Fangstrategie, die entweder nur kleine Fangraten zulässt oder das Überleben des Bestandes gefährdet.

Fangstrategie konstanter Aufwand

Nun wollen wir eine zweite Fangstrategie diskutieren: Anstatt die Fangrate konstant zu halten, soll nun der Aufwand, mit dem der Bestand bejagt wird, konstant gehalten werden. Man kann sich vorstellen, dass z. B. eine feste Zahl an Booten benutzt wird und auch die Zeiten, in denen die Boote auf Fischfang sind, konstant gehalten werden. Zunächst muss ein

neues mathematisches Modell aufgestellt, also eine neue Differentialgleichung gefunden werden. Es ist plausibel anzunehmen, dass die Fangrate bei gleich bleibendem Aufwand proportional zur Bestandsgröße ist: Wenn doppelt so viele Fische im See sind, werden auch doppelt so viele gefangen. Wir modellieren die Situation also mit der Differentialgleichung

$$\dot{P} = rP \left(1 - \frac{P}{K}\right) - EP =: f_E(P). \quad (5.11)$$

Durch die Proportionalitätskonstante E wird der betriebene Aufwand dargestellt. Wir fragen uns nun: Bei welchem Aufwand erhalten wir auf lange Sicht die maximale Fangrate? Wie stabil gegen kleine Störungen ist die Situation in diesem Fall?

Aus (5.11) und der Gleichung $\dot{P} = 0$ erhalten wir die beiden Fixpunkte $P_0^* = 0$ und $P_1^*(E) = K(1 - \frac{E}{r})$. Für $E < rK$ ist $P_1^*(E)$ positiv und es liegt dort ein Vorzeichenwechsel von $+$ nach $-$ vor. Somit ist $P_1^*(E)$ in diesem Fall stabil (siehe Abb. 5.6). Auf lange Sicht nimmt der Bestand also die Größe $P_1^*(E)$ an.

Für die Fangrate $F = F(E)$ ergibt sich

$$F(E) = EP_1^*(E) = EK \left(1 - \frac{E}{r}\right).$$

Die Fangrate hängt also nicht monoton wachsend von E ab, sondern ist maximal für $E_{\max} = \frac{r}{2}$ mit dem Wert $F_{\max} = \frac{rK}{4}$. Mit dieser Strategie erhält man also dieselbe maximale Fangrate, die bei der Strategie *konstante Fangrate* nur in der instabilen Grenzsituation erreicht werden kann, in der jede noch so kleine Störung in die falsche Richtung zum Aussterben des Bestandes führt.

Auch hier können wir wie oben die Rekonvaleszenzzeit ausrechnen und sehen, dass eine kleine Störung in linearer Näherung in der Zeit $t_R = \frac{2\ln 2}{r}$ um die Hälfte abklingt.

Zu diskutieren bleibt, wie denn nun praktisch die optimale Aufwandsrate $E_{\max}$ gefunden werden kann, da die Werte für r und K in der Regel nur grob abgeschätzt werden können: Man variiert den Aufwand E so langsam, dass man zu jeder Zeit davon ausgehen kann, dass sich die Population im stationären Zustand $P_1^*(E)$ befindet. Die zugehörige Fangrate $F(E)$ lässt sich direkt messen. Bei langsamer Variation beginnend von kleinem E steigt die Fangrate zunächst und beginnt irgendwann (wenn $E_{\max}$ passiert wird) wieder zu sinken. Sobald sie kleiner wird, reduziert man E wieder leicht und kann sich so dem Maximum annähern.

Fischt man allerdings mit einem Aufwand $E > rK$, so ist der stationäre Punkt $P_1^*(E)$ negativ. Es gibt dann keinen positiven stationären Punkt, stattdessen ist P_0^* stabil und der Bestand würde auf lange Sicht aussterben, siehe Abb. 5.6.

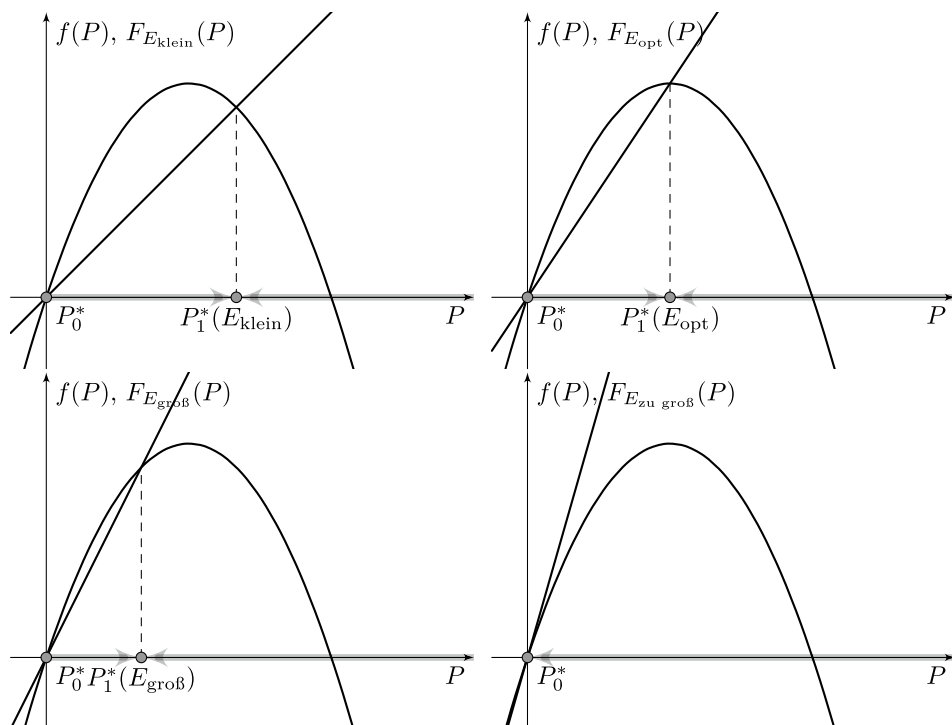


Abb. 5.6 Fischen mit konstantem Aufwand: kleiner, optimaler, größer und zu großer Aufwand

5.3.3 Schwellen- und Sättigungseffekte

Die Darstellung in diesem Kapitel orientiert sich an [6]. Gemäß der Fangstrategie *konstanter Aufwand* ändert sich die stabile Bestandsgröße $P_1^*(E)$ bei kleiner Änderung des Aufwands E auch nur geringfügig. In der Realität tritt jedoch mitunter das Phänomen auf, dass ein Fischbestand, der jahrelang stabil war, bei Befischung mit geringfügig größerem Aufwand (etwa ein paar Fischerboote mehr) in kurzer Zeit zusammenbricht und sich bei der anschließenden Reduzierung des Aufwandes auf den alten Stand nicht wieder zu alter Größe erholt. Man spricht auch von einem *Kipppunkt*. Wir werden nun ein Modell erläutern, das diesen Effekt erklären kann. Dazu wollen wir die Fangstrategie *konstanter Aufwand* genauer betrachten und zwei weitere Effekte berücksichtigen. Wir stellen uns vor, dass der Fischfang mit Booten betrieben wird. Auch bei beliebig großem Fischaufkommen kann ein Boot nicht beliebig viele Fische in einer bestimmten Zeitspanne einfangen. Irgendwann ist das Boot voll, muss zurück zum Hafen, abgeladen und gesäubert werden, bevor es wieder herausfahren kann. Jedes Boot hat also eine so genannte Sättigungsgrenze S , die angibt, welche Menge Fisch pro Zeiteinheit von einem Boot maximal gefischt werden kann, egal wie viel Fisch auch

vorhanden sein mag. Dazu gehört eine bestimmte Größe A des Fischbestands, bei der diese Grenze erreicht wird.

Ein anderer Effekt betrifft die Fischpopulation und wird Schwelleneffekt genannt. Eine sehr kleine Population hat ggf. die Möglichkeit, sich der Bejagung zu entziehen, indem sie Bereiche aufsucht (z.B. flache Bereiche des Sees oder sehr tiefe Bereiche), in denen kein Fischfang möglich ist. Bis zu einer Größe V kann sich die Population vollkommen verstecken. Erst wenn die Population größer als V ist, wird Fischfang erfolgreich.

Diese Überlegungen sollen jetzt mathematisiert werden. In dem einfachsten denkbaren Modell hängt die Fangrate $g(P)$ eines einzelnen Boots folgendermaßen von der Größe P des Fischbestands ab: Für $P \leq V$ ist sie null, für $V < P \leq A$ wächst sie linear von null bis zur Schwelle S , und für $P > A$ behält sie konstant den Wert S . Der zugehörige Funktionsterm lautet:

$$g(P) = g_{SAV}(P) = \begin{cases} 0, & \text{falls } P \leq V \\ \frac{S}{A-V} (P - V), & \text{falls } V < P \leq A \\ S, & \text{falls } A < P \end{cases}$$

Natürlich sind auch Modelle ohne Knickstellen möglich. Solche lassen sich z.B. mit gebrochen-rationalen Funktionen modellieren, siehe z.B. [65]. Die gesamte Fangrate $F(P)$ erhalten wir dann, indem wir $g(P)$ mit der Anzahl n der eingesetzten Boote multiplizieren, also $F(P) = ng(P)$. Insgesamt wird die Situation durch die Differentialgleichung

$$\dot{P} = rP \left(1 - \frac{P}{K}\right) - ng(P) \quad (5.12)$$

modelliert. Wir müssen nun weitere Annahmen bezüglich der Größenverhältnisse von V , A und K treffen. Wir wollen hier nur den interessantesten Fall diskutieren, in dem sich die oben angesprochenen Phänomene finden lassen. Wie sich herausstellen wird, ist das der Fall für $A < K/2$. Wir gehen also von kleinen Booten aus, die ihre Sättigungsgrenze bei einer Bestandsgröße erreichen, die kleiner ist als die halbe Kapazität. In den ersten fünf Diagrammen in Abb. 5.7 ist die Situation für unterschiedliche Bootsanzahlen n dargestellt. Für kleines n (siehe Abb. 5.7 links oben), gibt es nur einen stabilen Fixpunkt, den wir jetzt, anders als im alten Modell, mit $P_3^*(n)$ bezeichnen. Die zugehörige Fangrate ist relativ klein. Für solche n befindet sich der Bestand langfristig im Zustand $P_3^*(n)$. Wird n vergrößert, wird der Fixpunkt kleiner; die zugehörige Fangrate, die Höhe des Schnittpunkts, nimmt zu. Ab einer bestimmten Zahl $n = \underline{n}$ (in der Regel nicht ganzzahlig, siehe Abb. 5.7 oben rechts), gibt es den weiteren Fixpunkt A , der sich bei weiterer Vergrößerung von n in zwei Fixpunkte aufspaltet, $V < P_1^*(n) < A < P_2^*(n)$ (siehe Abb. 5.7, Mitte links). Dabei ist $P_1^*(n)$ stabil und $P_2^*(n)$ instabil, wie man leicht am Vorzeichenwechsel von $\dot{P}$ erkennt. Für solche n existieren zwei stabile stationäre Punkte, $P_1^*(n)$ und $P_3^*(n)$, und es kommt auf die Vorgeschichte an, in welchem Punkt sich die Populationsgröße auf lange Sicht einfindet: Ist sie zu Beginn größer als $P_2^*(n)$, entwickelt sie sich nach $P_3^*(n)$; startet sie kleiner als $P_2^*(n)$, strebt sie gegen $P_1^*(n)$. Sowohl für die Fischer als auch für die Fische ist P_3^* günstiger, denn die Fangrate ist höher und es gibt mehr Fische.

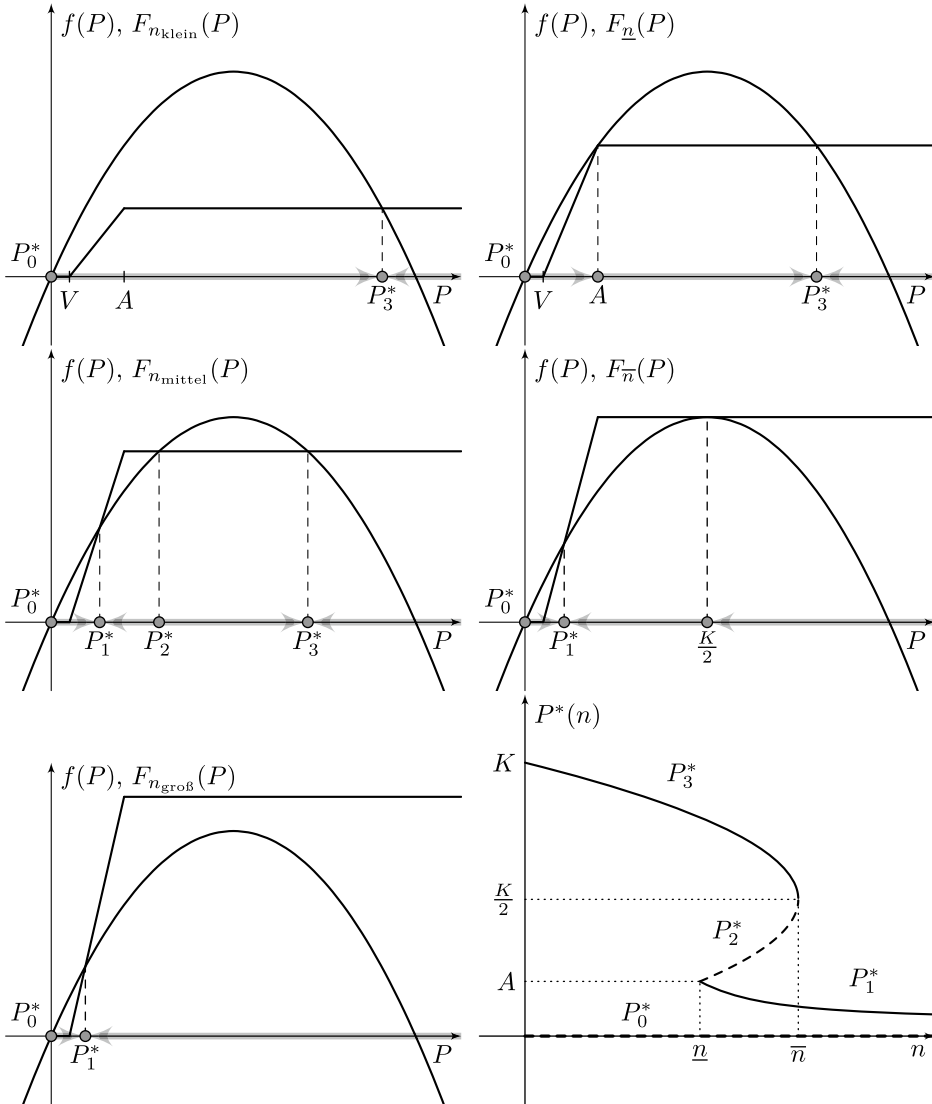


Abb. 5.7 Schwellen- und Sättigungseffekt: das Modell für kleine, mittlere und große Bootsanzahlen n , Bifurkationsdiagramm mit dem Bifurkationsparameter n

Wird die Bootsanzahl n weiter erhöht, bewegen sich die stationären Punkte $P_2^*(n)$ und $P_3^*(n)$ aufeinander zu. In der Grenzsituation $n = \bar{n}$ (siehe Abb. 5.7 Mitte rechts), wird gerade die maximale Fangrate $\frac{rK}{4}$ erreicht. Im Gegensatz zu dem einfacheren Modell *konstanter Aufwand* aus Abschn. 5.3.2 ist die Situation jetzt aber instabil: Eine kleine Störung kann zum Zusammenbrechen der Population auf die Größe $P_1^*(n)$ führen. Wird die

Bootsanzahl auch nur um ein weiteres Boot erhöht, verschwindet der größere stationäre Punkt und die Population entwickelt sich hin zur kleinen anziehenden Größe P_1^* zusammen mit der zugehörigen kleinen Fangrate (siehe Abb. 5.7 unten links).

Die Gesamtsituation ist in Abb. 5.7 rechts unten in der Form eines Bifurkationsdiagramms zusammengefasst. Die von n abhängigen stationären Punkte sind gegen den Bifurkationsparameter n aufgetragen. In diesem Diagramm verbirgt sich der in der Einleitung zu diesem Abschnitt angesprochene Effekt: Wir stellen uns vor, dass wir vom Zustand ohne Befischung so langsam die Anzahl der Boote erhöhen, dass sich der Bestand immer in einem zum aktuellen n gehörigen stabilen Zustand befindet. Der Bestand startet in K und wandert mit zunehmenden n den P_3^* -Ast herab. Beim Passieren von $n = \bar{n}$ verschwindet dieser stationäre Punkt und der Bestand fällt in kurzer Zeit auf den P_1^* -Ast herunter. Eine beliebig kleine Änderung von n von unten über die $\bar{n}$ -Grenze hinweg bewirkt also eine große Veränderung in der Bestandsgröße! Nun bewegt sich der Zustand auf dem P_1^* -Ast. Wird n nun wieder reduziert, wandert der Bestand den P_1^* -Ast hoch und springt erst beim Unterschreiten von $n = \underline{n}$ auf den P_3^* -Ast zurück. Man spricht von einem Hysterese-Effekt. Praktisch bedeutet das: Eine kleine Überschreitung von $\bar{n}$ ist ziemlich verheerend; es reicht nicht, n nur wieder ein wenig zu verkleinern, sondern der Bootsbestand muss erst bis unter die Grenze $\underline{n}$ reduziert werden, damit sich der Bestand wieder erholen kann.

Nach diesem Modell führt die Befischung eines Sees mit kleinen Booten also zu einer instabilen Situation, wenn man ihn in der Nähe der maximalen Fangrate $\frac{rK}{4}$ befischen will. Interessanterweise wird die Situation stabil, wenn mit größeren Booten gefischt wird, deren Sättigungsgrenze S erst ab einer Bestandsgröße $A > \frac{K}{2}$ erreicht wird. In diesem Fall lässt sich die maximale Fangrate in einer stabilen Art und Weise erreichen. Das soll in Aufg. 5.3 diskutiert werden. Gemäß diesem Modell wäre es also unter den Gesichtspunkten der Nachhaltigkeit der Befischung besser, wenige und dafür möglichst große Boote oder Schiffe einzusetzen.

5.4 Der Tannenwickler und die Balsamtanne – erster Teil

In diesem Kapitel und in Abschn. 7.3 soll ein mathematisches Modell für die Interaktion des Tannenwicklers *Choristoneura fumiferana*, einer nordamerikanischen Schmetterlingsart, mit seiner vornehmlichen Wirtspflanze, der Balsamtanne *Abies balsamea*, und seinen Fressfeinden, vornehmlich Vögel wie bestimmten Grasmücken-, Ammer- und Kernbeißerarten, entwickelt werden, siehe Abb. 5.8.

Die Tannenwickler-Populationen zeichnen sich dadurch aus, dass sie eine lange Zeit, 30 bis 40 Jahre, nur eine sehr geringe Populationsdichte aufweisen, so dass sie während dieser Zeit als selten bezeichnet werden können. Wenn dann aber die Tannenbestände eine günstige Größe erreicht haben, „explodiert“ die Populationsdichte der Tannenwickler innerhalb weniger Jahre, man spricht von einer *Epidemie*. Dabei schädigen sie ihre Wirtspflanze so stark, dass die Tannenbestände über große Flächen hinweg absterben, woraufhin auch

Abb. 5.8 Der Tannenwickler *Choristoneura fumiferana* als Schmetterling und als Raupe und die Balsamtanne *Abies balsamea*. (Fotos: U.S. Fish and Wildlife Service)



der Tannenwickler-Bestand wieder zusammenbricht. Auf den abgestorbenen Flächen entwickelt sich anschließend wieder eine neue Tannenpopulation. Beobachtungsdaten deuten darauf hin, dass solche Epidemien periodisch mit einer Periode zwischen 30 und 40 Jahren auftreten, siehe [16].

Weil die durch diese Epidemien hervorgerufenen forstwirtschaftlichen Schäden recht groß sind, ist man an effektiven Bekämpfungsstrategien gegen den Tannenwickler interessiert. Ein gutes mathematisches Modell kann hier helfen, unterschiedliche Bekämpfungsstrategien im Vorfeld zu simulieren und auf ihre Wirksamkeit hin abzuschätzen. Wir wollen hier ein einfaches Modell mit Differentialgleichungen vorstellen, das von Ludwig et al. [60] entwickelt wurde.

Das Modell besteht aus mehreren Teilen: Im Wesentlichen sollen die Populationsdichte der Tannenwickler (Anzahl Tannenwickler pro Fläche), der Zustand des Balsamtannen-Bestands und ihre Interaktionen modelliert werden. Um die Phänomene erklären zu können, müssen auch noch die Fressfeinde des Tannenwicklers berücksichtigt werden.

Dabei vollziehen sich die Veränderungen innerhalb der Tannenwickler-Population sehr viel schneller als die Veränderung im Zustand des Balsamtannen-Bestands: Die Dichte der Tannenwickler-Population kann innerhalb weniger Jahre um ein Mehrhundertfaches zunehmen, während die Tannen 7–10 Jahre benötigen, um Fraßschäden zu kompensieren. Die Lebensdauer der ungeschädigten Balsamtanne beträgt dabei 100 bis 150 Jahre.

Wir werden bei der Modellierung zweischrittig vorgehen: In diesem Kapitel betrachten wir die Tannenpopulation als statisch, modellieren also die „schnelle“ Entwicklung der Tannenwickler-Population vor dem Hintergrund eines konstanten Tannenbestandes. In Abschn. 7.3 werden wir dann die zeitliche Entwicklung der Balsamtannen mit hinzuziehen.

5.4.1 Die Modellierung der Populationsdichte der Tannenwickler

Die zu modellierende Größe ist also die Populationsdichte $P(t)$ des Tannenwicklers zur Zeit t in einem bestimmten Areal. Die Größe P hat damit die Dimension Anzahl pro Fläche. In einem Grundmodell setzen wir ein logistisches Wachstum an:

$$\dot{P} = r_P P \left(1 - \frac{P}{K_P} \right) \quad (5.13)$$

Dabei wird der Zustand des Tannenbestandes später unter anderem in den Kapazitätsparameter K_P eingehen: Ein gut ausgewachsener Tannenbestand mit einer großen Zweigfläche wird den Tannenwicklern eine größere Kapazität bieten als ein Bestand mit einer geringen Zweigfläche. Der Parameter r_P hat die Dimension einer inversen Zeit und bestimmt, wie schnell sich die Tannenwickler-Population entwickelt.

Die Tatsache, dass die Tannenwickler-Population über lange Zeiträume sehr klein bleibt und dann „schlagartig explodiert“, deutet darauf hin, dass die Entwicklung nicht nur durch die Futterversorgung bestimmt ist, denn dann müsste die Tannenwickler-Population ja gleichmäßig mitwachsen, sondern zunächst durch einen anderen Einfluss gehemmt wird. Für diesen hemmenden Einfluss kommen vor allem Fressfeinde und hier insbesondere die oben genannten Vogelarten infrage. Wir fügen also zur Gl. (5.13) einen weiteren Term $-g(P)$ hinzu, der die Bejagung durch Fressfeinde modellieren soll. Dabei gehen wir von den folgenden Annahmen aus:

- Die Fressfeinde sind Reviervögel. Das heißt, ihre Anzahl wird im Wesentlichen nicht durch das Futterangebot bestimmt, sondern durch die Anzahl passender Reviere, also hauptsächlich durch die Größe der Fläche. Insbesondere sind die Fressfeinde nicht auf die Tannenwickler als Nahrungsgrundlage angewiesen: Wenn Tannenwickler da sind, werden sie gefressen, wenn keine vorhanden sind, wird auf eine alternative Futterquelle ausgewichen. Die Anzahl der Fressfeinde ist also im Wesentlichen unabhängig von der Größe der Tannenwickler-Population.
- Es gibt einen Sättigungseffekt: Jeder Fressfeind kann maximal nur eine bestimmte Zahl an Tannenwicklern pro Zeit fressen.
- Es gibt einen Schwelleneffekt: Damit die Tannenwickler von den Fressfeinden als lohnende Nahrungsquelle wahrgenommen werden, muss die Tannenwickler-Dichte einen bestimmten Schwellenwert überschreiten. Ist die Tannenwickler-Dichte zu gering, so weichen die Fressfeinde auf alternative, leichter zu findende Nahrungsquellen aus.

Wir übersetzen diese Annahmen in Anforderungen an den Bejagungsterm $g(P)$, der übrigens die Dimension Anzahl pro Fläche pro Zeit haben muss.

- g ist monoton wachsend. Bei größerer Populationsdichte wird mit einer größeren Rate gefressen.

- g nähert sich für $P \rightarrow \infty$ einer Sättigungsgrenze S an, die angibt, mit welcher Rate die Fressfeinde maximal fressen können.
- g bleibt in einer Umgebung von $P = 0$ klein, beginnt also nicht direkt zu wachsen. Eine bequeme mathematische Modellierung besteht darin, g in einer Umgebung von $P = 0$ als quadratisch anzusetzen.

Eine einfache analytische Funktion, die diesen Anforderungen entspricht, ist gegeben durch

$$g(P) = \frac{SP^2}{V^2 + P^2} = S \cdot \left(1 - \frac{V^2}{V^2 + P^2}\right) = \frac{SP^2}{V^2} \cdot \frac{1}{1 + \frac{P^2}{V^2}}. \quad (5.14)$$

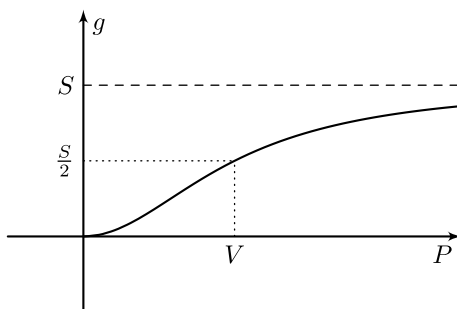
Dabei ist S die Sättigungsrate und V bezeichnet die Populationsdichte, bei der mit der halben Sättigungsrate gejagt wird. Aus den unterschiedlichen Darstellungen der Abbildungsvorschrift in (5.14) lässt sich gut ablesen, dass die unterschiedlichen Anforderungen an die Funktion – die Monotonie und das Grenzwertverhalten bei $P = \infty$ und $P = 0$ – erfüllt sind. Der Graph dieser Funktion ist in Abb. 5.9 dargestellt. Die konkrete Wahl der Funktion g ist zu einem gewissen Grad beliebig, wir könnten hier auch mit einem abschnittsweise definierten Term ähnlich wie in Abschn. 5.3.3 agieren. Die qualitativen Befunde sollten sich dadurch nicht ändern (was sie auch nicht tun, wie man überprüfen kann). Der Term (5.14) ist bei den mathematischen Biologen für ähnliche Situation recht beliebt, so dass er auch einen eigenen Namen erhalten hat: Es handelt sich um eine *Type-III S-shaped functional response*, siehe [46] (das muss man sich nicht merken).

Wir müssen jetzt also das Gesamtmodell

$$\dot{P} = r_P P \left(1 - \frac{P}{K_P}\right) - S \frac{P^2}{V^2 + P^2} \quad (5.15)$$

mit den vier Parametern r_P , K_P , S und V studieren. Insbesondere interessiert uns hierbei, wie die stationären Punkte von den Parametern abhängen.

Abb. 5.9 Graph des Bejagungsterms mit Sättigungsrate S und Populationsdichte bei halber Sättigungsrate V



5.4.2 Stationäre Punkte, Bifurkationen und die kritische Kurve

Um die Situation übersichtlicher zu gestalten, entdimensionalisieren wir zunächst und hoffen dadurch, weil wir zwei Skalen wählen können, die Anzahl der Parameter auf zwei zu reduzieren.

Entdimensionalisierung Nach dem bekannten Schema führen wir die dimensionslosen Variablen ein,

$$\tau = \frac{t}{\bar{t}} \quad \text{und} \quad p(\tau) = \frac{P(\bar{t}\tau)}{\bar{P}},$$

und berechnen die dimensionslose Differentialgleichung, siehe auch Abschn. 1.3,

$$\dot{p} = \bar{t} r_P p \left(1 - \frac{\bar{P}}{K_P} p \right) - \frac{\bar{t}}{\bar{P}} \cdot \frac{S \bar{P}^2}{\bar{P}^2} \cdot \frac{p^2}{\frac{V^2}{\bar{P}^2} + p^2}.$$

Um die Abhängigkeit der stationären Punkte von den Parametern einfach beschreiben zu können, ist es günstig, den mathematisch komplizierten Term von Parametern zu befreien, in diesem Fall also den Jagdterm. Man beachte, dass es sich hierbei nicht um eine inhaltliche Erwägung handelt, sondern um eine rein kalkülhafte. Wir wählen

$$\bar{P} = V \quad \text{und} \quad \bar{t} = \frac{V}{S},$$

also die Populationsgröße, bei der die Bejagungsrate die Hälfte der maximalen Rate beträgt, und die Zeit, welche die Vögel benötigen würden, diese Menge an Tannenwickler-Larven beim Fressen mit maximaler Fressrate zu vertilgen. Mit den dimensionslosen Parametern

$$a = \frac{V r_P}{S} \quad \text{und} \quad b = \frac{K_P}{V}$$

erhalten wir die Gleichung

$$\dot{p} = ap \left(1 - \frac{p}{b} \right) - \frac{p^2}{1 + p^2} = p \cdot \left(a \left(1 - \frac{p}{b} \right) - \frac{p}{1 + p^2} \right). \quad (5.16)$$

Stationäre Punkte und ihre Stabilität Aus Gl. (5.16) können wir ablesen, dass $p_0^* = 0$ ein instabiler stationärer Punkt ist, das Vorzeichen wechselt von $-$ nach $+$. Die weiteren positiven stationären Punkte erhalten wir aus den Lösungen der Gleichung

$$a \left(1 - \frac{p}{b} \right) = \frac{p}{1 + p^2}. \quad (5.17)$$

Dabei können wir die linke Seite $r(p) = a(1 - \frac{p}{b})$ als die Pro-Kopf-Reproduktionsrate aufgrund der Nahrungsgrundlage und die rechte Seite $s(p) := \frac{p}{1+p^2}$ als Pro-Kopf-Sterberate aufgrund der Bejagung durch die Fressfeinde interpretieren. Plotten wir die beiden

Funktionen in ein Diagramm, erhalten wir die positiven stationären Punkte gerade als die Schnittstellen der Funktionsgraphen. In Abb. 5.10 ist eine Parameterkombination (a, b) gewählt worden, so dass drei stationäre Punkte $p_1^* < p_2^* < p_3^*$ existieren, wobei aufgrund der Vorzeichensituation p_1^* und p_3^* stabil sind und p_2^* instabil ist. Je nach Wahl der Parameter a und b kann aber auch nur ein kleiner oder ein großer positiver stationärer Punkt auftreten, siehe Abb. 5.11. Ändern sich nun die Parameter (aufgrund des sich langsam ändernden Balsamtannen-Bestands), kann es passieren, dass zwei stationäre Punkte zusammenfallen und dann verschwinden bzw. dass zwei neue stationäre Punkte entstehen.

Wir verfolgen die unterschiedlichen stationären Punkte in der Situation, dass der Parameter b festgehalten und der Parameter a variiert wird. In Abb. 5.12 sind für $b = 14$ in einem Bifurkationsdiagramm die stationären Punkte über dem Parameter a aufgetragen. Stellen wir uns nun vor, dass der Parameter a so langsam von null beginnend vergrößert wird, dass sich die Tannenwickler-Population immer in einem stationären Punkt befindet, so startet das System auf dem p_1^* -Ast und wandert diesen mit zunehmendem a hinauf. Passiert der Parameter a von unten den kritischen Wert $\bar{a}$, so muss der stationäre Punkt den p_1^* -Ast verlassen. Die Populationsdichte wächst in kurzer Zeit stark an und „springt“ dabei auf den p_3^* -Ast hinauf.

Genau so ein Phänomen suchen wir: Die Tannenwickler-Population bleibt lange Zeit niedrig und ändert sich dann plötzlich innerhalb von kurzer Zeit von einer niedrigen auf eine hohe Dichte, obwohl sich die Umgebung, dargestellt durch die Parameter a und b , nur geringfügig geändert hat. Auch hier haben wir einen Hystherese-Effekt: Das System springt

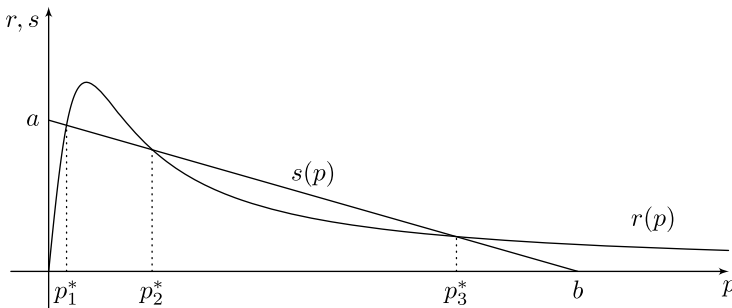


Abb. 5.10 Parameterkombination mit drei positiven stationären Punkten

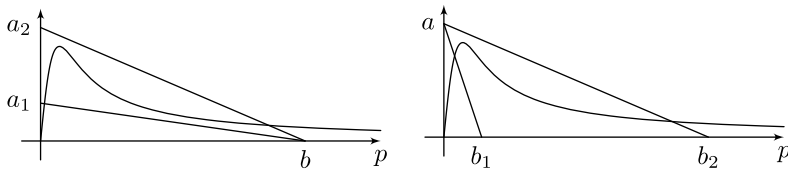


Abb. 5.11 Parameterkombinationen mit einem positiven stationären Punkt

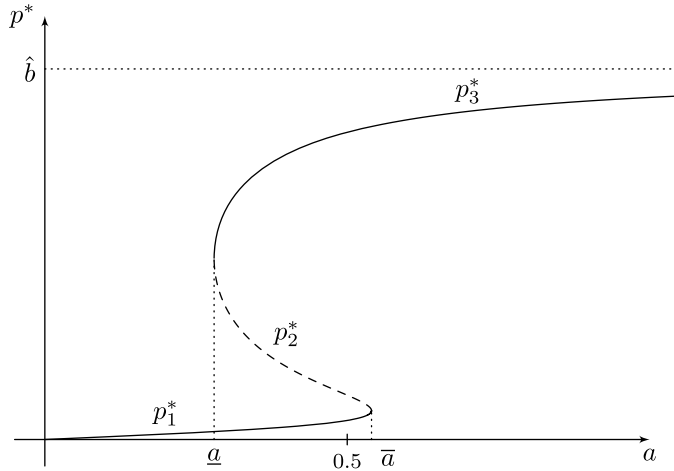


Abb. 5.12 Bifurkationsdiagramm für $b = 14$: Die p_1^* - und p_3^* - Äste sind stabil, der p_2^* -Ast ist instabil

erst wieder auf den p_1^* -Ast zurück, wenn der Parameter a unter den Wert $\underline{a}$ zurückgeführt wird.

Wir wollen jetzt allgemein (nicht nur für konstantes b) die kritischen Parameterpaare (b, a) bestimmen, an denen solche Übergänge von einem zu drei positiven stationären Punkten oder umgekehrt stattfinden, und in einem b - a -Diagramm darstellen. Das sind genau die Übergänge, an denen das System springen kann.

Die kritischen Kurven

Zwei stationäre Punkte fallen genau dann zusammen, wenn die Graphen von $r(p)$ und $s(p)$ im Schnittpunkt tangential zueinander sind, siehe Abb. 5.13. Für $p > 1$ müssen wir also das Gleichungssystem

$$r(p) = s(p) \quad \text{und} \quad r'(p) = s'(p),$$

also

$$\begin{aligned} a \left(1 - \frac{p}{b}\right) &= \frac{p}{1 + p^2} \\ -\frac{a}{b} &= \frac{1 - p^2}{(1 + p^2)^2} \end{aligned}$$

nach a und b auflösen und erhalten als Lösung die mit p parametrisierte *kritische Kurve*

$$(1, \infty) \ni p \mapsto \left(\frac{2p^3}{p^2 - 1}, \frac{2p^3}{(1 + p^2)^2} \right)^t,$$

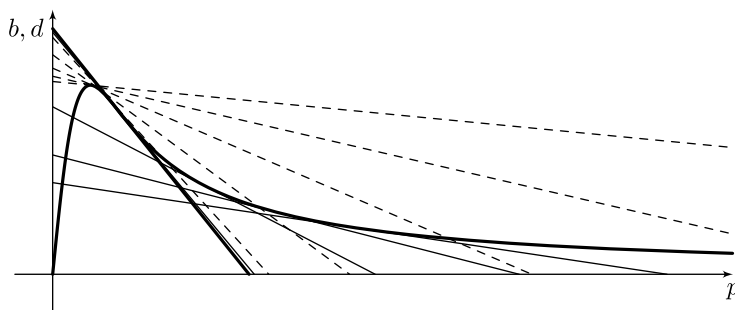


Abb. 5.13 Parameterkombinationen mit zwei zusammenfallenden positiven stationären Punkten: Die Graphen der Pro-Kopf-Reproduktionsrate und der Pro-Kopf-Sterberate sind in einem der beiden Schnittpunkte tangential zueinander. Die Parameterpaare zu den gestrichelten Tangenten an Schnittstellen kleiner als die Wendestelle $p_W = \sqrt{3}$ bilden den oberen Zweig, Parameterpaare zu den durchgezogenen Tangenten an Schnittstellen größer der Wendestelle bilden den unteren Zweig der kritischen Kurve in Abb. 5.14. Die Spitze wird durch das Parameterpaar der Wendetangente (in dieser Abbildung gefettet dargestellt) gebildet

die in Abb. 5.14 geplottet ist. Für $p = 1$ hat das Gleichungssystem keine Lösung und für $0 < p < 1$ wird der Lösungsparameter b negativ. Im Inneren der kritischen Kurve gibt es drei positive stationäre Punkte p_1^* , p_2^* und p_3^* . Auf dem oberen Zweig fallen p_1^* und p_2^* zusammen, oberhalb gibt es nur den Punkt p_3^* . Auf dem unteren Zweig fallen die stationären Punkte p_2^* und p_3^* zusammen, unterhalb existiert nur der Punkt p_1^* . In der Spitze der Kurve fallen alle drei positiven stationären Punkte zusammen (das ist genau der Fall, wenn der Graph der Pro-Kopf-Reproduktionsrate r die Wendetangente des Graphen der Pro-Kopf-Sterberate s ist). Beim Passieren des oberen kritischen Zweigs von unten nach oben springt das System vom p_1^* - auf p_3^* -Ast, eine Epidemie bricht aus. Beim Passieren des unteren Zweigs der kritischen Kurve von oben nach unten springt das System von p_3^* auf p_1^* . Die Epidemie erlischt, man spricht von einem *endemischen Zustand* oder einer *Endemie*.

Die Kurve in Abb. 5.14 ist übrigens in den 70er Jahren des letzten Jahrhunderts unter dem Namen *Cusp-Katastrophe* zu einer etwas zweifelhaften Popularität gekommen: Alle möglichen Gegensatzpaare mit diskontinuierlichen Übergängen wie Flucht und Angriff oder Genie und Wahnsinn wurden mit Hilfe der Cusp-Katastrophe „modelliert“, siehe z. B. [95]. Das führte natürlich schnell zu berechtigter Kritik: Der geneigten Leserin mit einem Faible für fundierte Verisse sei [86] empfohlen. Eine interessante Darstellung der Hintergründe, legitimen Anwendungen und Beschränkungen der auf R. Thom zurückgehenden sogenannten Katastrophentheorie findet sich in [4].

Wir werden in Abschn. 7.3 ebenfalls die Entwicklung des Balsamtannen-Bestands modellieren und die beiden Modelle miteinander koppeln. Dabei wird der Zustand des Tannenbestandes über die Parameter K_P und V die dimensionslosen Parameter a und b verändern und dadurch an einer bestimmten Stelle der Entwicklung die Tannenwickler-Population in einen

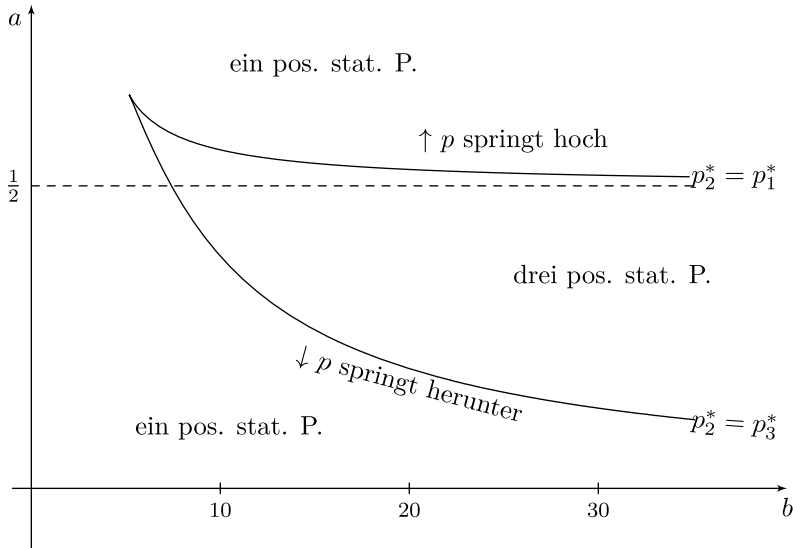


Abb. 5.14 Die kritische Kurve für die Gl. (5.16)

epidemischen Zustand springen lassen. Umgekehrt wird auch der die Populationsgröße der Tannenwickler-Population in das Modell für den Tannenbestand eingehen.

5.5 Eine einfache Lösungstheorie für Differentialgleichungen – die Methode von der Trennung der Variablen

In diesem Kapitel soll ein elementarer Existenz- und Eindeutigkeitssatz für autonome Differentialgleichungen vorgestellt und bewiesen werden, der unter dem Stichwort *Trennung der Variablen* bekannt ist. Während in den meisten Lehrbüchern diese Methode lediglich mit dem Ziel eingesetzt wird, explizite Lösungen zu berechnen, soll sie hier vorrangig unter theoretischen Gesichtspunkten behandelt werden: Sie liefert das absichernde Fundament für die Benutzung der qualitativen Methoden. Im Gegensatz zu dem modernen Existenz- und Eindeutigkeitssatz von Picard-Lindelöf, den wir in Abschn. 6.2 in Verbindung mit Differentialgleichungssystemen (ohne Beweis) vorstellen werden, liefert die Methode der Trennung der Variablen für die einfache Situation zeitautonomer Differentialgleichungen tatsächlich ein notwendiges und hinreichendes Kriterium für die Existenz und Eindeutigkeit von Lösungen, während der Satz von Picard-Lindelöf nur ein hinreichendes Kriterium liefert.

5.5.1 Die Existenz- und Eindeutigkeitsaussagen

Die Existenz von Lösungen zu (5.8) für stetiges f lässt sich relativ leicht mit dem Ansatz von der Trennung der Variablen zeigen. Für den Spezialfall $f(x) = x$ entspricht das übrigens auch gerade dem von Felix Klein vorgeschlagenen Zugang zur Exponentialfunktion, siehe [53]. Delikater ist die Frage nach der Eindeutigkeit von Lösungen: Bekanntlich hat das Anfangswertproblem (5.8), wenn f nur als stetig vorausgesetzt wird, nicht in jedem Fall eindeutige Lösungen. Ein gängiges Gegenbeispiel ist das Anfangswertproblem

$$\dot{x} = 2\sqrt{|x|}, \quad x(0) = 0 \quad (5.18)$$

mit der Lösungsschar $x_c(t) = (t-c)^2$ für $t \geq c$ und $x_c(t) = 0$ sonst. Dabei ist $c \in \mathbb{R}$ beliebig, siehe Abb. 5.15. Für gewöhnlich wird Eindeutigkeit durch die lokale Lipschitz-Stetigkeit von f erzwungen, siehe Definition 5.1. Insbesondere ist die lokale Lipschitz-Stetigkeit erfüllt, wenn f stetig differenzierbar ist. Interessanterweise lässt sich die eindeutige Lösbarkeit von (5.8) aber auch durch eine schwächere Bedingung erreichen, die lediglich den Begriff des uneigentlichen Integrals benutzt. Im Wesentlichen darf das uneigentliche Integral von $1/f$ an den Nullstellen von f gerade **nicht** endlich sein. Diese Tatsache wollen wir im Folgenden darstellen. Dabei orientieren wir uns grob an der Darstellung in [87].

Für eine stetige Funktion f von $\mathbb{R}$ nach $\mathbb{R}$ und $x \in \mathbb{R}$ definieren wir $x_+(x) \leq \infty$ als die nächstgrößere Nullstelle, formal $x_+(x) := \min \{y \geq x \mid f(y) = 0\}$, und genauso $-\infty \leq x_-(x) := \max \{y \leq x \mid f(y) = 0\}$ als die nächstkleinere Nullstelle. Wie gewöhnlich sei dabei $\min \emptyset = +\infty$ und $\max \emptyset = -\infty$. Falls x selber eine Nullstelle von f ist, gilt offensichtlich $x_-(x) = x = x_+(x)$. Wir sagen, eine Lösung von (5.8) ist *eindeutig*, wenn sie auf jedem Intervall um t_0 , auf dem überhaupt eine Lösung existiert, eindeutig ist. Eine Lösung heißt *maximal*, wenn sie nicht fortgesetzt werden kann, und eine Lösung heißt *global*, wenn sie für alle Zeiten $t \in \mathbb{R}$ existiert.

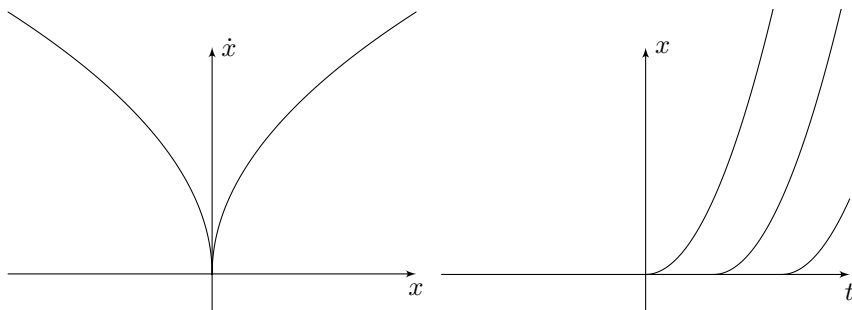


Abb. 5.15 Ein Gegenbeispiel zur Eindeutigkeit von Lösungen des Anfangswertproblems $\dot{x} = f(x)$, $x(t_0) = 0$

Satz 5.1

Es sei f eine auf $\mathbb{R}$ definierte reellwertige stetige Funktion und für alle $x \in \mathbb{R}$ mit $f(x) \neq 0$ gelte

$$\left| \int_x^{x_{\pm}(x)} \frac{1}{f(z)} dz \right| = \infty, \quad \text{falls } |x_{\pm}(x)| < \infty. \quad (5.19)$$

Dann hat für alle $(t_0, x_0) \in \mathbb{R}^2$ das Anfangswertproblem (5.8) eine eindeutige Lösung. Die maximale Lösung ist unbeschränkt oder global, ferner ist sie konstant oder streng monoton.

Beweis 5.1

1. Schritt: Ohne Einschränkung können wir $t_0 = 0$ annehmen, denn $x(t)$ ist genau dann eine Lösung von (5.8) auf dem Intervall I um t_0 , wenn $y(t) := x(t + t_0)$ eine Lösung zu $\dot{y} = f(y)$ und $y(0) = x_0$ auf dem verschobenen Zeitintervall $I - t_0$ um 0 ist. Sieht man diese Aussage vor einem Modellierungshintergrund, spiegelt sich darin die Unabhängigkeit der Umwelt von der Zeit.
2. Schritt: Es gelte zunächst $f(x_0) \neq 0$ und es sei $J_{x_0} := (x_-(x_0), x_+(x_0))$ das maximale offene Intervall um x_0 , auf dem f keine Nullstelle hat. Weil die Funktion f stetig ist, hat sie nach dem Zwischenwertsatz auf J_{x_0} ein einheitliches Vorzeichen. Wir nehmen an, dass f auf dem Intervall J_{x_0} positiv ist. Die Argumentation im anderen Fall verläuft analog.
3. Schritt (Eindeutigkeit): Wir nehmen nun an, dass ein offenes Zeitintervall I um 0 und eine Lösung $x(t)$ von (5.8) auf I mit $x(t) \in J_{x_0}$ für alle $t \in I$ gegeben sind. Wir zeigen, dass die Lösung eindeutig über I ist, indem wir eine Lösungsformel herleiten: Dazu trennen wir (5.8) nach den Variablen (auf der rechten Seite erscheint die Variable x nicht mehr, auf der linken Seite erscheint die Variable t nur noch als Argument der Funktion x bzw. ihrer Ableitung)

$$\frac{\dot{x}(t)}{f(x(t))} = 1,$$

integrieren bezüglich der Zeitvariable von 0 bis t

$$\int_0^t \frac{\dot{x}(\tau)}{f(x(\tau))} d\tau = \int_0^t 1 d\tau = t,$$

und erhalten mit der Substitution $z = x(s)$

$$\int_{x_0}^{x(t)} \frac{1}{f(z)} dz = t. \quad (5.20)$$

Für $x \in J_{x_0}$ definieren wir

$$T(x) := \int_{x_0}^x \frac{1}{f(z)} dz. \quad (5.21)$$

Dann ist T stetig differenzierbar mit Ableitung $T'(x) = 1/f(x) > 0$, streng monoton wachsend und hat den Wertebereich (T_-, T_+) , wobei

$$\begin{aligned} -\infty \leq T_- &:= \lim_{x \searrow x_-(x_0)} T(x) = \int_{x_0}^{x_-(x_0)} \frac{1}{f(z)} dz < 0 \\ 0 < T_+ &:= \lim_{x \nearrow x_+(x_0)} T(x) = \int_{x_0}^{x_+(x_0)} \frac{1}{f(z)} dz \leq \infty. \end{aligned}$$

Wegen $T' > 0$ besitzt T eine streng monoton wachsende differenzierbare Umkehrfunktion

$$X := T^{-1} : (T_-, T_+) \longrightarrow (x_-(x_0), x_+(x_0)),$$

für die gilt:

$$\dot{X}(t) = \frac{1}{T'(X(t))} = \frac{1}{\frac{1}{f(X(t))}} = f(X(t)) \quad (5.22)$$

Damit ist die Eindeutigkeit auf I gezeigt: Wir können (5.20) schreiben als $T(x(t)) = t$ und erhalten als eindeutige Darstellung von $x(t)$ dann $x(t) = X(t)$.

Kurzer reflektierender Einschub: Obwohl die Existenz einer Lösung gezeigt werden soll, geht der Beweis zunächst von einer vorhandenen Lösung aus und versucht diese genauer zu beschreiben. Das ist eine bekannte mathematische Denkfigur und wird *Analyse* genannt. In konkreten Beispielen wird es nicht immer möglich sein, die Umkehrfunktion $X = T^{-1}$ explizit anzugeben. In solchen Fällen kommen wir rechnerisch nicht über die Darstellung (5.20) hinaus. Aber auch die ist bereits nützlich: Wir können zwar noch nicht für gegebenes t den zugehörigen x -Wert berechnen, haben aber zu gegebenem x -Wert zwischen x_- und x_+ eine Darstellung der zugehörigen Zeit, zu der dieser x -Wert angenommen wird, nämlich gerade den orientierten Flächeninhalt von $1/f$ bezüglich des Startpunkts x_0 , siehe Abb. 5.16.

4. Schritt: Definieren wir eine Funktion $x(t)$ gemäß $x(t) = X(t)$, so liefert diese Funktion nach (5.22) und (5.21) eine Lösung des Anfangswertproblems auf (T_-, T_+) , die nach dem eben Gezeigten auf diesem Intervall eindeutig ist.

Kurzer reflektierender Einschub: In diesem Schritt wird auf der Grundlage des Ergebnisses der Analyse (Schritt 3), in der die Existenz einer Lösung angenommen wurde, die Existenz dieser Lösung gezeigt. Man spricht von der *Synthese*. Diese Abfolge von Analyse und Synthese ist ein häufig auftretendes Muster des mathematischen Arbeitens.

5. Schritt: Zu beantworten bleibt die Frage, ob die Lösung eventuell über das Intervall (T_-, T_+) hinaus fortsetzbar und ob diese Fortsetzung ebenfalls eindeutig ist. Dazu unterscheiden wir drei Fälle:
 - $T_+ = \infty$: Die Lösung existiert in positiver Zeitrichtung für alle Zeiten, die Frage nach Fortsetzbarkeit über T_+ hinaus stellt sich nicht.

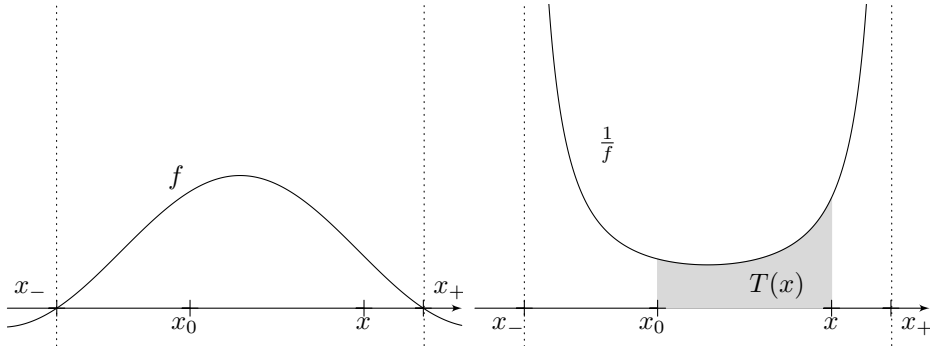


Abb. 5.16 Nullstellen und Integraldarstellung von $T(x)$

- $T_+ < \infty$ und $x_+ = \infty$: Die Lösungsfunktion hat bei T_+ eine Polstelle und kann als Funktion nicht stetig in T_+ fortgesetzt werden. Insbesondere kann sie nicht als Lösung von (5.8) fortgesetzt werden.
- $T_+ < \infty$ und $x_+ < \infty$: Das ist der kritische Fall. Hier gilt $f(x_+) = 0$ und die Lösung kann durch $x(t) = x_+$ für alle $t \geq T_+$ fortgesetzt werden. Die Eindeutigkeit der Lösungen zum Anfangswert $x_0 = x_+$ geht dabei verloren. Neben unserer Funktion wäre ja auch die komplett konstante Funktion eine Lösung. Durch die Bedingung (5.19) ist diese Möglichkeit aber ausgeschlossen.

Die analogen Betrachtungen lassen sich für T_- zusammen mit x_- anstellen. Damit ist die Lösung $X : (T_-, T_+) \rightarrow (x_-(x_0), x_+(x_0))$ maximal, eindeutig, streng monoton wachsend sowie unbeschränkt oder global.

6. Schritt: Gelte nun $f(x_0) = 0$. Dann hat (5.8) die konstante und globale Lösung $x(t) = x_0$ für alle $t \in \mathbb{R}$. Angenommen es gibt eine weitere Lösung $y(t)$ von (5.8) und ein t_1 mit $y(t_1) = x_1$ und $f(x_1) \neq 0$. Nach dem eben Gezeigten kann y nicht in eine Nullstelle von f hineinlaufen bzw. aus einer Nullstelle von f herauskommen, im Widerspruch zu $y(t_0) = x_0$ und $f(x_0) = 0$. Die Lösung muss also in der Nullstellenmenge von f verharren. Es gilt also $\dot{y}(t) = 0$ und somit $y(t) = x_0$ für alle t .

Wir ergänzen den Satz und Beweis mit drei Bemerkungen.

Bemerkung 2

- 1.) Die Voraussetzung (5.19) ist eine schwächere Bedingung als die Dehnungsbeschränktheit in allen Nullstellen von f . Dabei heißt eine Funktion f dehnungsbeschränkt in x_0 , wenn es ein $\delta > 0$ und ein $L \geq 0$ gibt, sodass für alle x mit $|x - x_0| < \delta$ die Abschätzung $|f(x) - f(x_0)| < L|x - x_0|$ folgt. Als Beispiel sei die Funktion f mit $f(x) = |\ln(|x|)|x$ für $x \neq 0$ und $f(0) = 0$ genannt. Diese Funktion ist stetig auf

- $\mathbb{R}$, erfüllt das Kriterium (5.19), ist aber nicht dehnungsbeschränkt in $x_0 = 0$ (siehe Aufg. 5.13), denn $\lim_{x \rightarrow 0} f'(x) = \infty$, andererseits gilt z. B. $\int_{1/2}^0 \frac{1}{f(z)} dz = -\infty$.
- 2.) Die Voraussetzung (5.19) ist erfüllt, wenn f dehnungsbeschränkt in allen Nullstellen von f ist, insbesondere also dann, wenn f lokal Lipschitz-stetig ist, und wiederum insbesondere dann, wenn f stetig differenzierbar ist (siehe Aufg. 5.13).
 - 3.) Der Beweis lässt sich mutatis mutandis auch auf nichtautonome Differentialgleichungen mit getrennten Variablen $\dot{x} = f(x)g(t)$ mit einer stetigen Funktion g übertragen. Natürlich sind die Lösungen dann i.A. nicht monoton (siehe Aufg. 5.15).
 - 4.) Mit der Methode der Trennung der Variablen lassen sich auch viele Differentialgleichungen explizit lösen, unter anderen auch die logistische Differentialgleichung (5.9). Entscheidend dabei ist, ob sich der Ausdruck $T(x)$ integralfrei darstellen und die Umkehrfunktion von T explizit berechnen lässt. Einige Beispiele sollen in Aufg. 5.16 berechnet werden.

5.5.2 Grobe Bestimmung der Zeitabhängigkeit und Vergleichssätze

Auch wenn es häufig hinreichend, ist den stabilen Endzustand eines Systems zu kennen, möchte man doch auch manchmal wissen, wie lange es z. B. dauert, bis man sich diesem Endzustand hinreichend gut angenähert hat. Das Ziel ist also, die qualitativen Betrachtungen, die lediglich die stationären Punkte und ihre Stabilität liefern, mit quantitativen Betrachtungen zu verknüpfen, die etwas über die Zeitabhängigkeit der Lösungen aussagen, ohne dass die Lösung explizit berechnet werden muss.

Eine andere Frage qualitativer Natur, die wir noch nicht beantworten können, ist die nach der Lebensdauer einer Lösung, die nicht gegen einen stationären Punkt konvergiert. Wir haben in Abschn. 5.5 gesehen, dass z. B. im Fall $T_+ < \infty$ und $x_+ = \infty$ die Lösung einer Differentialgleichung in endlicher Zeit unendlich groß werden kann. Man spricht von einem *Blow-up in endlicher Zeit*. Im Modellierungskontext bedeutet dies, dass eine Größe in endlicher Zeit unendlich groß wird. In der Regel, aber nicht immer, deutet das auf eine Limitierung des Modells hin. Wir stellen uns in diesem Kapitel die Frage, wie wir erkennen können, ob eine Lösung einer Differentialgleichung einen Blow-up in endlicher Zeit aufweist, ohne die Lösung explizit berechnen zu müssen.

Die Berechnung der Rekonvaleszenzzeit in Abschn. 5.3.1 hat einen Schritt in Richtung der Beantwortung der ersten Frage unternommen, den wir jetzt systematisieren wollen. Dazu wollen wir Lösungen von (5.8) mit geschickt gewählten Lösungen linearer und anderer einfacher Gleichungen vergleichen, deren Lösungen wir explizit berechnen können. Wir werden ein Beispiel durcharbeiten, dabei den benötigten Satz angeben und anschließend beweisen. Wir starten dazu wieder mit dem logistischen Grundmodell (5.9) und fragen uns, wie lange eine Lösung braucht, die klein mit dem Wert x_0 startet, um sehr nah an den stabilen stationären Punkt K zu kommen. Dazu stellen wir zwei lineare Funktionen g_1 und g_2 auf, so dass $g_1(x) \leq f(x)$ für alle $x_0 \leq x \leq \frac{K}{2}$ und $g_2(x) \leq f(x)$ für alle

$\frac{K}{2} \leq x \leq K$, z. B. $g_1(x) = r/2x$ und $g_2(x) = r/2(K - x)$ (siehe Abb. 5.17). Wir benötigen nun einen Satz, der uns erlaubt, Lösungen von Differentialgleichungen mit sogenannten Unter- und Oberlösungen zu vergleichen. Dabei nennen wir eine Funktion $I \ni t \mapsto y(t)$ eine Unterlösung der Differentialgleichung $\dot{x} = f(t, x)$ auf dem Intervall I , wenn für alle $t \in I$ die Ungleichung $\dot{y}(t) \leq f(t, y(t))$ gilt, und sprechen von einer Oberlösung, wenn $\dot{y}(t) \geq f(t, y(t))$ für alle $t \in I$. Wir würden jetzt gerne schließen, dass eine Unterlösung immer unterhalb einer Oberlösung bleibt, wenn sie unterhalb der Oberlösung startet.

Das Gegenbeispiel $\dot{x} = 2\sqrt{|x|}$, $x(0) = 0$, zur eindeutigen Lösbarkeit aus Abschn. 5.5 zeigt aber, dass es für diese Eigenschaft nicht hinreichend ist, lediglich die Stetigkeit von f zu fordern, es müsste sonst dieses Anfangswertproblem ebenfalls eindeutig lösbar sein: Man könnte zwei Lösungen y_1 und y_2 als Unter- und Oberlösung gegeneinander antreten lassen und erhielte $y_1(t) \leq y_2(t)$. Durch Vertauschen der Rollen von y_1 und y_2 könnte man $y_2(t) \leq y_1(t)$ folgern, also $y_1 = y_2$.

Die „richtige“ Eigenschaft wird durch den Begriff der *lokalen Lipschitz-Stetigkeit* definiert. Um nicht nur autonome Differentialgleichungen behandeln zu können, definieren wir direkt etwas allgemeiner.

Definition 5.1

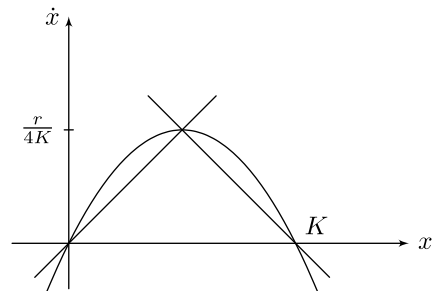
Ein stetige Funktion $f : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ heißt lokal Lipschitz-stetig in der zweiten Variablen, wenn es für jede kompakte Menge $K \subset \mathbb{R} \times \mathbb{R}$ eine Konstante L_K gibt, sodass für alle $(t, x_1), (t, x_2) \in K$ die folgende Abschätzung gilt:

$$|f(t, x_1) - f(t, x_2)| \leq L_K |x_1 - x_2|$$



Man könnte auch sagen, die Funktion f ist lokal sublinear, lässt sich also lokal immer gegen eine geeignete lineare Funktionen abschätzen. Ein einfaches hinreichendes Kriterium für lokale Lipschitz-Stetigkeit liefert die stetige Differenzierbarkeit. Das soll in Aufg. 5.13 gezeigt werden.

Abb. 5.17 Logistisches Modell mit Vergleichsfunktionen



Satz 5.2 (Vergleichssatz für Ober- und Unterlösungen)

Gegeben seien eine stetige Funktion $f : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$, die lokal Lipschitz-stetig bezüglich der zweiten Variable ist, ein Intervall $I = [t_0, t_1]$ und zwei stetig differenzierbare Funktionen $y, z : I \rightarrow \mathbb{R}$ mit

$$\begin{aligned} \dot{y}(t) &\leq f(t, y(t)) \quad \text{für alle } t \in I, \quad y \text{ ist also eine Unterlösung von } \dot{x} = f(t, x), \\ \dot{z}(t) &\geq f(t, z(t)) \quad \text{für alle } t \in I, \quad z \text{ ist also eine Oberlösung von } \dot{x} = f(t, x), \\ y(t_0) &\leq z(t_0). \end{aligned}$$

Dann gilt

$$y(t) \leq z(t) \quad \text{für alle } t \in I.$$

Bevor wir diesen Satz beweisen, wollen wir ihn erst einmal anwenden: Wir betrachten die Lösung $x(t)$ von (5.9) zu einem kleinen positiven Startwert $x_0 < \frac{K}{2}$ zur Zeit $t_0 = 0$. Wir wissen, dass x streng monoton gegen K konvergiert. Es gibt also ein $t_1 > 0$ mit $x(t_1) = \frac{K}{2}$. Für alle $t \in [0, t_1]$ gilt $\dot{x}(t) = f(x(t)) \geq g_1(x(t))$. Andererseits ist die Funktion $y(t) = x_0 e^{r/2t}$ die Lösung von $\dot{y} = g_1(y)$ zum Anfangswert $y(0) = x_0$. Nach dem Vergleichssatz gilt $x(t) \geq x_0 e^{r/2t}$, solange $x(t) \leq \frac{K}{2}$. Die Unterlösung y erreicht $\frac{K}{2}$ zur Zeit $t_2 = \frac{2}{r} \ln(\frac{K}{2x_0})$ und t_1 muss dementsprechend kleiner sein.

Starten wir nun zur Zeit $t = 0$ bei einem $x_0 \geq \frac{K}{2}$, können wir gegen die Lösung von $\dot{y} = g_2(y)$, $y(0) = x_0$ abschätzen: $K > x(t) \geq K - (K - x_0)e^{-\frac{r}{2}t}$ für alle $t \geq 0$.

Nun zum Beweis von Satz 5.2. Zur Vorbereitung definieren wir den positiven Teil einer reellen Funktion: Für irgendeine Menge M und eine Funktion $\varphi : M \rightarrow \mathbb{R}$ definieren wir $\varphi_+(m) = \varphi(m)$, falls $\varphi(m) \geq 0$, und $\varphi_+(m) = 0$ sonst.

Beweis 5.2 Es sei K die Bahn der beiden Lösungen,

$$K = \{(t, y(t)) : t \in I\} \cup \{(t, z(t)) : t \in I\}.$$

Dann ist K kompakt und L_K sei eine zugehörige Lipschitz-Konstante von f . Wir definieren für $t \in I$

$$d(t) = e^{-2Lt} |y(t) - z(t)|_+^2.$$

Dann gilt $d(t_0) = 0$. Ferner ist d differenzierbar (man beachte, dass $|\cdot|_+^2$ auch in 0 differenzierbar ist) und für die Ableitung gilt:

$$\begin{aligned} \dot{d}(t) &= e^{-2Lt} \left(-2L |y(t) - z(t)|_+^2 + 2 |y(t) - z(t)|_+ \cdot (\dot{y}(t) - \dot{z}(t)) \right) \\ &\leq e^{-2Lt} \left(-2L |y(t) - z(t)|_+^2 + 2 |y(t) - z(t)|_+ \cdot (f(t, y(t)) - f(t, z(t))) \right) \\ &\leq e^{-2Lt} \left(-2L |y(t) - z(t)|_+^2 + 2L |y(t) - z(t)|_+ \cdot |y(t) - z(t)| \right) \\ &= 0 \end{aligned}$$

Also folgt $d(t) = 0$ für alle $t \in I$ und damit die Behauptung des Satzes.

Lösungen mit und ohne Blow-up

Um mit Hilfe des Vergleichssatzes zu entscheiden, ob eine Lösung einer Differentialgleichung einen Blow-up in endlicher Zeit entwickelt, ist es hilfreich, geeignete Differentialgleichungen mitsamt Lösungen an der Hand zu haben, die gerade einen Blow-up entwickeln oder eben gerade nicht. Prototypen sind die Lösungen der Differentialgleichungen

$$\dot{x} = x^\alpha, \quad x(0) = x_0 > 0$$

mit $0 < \alpha$. Diese können wir mit der Methode der Trennung der Variablen (die wir bis jetzt nur für theoretische Belange genutzt haben) explizit berechnen. Für die Lösung finden wir

$$\begin{aligned} t &= \int_0^t \frac{\dot{x}(s)}{x(s)^\alpha} ds = \int_{x_0}^{x(t)} z^{-\alpha} dz \\ &= \begin{cases} \ln\left(\frac{x(t)}{x_0}\right) & \text{falls } \alpha = 1 \\ \frac{1}{1-\alpha} \left(x^{1-\alpha}(t) - x_0^{1-\alpha}\right) & \text{falls } \alpha \neq 1. \end{cases} \end{aligned}$$

In beiden Fällen können wir nach $x(t)$ auflösen und erhalten

$$x(t) = \begin{cases} x_0 e^t & \text{falls } \alpha = 1 \\ ((1-\alpha)t + x_0^{1-\alpha})^{\frac{1}{1-\alpha}} & \text{falls } \alpha \neq 1. \end{cases}$$

Im Fall $\alpha \leq 1$ existiert die Lösung für alle $t \geq 0$, ist also global. Im Fall $1 < \alpha$ entwickelt die Lösung einen Blow-up in endlicher Zeit: Mit $\hat{t} := \frac{x_0^{1-\alpha}}{\alpha-1}$ gilt $\lim_{t \rightarrow \hat{t}} x(t) = \infty$.

Ist jetzt z. B. eine Lösung $t \mapsto x(t)$ der Differentialgleichung $\dot{x} = f(x)$ mit dem Anfangsdatum $x(0) = x_0 > 0$ gegeben und gilt $f(x) \geq x^2$ für alle $x > 0$, dann muss diese Lösung nach dem Vergleichssatz einen Blow-up entwickeln. Für den Zeitpunkt $\tilde{t}$ des Blow-up muss außerdem $\tilde{t} \leq x_0^{-1}$ gelten.

5.6 Ein paar ausgewählte Lösungsverfahren

Die Herangehensweise in diesem Buch besteht weniger darin, explizite Lösungen von Differentialgleichungen auszurechnen, was in vielen Fällen auch gar nicht möglich ist, sondern qualitative Aussagen über die Lösungen von Differentialgleichungen zu treffen. Trotzdem werden einige Lösungen von einfachen Differentialgleichungen benötigt, z. B. zur Verwendung des Vergleichssatzes 5.2 als Unter- oder Oberlösung. In diesem Kapitel sollen explizite Lösungen für inhomogenen lineare Differentialgleichungen entwickelt, sowie – als ein Beispiel für die Mannigfaltigkeit der Substitutionsmethoden – die Lösungen der Bernoulli'schen Differentialgleichungen vorgestellt werden.

5.6.1 Die inhomogene lineare Differentialgleichung

Wir betrachten Differentialgleichungen der Form

$$\dot{x} = a(t)x + s(t).$$

Dabei sollen a und s stetige Funktionen von $\mathbb{R}$ nach $\mathbb{R}$ sein. Auch so eine Gleichung passt in das lineare Lösungsschema in Abschn. 3.4.1.

Die Lösungen der homogenen linearen Differentialgleichung

Wir gehen also in den bewährten Schritten vor und bestimmen zunächst die Lösungen der homogenen Gleichungen

$$\dot{x}^{\text{hom}} = a(t)x^{\text{hom}}. \quad (5.23)$$

Dazu nutzen wir die Methode von der Trennung der Variablen. Die Lösung $t \mapsto x(t)$ mit den Anfangsdaten (t_0, x_0) erfüllt dann die Gleichung

$$\int_{t_0}^t a(s) ds = \int_{t_0}^t \frac{\dot{x}(s)}{x(s)} ds = \ln \left(\frac{x(t)}{x_0} \right). \quad (5.24)$$

Auflösen nach $x(t)$ liefert

$$x(t) = x_0 \exp \left(\int_{t_0}^t a(s) ds \right).$$

Aus diesem Ergebnis folgt, dass die Gesamtheit der Lösungen des homogenen Problems (5.23) folgendermaßen ausgedrückt werden kann:

$$x^{\text{hom}}(t) = ce^{A(t)},$$

wobei $c \in \mathbb{R}$ und A eine Stammfunktion von a ist.

Eine Lösung der inhomogenen linearen Differentialgleichung

In vielen Fällen, insbesondere wenn die Funktion a konstant ist, lässt sich eine partikuläre Lösung mit Hilfe eines Ähnlichkeitsansatzes finden. Wir diskutieren zwei häufig auftretende Fälle. Sollte kein Ähnlichkeitsansatz erfolgreich sein, so funktioniert doch immer die *Methode von der Variation der Konstanten*, die wir anschließend vorstellen.

Ähnlichkeitsansätze Wir betrachten die inhomogene lineare Differentialgleichung mit konstanten Koeffizienten $\dot{x} = ax + s$ mit $a, s \in \mathbb{R}$. Nach dem eben Gesagten sind die homogenen Lösungen durch $x^{\text{hom}}(t) = ce^{at}$ mit $c \in \mathbb{R}$ gegeben. Die Inhomogenität ist konstant. Wir setzen nach dem Ähnlichkeitsprinzip auch die partikuläre Lösung als Konstante an $x^{\text{part}}(t) = k$ mit $k \in \mathbb{R}$. Wir setzen diesen Ansatz in die Differentialgleichung ein,

$$0 = \dot{x}^{\text{part}}(t) \stackrel{!}{=} ax^{\text{part}}(t) + s = ak + s, \quad (5.25)$$

und sehen, dass diese Gleichung genau dann erfüllt ist, wenn $k = -\frac{s}{a}$ gilt. Die allgemeine Lösung ist also gegeben durch

$$x^{\text{all}}(t) = ce^{at} - \frac{s}{a} \quad \text{mit } c \in \mathbb{R}.$$

Für ein zweites Beispiel betrachten wir die Gleichung

$$\dot{x} = ax + b \exp(ct)$$

mit Konstanten $a, b, c \in \mathbb{R}$. Die Lösungen der homogenen Gleichung kennen wir schon. Für die partikuläre Lösung machen wir nun nach dem Ähnlichkeitsprinzip den Ansatz $x^{\text{part}} = k \exp(ct)$. Einsetzen liefert

$$kc \exp ct = \dot{x}^{\text{part}}(t) \stackrel{!}{=} ak \exp(ct) + b \exp(ct) \quad \text{für alle } t \in \mathbb{R}.$$

Diese Gleichung ist genau dann erfüllt, wenn $kc = ak + b$, im Fall $c \neq a$ also genau dann, wenn $k = \frac{b}{c-a}$. In diesem Fall ist die allgemeine Lösung gegeben durch

$$x^{\text{all}}(t) = \tilde{c} \exp(at) + \frac{b}{c-a} \exp(ct) \quad \text{mit } \tilde{c} \in \mathbb{R}.$$

Im Fall $a = c$ funktioniert dieser Ansatz jedoch nicht (das gleiche Phänomen ist uns auch schon bei den inhomogenen linearen Differenzengleichungen begegnet), weil die Ansatzfunktion bereits die homogene Gleichung löst. Wir gehen entsprechend zum Differenzengleichungsfall vor und modifizieren unseren Ansatz zu

$$x^{\text{part}}(t) = kt \exp(at).$$

Einsetzen in die Differentialgleichung liefert

$$\exp(at)(akt + k) = \dot{x}^{\text{part}}(t) \stackrel{!}{=} akt \exp(at) + b \exp(at).$$

Diese Gleichung ist genau dann erfüllt, wenn $k = b$ gilt. Wir haben also die Lösungsgesamtheit

$$x^{\text{all}}(t) = c \exp(at) + bt \exp(at) \quad \text{mit } c \in \mathbb{R}.$$

Das Prinzip von der Variation der Konstanten Wenn keine Ähnlichkeitsansätze zum Erfolg führen, funktioniert doch immer das Prinzip von der Variation der Konstanten. Auch dieses fußt auf einem Ansatz: Für die partikuläre Lösung nehme man die homogene Lösung, wähle aber für jeden Zeitpunkt t geschickt eine Konstante $c(t)$, $x^{\text{part}}(t) = c(t) \exp(A(t))$, wobei A eine fest gewählte Stammfunktion von a ist (gleichbedeutend könnte man auch die Stammfunktion variieren). Wir setzen diesen Ansatz in die Differentialgleichung ein und erhalten

$$\exp(A(t))(\dot{c}(t) + a(t)c(t)) = \dot{x}^{\text{part}}(t) \stackrel{!}{=} a(t)x^{\text{part}}(t) + b(t) = a(t)c(t)\exp(A(t)) + b(t)$$

für alle $t \in \mathbb{R}$. Diese Gleichung ist genau dann erfüllt, wenn

$$\dot{c}(t)\exp(A(t)) = b(t) \quad \text{bzw.} \quad \dot{c}(t) = b(t)\exp(-A(t)) \quad \text{für alle } t \in \mathbb{R}$$

gilt. Ist also c irgendeine Stammfunktion von $t \mapsto b(t)\exp(-A(t))$, so können wir die Lösungsgesamtheit ausdrücken als

$$x^{\text{all}}(t) = \tilde{c}\exp(A(t)) + c(t)\exp(A(t)) \quad \text{mit } \tilde{c} \in \mathbb{R}. \quad (5.26)$$

Für das Anfangswertproblem $x(t_0) = x_0$ können wir auch direkt die richtigen Stammfunktionen für A und c wählen und erhalten die Lösungsformel

$$x(t) = x_0 \exp\left(\int_{t_0}^t a(s) ds\right) + \int_{t_0}^t b(s) \exp\left(-\int_{t_0}^s a(\sigma) d\sigma\right) ds.$$

5.6.2 Die Bernoulli'sche Differentialgleichung als Beispiel für ein Substitutionsverfahren

Mitunter lassen sich die Lösungen nichtlinearer Differentialgleichungen durch eine geschickte Substitution auf die Lösung linearer Differentialgleichungen zurückführen. Wir betrachten hier als prominentes Beispiel die Bernoulli'schen Gleichungen, welche die Bauart

$$\dot{x} = a(t)x + b(t)x^\alpha \quad (5.27)$$

mit stetigen Funktionen a und b und einer Konstante $\alpha \neq 1$ aufweisen. Wir nehmen nun an, wir hätten eine Lösung $t \mapsto x(t)$ dieser Differentialgleichung auf einem Intervall I für alle $t \in I$. Wir definieren die Funktion z vermöge $z(t) := (x(t))^{1-\alpha}$ (im Fall $\alpha > 1$ müssen wir $x(t) \neq 0$ für alle $t \in I$ voraussetzen) und rechnen nach, welche Differentialgleichung z löst (stillschweigend unterdrücken wir, wie es häufig üblich ist, das Argument in den Funktionen x und z , aber nicht in den Koeffizientenfunktionen a und b). Mit der Kettenregel folgt

$$\dot{z} = (1-\alpha)x^{-\alpha}\dot{x} = (1-\alpha)a(t)x^{1-\alpha} + (1-\alpha)b(t) = (1-\alpha)a(t)z + (1-\alpha)b(t). \quad (5.28)$$

Die Funktion z löst also eine inhomogene lineare Differentialgleichung, die wir auch explizit lösen könnten, wenn $x(t)$ noch nicht vorläge. Umgekehrt gilt: Löst z die Differentialgleichung $\dot{z} = a(t)z + b(t)$ auf einem Intervall I , dann löst die Funktion $x(t) := (z(t))^{\frac{1}{1-\alpha}}$ die ursprüngliche Differentialgleichung (5.27). Auch hier müssen wir im Fall $\alpha > 2$ die Voraussetzung $z(t) \neq 0$ fordern. Das ist die interessante Richtung: Um die Lösung x von (5.27) mit $x(t_0) = x_0$ zu erhalten, berechne man die Lösung z von (5.28) mit $z(t_0) = x_0^{1-\alpha}$ und setze $x(t) = z(t)^{1/(1-\alpha)}$.

5.7 Ein einfaches Verfahren zur numerischen Lösung von Anfangswertproblemen

In einigen Fällen werden uns rein qualitative Aussagen über die Lösungen von Differentialgleichungen nicht ausreichen, wir werden aber auch nicht in der Lage sein, explizite Lösungen mit analytischen und algebraischen Methoden zu berechnen. In diesem Fall sind wir auf numerische Verfahren angewiesen. Für die Belange dieses Buches begnügen wir uns mit den beiden einfachsten Verfahren, dem expliziten und dem impliziten Euler-Verfahren. Für alle Fragen bezüglich der Konvergenz und Stabilität solcher Verfahren verweisen wir auf die Literatur, z. B. [19]. Auch hier betrachten wir direkt den allgemeineren nichtautonomen Fall einer Differentialgleichung

$$\dot{x} = f(t, x) \quad \text{mit Anfangswert} \quad x(t_0) = x_0, \quad (5.29)$$

der numerisch keine zusätzlichen Schwierigkeiten bereitet. Die Lösung dieses Problems lässt sich in einem sogenannten *Richtungsfeld* veranschaulichen. Dazu wird an jeden Punkt $(\hat{t}, \hat{x})$ in der t - x -Ebene ein Streckenstück mit Einheitslänge und der Steigung $f(\hat{t}, \hat{x})$ angehängt. Aufgrund der Differentialgleichung verläuft die Lösung $t \mapsto x(t)$ von (5.29) jederzeit tangential an dieses Richtungsfeld (siehe Abb. 5.18).

Beim Euler-Verfahren wird die Lösung aus dem Richtungsfeld (also aus der Differentialgleichung) approximativ rekonstruiert: Hat die Lösung zur Zeit t_0 den Wert x_0 , so wird sie sich einen kurzen Zeitschritt Δt später, zur Zeit $t_1 = t_0 + \Delta t$, nach der linearen Approximation ungefähr an der Stelle $x_1 = x_0 + \Delta t \cdot \dot{x}(t_0) = x_0 + \Delta t f(t_0, x_0)$ befinden. Dieses Verfahren lässt sich iterieren: Wir berechnen Annäherungen x_n an den Stützstellen $t_n := t_0 + n\Delta t$ über das Iterationsschema (eine Differenzengleichung)

$$x_{n+1} = x_n + \Delta t f(t_n, x_n). \quad (5.30)$$

Unangenehmerweise geht so der Fehler im n -ten Schritt direkt als Fehler im Anfangsdatum in den $n + 1$ -ten Schritt ein. Fehler summieren sich also mit der Zeit auf. Um den Fehler klein halten zu können, muss also eine kleine Schrittweite gewählt werden, was viele Rechenschritte erfordert, um die Lösung über ein vorgegebenes Zeitintervall zu verfolgen. Das Iterationsschema (5.30) definiert das *explizite Euler-Verfahren*.

Das *implizite Euler-Verfahren* unterscheidet sich dadurch, dass das Richtungsfeld an der neuen Stelle x_{n+1} ausgewertet wird:

$$x_{n+1} = x_n + \Delta t f(t_n, x_{n+1}) \quad (5.31)$$

Durch diese Gleichung ist x_{n+1} nur *implizit* bestimmt. In der Regel muss diese implizite Gleichung in jedem Iterationsschritt noch einmal iterativ, z. B. mit dem Newton-Verfahren, bestimmt werden. Das implizite Euler-Verfahren hat den Vorteil, dass es insbesondere in der Umgebung von stationären Punkten größere Schrittweiten Δt zulässt. Dieser Vorteil

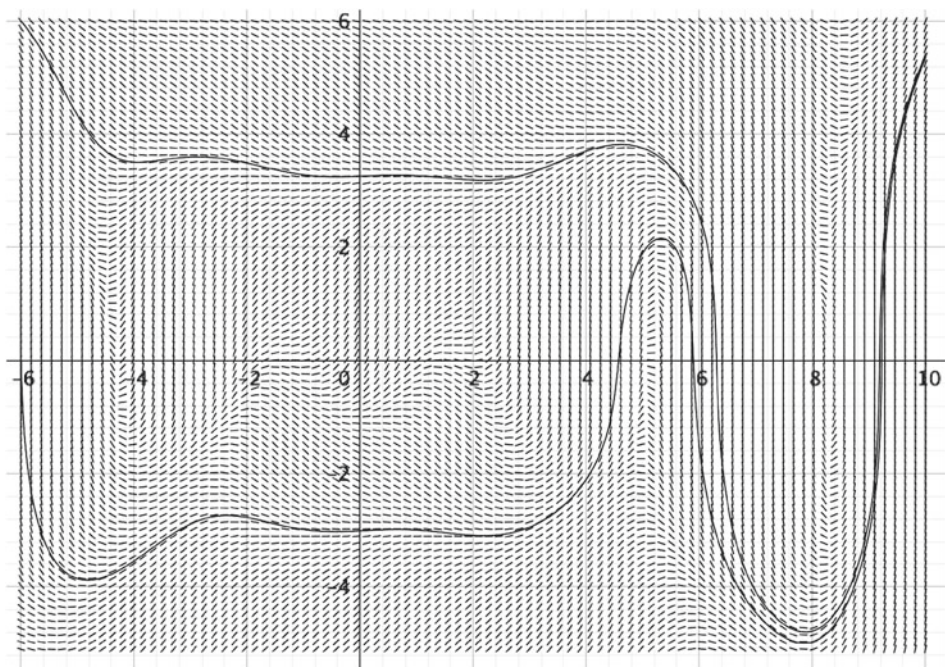


Abb. 5.18 Das Richtungsfeld der Differentialgleichung $f(t, x) = (t^2 - 2) \frac{\sin(t-2)}{1+x^2} + \sin(x)$, zusammen mit den numerisch berechneten Approximationen zu den Anfangswerten $x(-6) = 0$ und $x(-6) = 6$; Schrittweite $\Delta t = 0,01$ (Abbildung erstellt mit GeoGebra)

wiegt häufig die zusätzlichen inneren Iterationen in jedem äußeren Iterationsschritt auf. In Aufg. 5.19 sollen Sie etwas Erfahrung mit diesen beiden Verfahren erlangen.

5.8 Modellierungen in Abituraufgaben – Blutalkohol

Heutzutage spielen Kontextbezüge in den Mathematikaufgaben des Zentralabiturs eine prominente Rolle, sind jedoch nicht unumstritten. Auf der einen Seite werden sie als „beispielhaft für eine an einem breiteren Kompetenzspektrum orientierte und somit sowohl anspruchsvollere als auch sinnhaftere Aufgabenkultur“ angesehen, siehe [10]. Andererseits wird von „sogenannte[n] Textaufgaben oder auch ‚Sachaufgaben‘ mit zum Teil sehr langem Text und mit gezielt und fast krampfhaft herangezogenen Anwendungssituationen“ gesprochen, siehe [49], für deren Bearbeitung wirkliche Modellierungskompetenzen nicht erforderlich seien, Standardverfahren seien völlig ausreichend, siehe [57].

Wir wollen hier diese Diskussion nicht weiter vertiefen, aber den Kontext einer Anwendungsaufgabe aus dem Abitur für den Leistungskurs Mathematik in NRW aus dem Jahre 2012 heranziehen, um ein eigenständiges Modell aufzustellen und dieses mit dem in der

Abituraufgabe verwendeten Modell zu vergleichen. Weitere Kontexte aus Abituraufgaben werden in diesem Sinne in [20] untersucht.

5.8.1 Die Formulierung der Abituraufgabe

Die Abituraufgabe wird zunächst in den interessierenden Auszügen wiedergegeben.

„Bei einem medizinischen Test leert eine Versuchsperson ein Glas Wein in einem Zug. Anschließend wird der zeitliche Verlauf der Blutalkoholkonzentration (in Promille) aufgezeichnet. Diese wird im hier verwendeten Modell zunächst durch eine Funktion f_a mit der Gleichung

$$f_a(t) = \frac{a}{60} \left(1 - e^{-\frac{1}{20}t} \right) - \frac{1}{600}t$$

beschrieben. Dabei ist a die Alkoholmenge in Gramm und t die Zeit in Minuten, die nach der Alkoholaufnahme vergangen ist. Die Funktion f_a ist für alle $a \in \mathbb{R}$ und alle $t \in \mathbb{R}$ definiert. Zur Modellierung ist die Funktion für $a > 2$ und eine gewisse Zeitspanne geeignet. In der Abbildung 1 [hier Abb. 5.19] ist der Graph der Funktion f_{20} dargestellt.“

Nach einigen Aufgaben, in denen der Zeitpunkt der höchsten Alkoholkonzentration (Maximalstelle), die höchste Alkoholkonzentration (Maximum), die mittlere Alkoholkonzentration (Integralmittel) und die Konzentration nach 140 min (Funktionswert) berechnet werden sollen, wird die Modellfunktion $f = f_{20}$ für $t \geq 140$ modifiziert:

„Aus biologischen Gründen wird nach 140 Min. die Blutalkoholkonzentration der Versuchsperson durch die Funktion f nicht mehr richtig beschrieben. Für die Modellierung besser geeignet ist die an der Stelle $t = 140$ zusammengesetzte Funktion h mit der Gleichung

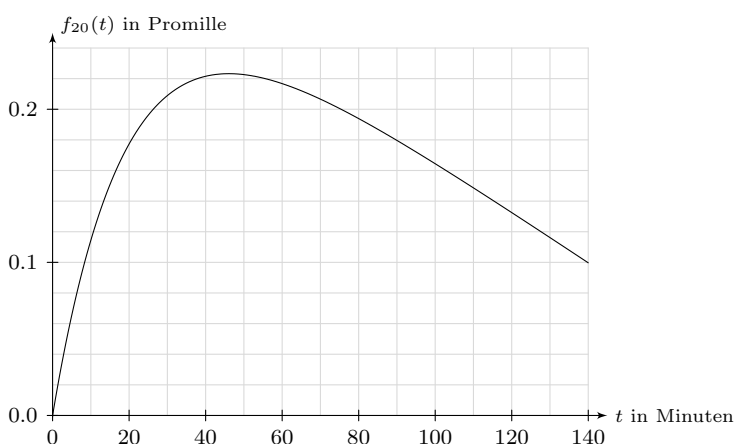


Abb. 5.19 nach Abb. 1 aus der Abituraufgabe

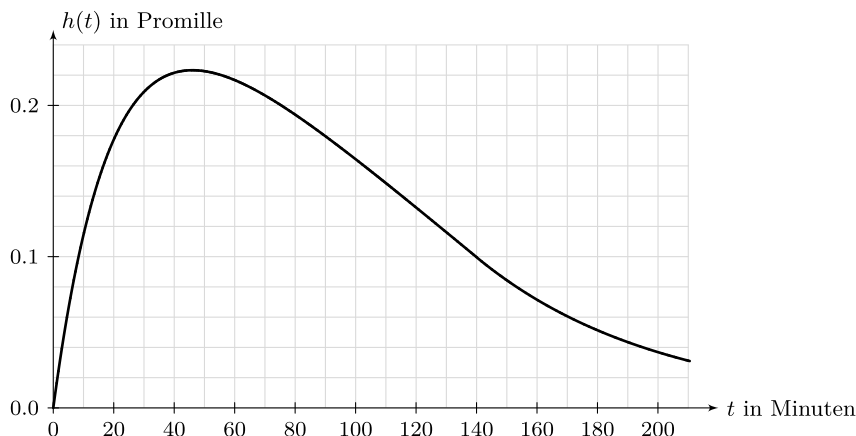


Abb. 5.20 nach Abb. 2 aus der Abituraufgabe

$$h(t) = \begin{cases} f(t), & 0 \leq t \leq 140 \\ g(t), & 140 < t \end{cases}$$

mit $g(t) = u \cdot e^{-\nu t}$, wobei $u > 0$ und $\nu > 0$ geeignet zu wählen sind (siehe Abbildung 2 [hier Abb. 5.20]).“

In der Aufgabe müssen nun u und ν so bestimmt werden, dass die zusammengesetzte Funktion differenzierbar ist, die Kontrollwerte $u \approx 1,01357$ und $\nu \approx 0,01657$ werden angegeben. Anschließend wird die Frage nach dem Zeitpunkt der schnellsten Konzentrationsänderung gestellt.

Aus Sicht des mathematischen Modellierens ist eigentlich eine andere Frage interessant, nämlich: Durch welche Überlegungen kommt man auf die angeführte Modellfunktion? Zur Beantwortung dieser Frage modellieren wir die angeführte Situation mit Hilfe von Differentialgleichungen und prüfen, ob wir auf eine ähnliche Modellfunktion kommen.

5.8.2 Die Modellierung der Alkoholkonzentration

Wir stellen zunächst einige sehr prinzipielle Überlegungen an, und zwar, mit welchen Mechanismen sich die Menge eines bestimmten Stoffes (hier von Alkoholmolekülen) in einem bestimmten Bereich des Körpers ändern kann. Wir unterteilen dazu den gesamten Körper in passende Subsysteme. Für den aktuellen Kontext sind die beiden folgenden Subsysteme sinnvoll:

- der Magen-Darm-Trakt. In dieses Subsystem wird der Alkohol durch Trinken eingebracht.

- die Körperflüssigkeit. Da sich die Alkoholmoleküle auf die gesamte Körperflüssigkeit verteilen, das sogenannte *Verteilungsvolumen*, betrachten wird dieses System und nicht nur den Blutkreislauf.

Der Rest des Körpers, Knochen, Sehnen und Ähnliches, spielt für die Entwicklung der Alkoholkonzentration keine Rolle. Bezogen auf ein Subsystem gibt es drei grundsätzliche Mechanismen, durch die sich die Anzahl der Alkoholmoleküle und damit die Alkoholkonzentration in dem betrachteten Subsystem ändern kann:

1. Zufuhr und Abtransport in bzw. aus dem Subsystem von bzw. nach außen,
2. Abbau und Erzeugung von Alkoholmolekülen innerhalb eines Subsystems,
3. Übergang von Alkoholmolekülen von einem Subsystem in das andere durch die abgrenzenden Trennflächen.

Bevor wir die einzelnen Mechanismen bezogen auf die beiden Subsysteme diskutieren, treffen wir die folgende *Grundannahme der Pharmakokinetik*: In jedem Subsystem verteilen sich die einzelnen Stoffe, in unserem Fall die Alkoholmoleküle, instantan gleichmäßig (instantan bedeutet: ohne dass Zeit vergeht). Offenkundig ist das nicht ganz wahr: Natürlich wird es einige Zeit dauern, bis Alkohol auch in der Flüssigkeit innerhalb der letzten Fußspitze angekommen ist. Auf der anderen Seite vereinfachen wir uns die Aufgabe gewaltig, weil wir innerhalb der Subsysteme von der Ortsabhängigkeit der Alkoholkonzentration absehen können. Wir führen nun unsere beiden Modellgrößen ein:

- $D(t)$: die Masse des Alkohols im Magen-Darm-Trakt zur Zeit t ,
- $F(t)$: die Konzentration des Alkohols in der Körperflüssigkeit zur Zeit t . Diese geben wir, wie es üblich ist, als Gewichtsanteil an: Eine Konzentration von 1 ‰ bedeutet etwa, dass eine Masse von 1 Gramm Alkohol auf 1 Kilogramm Körperflüssigkeit kommt.

Warum benutzen wir im Magen-Darm-Trakt die Masse und nicht die Konzentration? Weil im Aufgabenkontext die anfängliche Alkoholdosis auch als Masse angegeben wird!

Wir diskutieren nun die angesprochenen drei Mechanismen:

1. Für den ersten Mechanismus nehmen wir an, dass er nach der anfänglichen Aufnahme in beiden Subsystem keine Rolle mehr spielt, es also zu keiner weiteren relevanten Nahrungs- oder Alkoholaufnahme oder zu Erbrechen/Stuhlgang kommt (Magen-Darm-Trakt) bzw. keine Infusionen stattfinden oder Alkohol und Körperflüssigkeit in relevanten Mengen durch Urin oder Schweiß das Subsystem verlassen.
2. Wir gehen davon aus, dass Alkohol in keinem der beiden Subsysteme aus anderen Stoffen neu synthetisiert wird. Es spielen also nur Abbauvorgänge eine Rolle.

- Ein wichtiger Abbauvorgang ist der *spontane Zerfall*, bei dem ein Stoff ohne weitere Einwirkungen von außen in andere Stoffe zerfällt. Ein Beispiel dafür haben wir bei dem radioaktiven Zerfall kennengelernt. Die zugehörige Modellierung mit Hilfe einer Differentialgleichung für einen Stoff X , dessen Menge durch die Größe x modelliert wird, lautet in der Regel $\dot{x} \sim x$, also

$$\dot{x} = -\alpha x \quad (5.32)$$

mit einer positiven Konstante α : Wenn die doppelte Menge an Stoff da ist, zerfällt auch die doppelte Menge pro Zeiteinheit. Da das Alkoholmolekül jedoch recht stabil ist, spielt dieser Mechanismus in unserem Kontext keine Rolle.

- In der Körperflüssigkeit findet eine *enzymatische Reaktion* statt: In der Leber treffen Alkoholmoleküle auf das dort vorliegende Enzym Alkoholhydrogenase. Die beiden Moleküle verbinden sich beim Aufeinandertreffen zu einem so genannten Komplex, der wiederum instabil ist und spontan zerfällt, und zwar wieder in das Enzym, das also unverändert aus der Reaktion hervorgeht, und Abbauprodukte des Alkoholmoleküls. Im Rahmen der Reaktionskinetik werden wir diese Reaktion in Abschn. 7.2 genauer analysieren, wollen hier jedoch ein anschaulich plausibles Modell aufstellen: Die Rate, mit der Alkohol durch diese enzymatische Reaktion abgebaut wird, nennen wir A . Dabei ist A abhängig von der vorliegenden Alkoholkonzentration F : Wenn die Alkoholkonzentration niedrig ist, also das Verhältnis von der Anzahl der Enzyme zur Anzahl der Alkoholmoleküle in der Leber groß ist, werden bei Verdoppelung der Alkoholmoleküle doppelt so viele pro Zeiteinheit gespalten, jedes Alkoholmolekül findet direkt ein freies Enzym. Für kleine Konzentrationen nehmen wir also einen proportionalen Zusammenhang an, $A \sim F$. Bei großen Alkoholkonzentrationen hingegen sind zu jeder Zeit so gut wie alle Enzyme in einem Komplex gebunden, und dieser muss erst zerfallen, bevor das Enzym sich mit einem neuen Alkoholmolekül verbinden kann. Eine weitere Erhöhung der Alkoholkonzentration resultiert dann nicht in einer höheren Abbaurrate, sondern die maximale Abbaurrate, die wir $A_{\max}$ nennen, ist erreicht. Sie ist durch die Anzahl der in der Leber vorhandenen Enzyme und durch die durchschnittliche Lebensdauer des Komplexes bestimmt.

Man kann die Situation mit einem Supermarkt vergleichen: Bei geringer Besucherzahl ist die Anzahl der Besucher, die pro Zeiteinheit die Kassen passieren, proportional zur Anzahl der Besucher. Bei einer großen Zahl von Besuchern bilden sich Schlangen an den Kassen und die Zahl der Besucher, die pro Zeiteinheit die Kassen passieren, ist bestimmt durch die Anzahl der Kassen und die durchschnittliche Zeit, die ein Besucher für den Durchgang durch die Kasse benötigt.

Wir fügen jetzt diese beiden gewünschten Eigenschaften so einfach wie möglich zusammen, nämlich durch eine abschnittsweise lineare stetige Funktion:

$$A(F) = \begin{cases} \frac{A_{\max}}{F_{\text{satt}}} \cdot F & \text{falls } F \leq F_{\text{satt}} \\ A_{\max} & \text{falls } F > F_{\text{satt}} \end{cases} \quad (5.33)$$

Hier gibt F_{satt} die Konzentration an, ab der mit maximaler Rate abgebaut wird. Über den Sättigungswert F_{satt} von F ist sich die einschlägige Literatur nicht einig: Wir finden Werte zwischen 0,2 ‰ und 0,1 ‰, vgl. [83] bzw. [47]. Für die maximale Abbaurate findet man als Richtwert 0,1 ‰/h, vgl. z. B. [54]. Die enzymatische Reaktion mit Alkoholhydrogenase liefert also einen Beitrag $-A(F)$ zur Änderungsrate $\dot{F}$ der Konzentration.

Weitere Prozesse des Erzeugens oder Vernichtens von Alkoholmolekülen innerhalb der Subsysteme scheinen keine große Rolle zu spielen.

3. Der Übergang von Alkoholmolekülen von einem Subsystem in das andere durch die gemeinsamen Trennflächen findet häufig durch sogenannte *Diffusion* statt: Moleküle stoßen im Rahmen ihrer thermischen Bewegungen zufällig von der einen oder anderen Seite gegen die Trennfläche und treten mit einer gewissen Wahrscheinlichkeit durch die Trennfläche in das andere Subsystem über. Es ist plausibel anzunehmen, dass die Anzahl der Stöße auf der jeweiligen Seite der Trennwand – und damit die Anzahl der Übertritte in das benachbarte Subsystem – proportional zur Konzentration auf dieser Seite ist. Subsystem 1 verliert also mit einer Rate proportional zur Konzentration in Subsystem 1 Moleküle an Subsystem 2 und gewinnt Moleküle aus Subsystem 2 mit einer Rate, die proportional zur Konzentration in Subsystem 2 ist.

In unserem Kontext ist die Situation einfach: Weil das System Körperflüssigkeit sehr viel mehr Masse umfasst als das System Magen-Darm-Trakt, ist die Alkoholkonzentration in der Körperflüssigkeit verschwindend gering im Vergleich zur Alkoholkonzentration im Magen-Darm-Trakt. Übertritte von Alkoholmolekülen aus der Körperflüssigkeit in den Magen-Darm-Trakt sind also so selten, dass wir sie vernachlässigen können. Auf der anderen Seite ist die Rate der Übertritte vom Magen-Darm-Trakt in die Körperflüssigkeit proportional zur Konzentration im Magen-Darm-Trakt und damit, weil während des Beobachtungszeitraums dem Magen-Darm-Trakt keine Stoffe zu- oder abgeführt werden, zur Masse des Alkohols $D(t)$. Wir erhalten also einen Beitrag $-k_1 D(t)$ zur Änderungsrate $\dot{D}$ der Alkoholmasse im Magen-Darm-Trakt und einen Beitrag $+k_1/M_F D(t)$ zur Änderungsrate $\dot{F}$ der Konzentration in der Körperflüssigkeit. Dabei ist M_F die Masse der Körperflüssigkeit.

Fassen wir alles zusammen, so erhalten wir das folgende *Differentialgleichungssystem*

$$\begin{aligned}\dot{D} &= -k_1 D \\ \dot{F} &= +\frac{k_1}{M_F} D - A(F),\end{aligned}\tag{5.34}$$

das wir mit den folgenden Anfangswerten versehen:

$$D(0) = D_0 \quad \text{und} \quad F(0) = 0.\tag{5.35}$$

Diese spiegeln wider, dass der Proband zu Beginn der Beobachtung nüchtern ist.

5.8.3 Die Berechnung der Modellfunktion

Glücklicherweise ist die erste Gleichung von (5.34) von der zweiten Gleichung unabhängig und wir können lösen: $D(t) = D_0 e^{-k_1 t}$. Diese Lösung können wir nun in die zweite Gleichung einsetzen. Für die zweite Gleichung müssen wir drei Phasen unterscheiden: Für die erste Phase müssen wir mit $k_2 := \frac{A_{\max}}{F_{\text{satt}}}$ die Differentialgleichung

$$\dot{F} = \frac{k_1}{M_F} D(t) - k_2 F \quad (5.36)$$

zum Anfangswert $F(0) = 0$ betrachten. Dieses Anfangswertproblem hat die Lösung

$$F_1(t) = \frac{k_1 D_0}{M_F(k_2 - k_1)} \left(e^{-k_1 t} - e^{-k_2 t} \right) \quad (5.37)$$

(siehe Abschn. 5.6). Diese Modellfunktion ist gültig, bis die Konzentration F_{satt} erstmalig erreicht wird, also bis zum ersten Zeitpunkt $t_1 > 0$ mit

$$F_1(t_1) = F_{\text{satt}}. \quad (5.38)$$

Die Lösung dieser Gleichung ist nicht elementar angebbbar; wir werden den Zeitpunkt später, wenn wir unsere Parameterwerte spezifiziert haben, numerisch bestimmen.

In der zweiten Phase müssen wir die Differentialgleichung

$$\dot{F} = \frac{k_1}{M_F} D_0 e^{-k_1 t} - A_{\max} \quad (5.39)$$

mit dem Anfangswert $F(t_1) = F_{\text{satt}}$ betrachten. Eine einfache Integration ergibt

$$\begin{aligned} F_2(t) &= F(t_1) + \int_{t_1}^t \dot{F}(s) ds = F_{\text{satt}} + \int_{t_1}^t \frac{k_1}{M_F} D_0 e^{-k_1 s} - A_{\max} ds \\ &= F_{\text{satt}} + \frac{D_0}{M_F} e^{-k_1 t_1} \left(1 - e^{-k_1(t-t_1)} \right) - A_{\max}(t - t_1). \end{aligned} \quad (5.40)$$

Diese Lösung ist gültig, solange $F_2 > F_{\text{satt}}$ gilt. Mit t_2 bezeichnen wir den ersten Zeitpunkt $t_2 > t_1$ mit $F_2(t_2) = F_{\text{satt}}$. Auch diesen Zeitpunkt werden wir später für konkret gegebene Parameterwerte numerisch approximieren.

In der dritten Phase, wenn die Alkoholkonzentration wieder unter F_{satt} gesunken ist, müssen wir wieder die Differentialgleichung (5.36), aber diesmal mit dem Anfangswert $F_3(t_2) = F_{\text{satt}}$, lösen. Wir erhalten

$$F_3(t) = \left(F_{\text{satt}} - \frac{k_1}{M_F(k_2 - k_1)} D(t_2) \right) e^{-k_2(t-t_2)} + \frac{k_1}{M_F(k_2 - k_1)} D(t). \quad (5.41)$$

Insgesamt erhalten wir also die Modellfunktion

$$F(t) = \begin{cases} F_1(t) & \text{für } t \leq t_1 \\ F_2(t) & \text{für } t_1 < t \leq t_2 \\ F_3(t) & \text{für } t_2 < t, \end{cases} \quad (5.42)$$

wobei F_1 , F_2 und F_3 durch die Funktionsterme in (5.37), (5.40) bzw. (5.41) gegeben sind.

5.8.4 Vergleich der beiden Modellfunktionen

Wir wollen nun unsere Modellfunktion mit der Modellfunktion aus der Abituraufgabe vergleichen. Dabei bemerken wir, dass diese konkrete Parameterwerte benutzt und zudem dimensionalisiert vorliegt: Es wird die Zeitskala $\bar{t} = 1 \text{ min}$ benutzt und die Konzentration wird in Promille angegeben. Dimensionalisieren wir die angegebene Funktion, haben wir es für $a = 20$ mit der Funktion

$$G(t) = \begin{cases} \frac{\%_0}{3} \left(1 - e^{-\frac{t}{20 \text{ min}}}\right) - \frac{t \%_0}{600 \text{ min}} & \text{für } 0 \leq t \leq 140 \text{ min} \\ u \%_0 \cdot e^{-\frac{v \cdot t}{\text{min}}} & \text{für } 140 \text{ min} < t \end{cases} \quad (5.43)$$

zu tun, wobei u und v bestimmt werden müssen. Als Kontrolllösungen sind $u \approx 1,01357$ und $v \approx 0,01657$ angegeben. Der erste Teil dieser Lösung hat dieselbe Bauart wie der zweite Teil unserer Lösung. Anscheinend wird der Modellierung aus der Abituraufgabe zufolge sofort mit der maximalen Rate Alkohol abgebaut.

Zum Vergleich lösen wir ebenfalls die Gleichung (5.39) zum Anfangswert $F(0) = 0$, erhalten die Lösung

$$\tilde{F}(t) = \frac{D_0}{M_F} \left(1 - e^{-k_1 t}\right) - A_{\max} t \quad (5.44)$$

und erkennen darin tatsächlich den ersten Teil der Modellfunktion aus der Abituraufgabe. Offenbar werden die maximale Abbaurate $A_{\max} = \frac{1 \%_0}{600 \text{ min}} = 0.1 \%_0/\text{h}$ sowie die Diffusionskonstante $k_1 = 1/20 \frac{1}{\text{min}}$ benutzt. Gehen wir von einem Wasseranteil von 60 %, einer Gesamtkörpermasse von 100 kg und einer Anfangsration $D_0 = a \text{ g}$ aus, erhalten wir tatsächlich die Konstante $\frac{D_0}{M_F} = \frac{a}{60} \%_0$ und $1/3$ für $a = 20$.

Der Wechsel zur nächsten Phase vollzieht sich in der Modellfunktion der Abituraufgabe zur Zeit $\tilde{t}_2 = 140 \text{ min}$. Es gilt $G(\tilde{t}_2) \approx 0.1 \%_0$. Die Aufgabensteller haben offensichtlich den auch in der Literatur diskutierten Wert $F_{\text{satt}} = 0.1 \%_0$ gewählt. Die Fortsetzung der Modellfunktion aus der Abituraufgabe löst jetzt allerdings nur die Differentialgleichung

$$\dot{F} = -k_3 F \quad (5.45)$$

anstelle der Differentialgleichung

$$\dot{F} = \frac{k_1}{M_F} D - k_2 F. \quad (5.46)$$

Hierbei ist k_3 so bestimmt, dass die zusammengesetzte Funktion differenzierbar ist. Nachrechnen liefert

$$k_3 = k_2 - \frac{k_1 D(\tilde{t}_2)}{M_F}.$$

Es wird also die Lösungsfunktion

$$F(t) = F_{\text{satt}} e^{-k_3(t-t_2)} \quad (5.47)$$

von (5.45) zusammen mit dem Anfangswert $G(\tilde{t}_2) = F_{\text{satt}}$ betrachtet, anstelle der Lösungsfunktion

$$F(t) = \left(F_{\text{satt}} - \frac{k_1}{M_F(k_2 - k_1)} D(t_2) \right) e^{-k_2(t-t_2)} + \frac{k_1}{M_F(k_2 - k_1)} D(t) \quad (5.48)$$

von (5.46) zum selben Anfangswert. Da aber zu diesem Zeitpunkt bei den benutzten Parameterwerten bereits der überwiegende Anteil der Alkoholmenge im Körperwasser resorbiert ist ($D(\tilde{t}_2) = 20e^{-7} \text{ g} \approx 0,2 \text{ g}$), spielt diese Vereinfachung keine allzu große Rolle. Wir wollen nun vergleichen, wie stark sich die beiden Modellfunktionen bei den gegebenen Parameterwerten unterscheiden. Dabei nutzen wir die Parameterwerte aus der Abituraufgabe, also $D_0 = 20 \text{ g}$, $M_F = 60 \text{ kg}$, $k_1 = 1/20 \text{ min}^{-1}$, $A_{\text{max}} = 0,1 \text{ ‰/h}$, $F_{\text{satt}} = 0,1 \text{ ‰}$, also $k_2 = 1/60 \text{ min}^{-1}$. Mit dem Newton-Verfahren approximieren wir die Zeitpunkte t_1 und t_2 . Wir müssen also zunächst die kleinste positive Lösung der Gleichung

$$f_1(\tau) := -\frac{1}{2} \left(e^{-\frac{1}{20}\tau} - e^{-\frac{1}{60}\tau} \right) = 0, 1 \quad (5.49)$$

bestimmen. Dazu implementieren wir das Newton-Iterationsschema $\tau_{n+1} = \tau_n - \frac{f_1(\tau_n) - 0,1}{f_1'(\tau_n)}$ in einem Tabellenkalkulationsprogramm. Als Anfangswert können wir z. B. $\tau_0 = 0$ wählen. Nach wenigen Iterationsschritten erhalten wir $\tau_1^* = 7,7461$ und damit $t_1 = 7,7461 \text{ min}$. Mit diesem Wert für t_1 bestimmen wir nun t_2 mit Hilfe der zweiten positiven Lösung der Gleichung

$$f_2(\tau) := \frac{1}{3} e^{-\frac{\tau_1^*}{20}} \left(1 - e^{-\frac{\tau - \tau_1^*}{20}} \right) - 0, 1(\tau - \tau_1^*) = 0, 1 \quad (5.50)$$

wieder mit Hilfe des Newton-Verfahrens. Wir müssen nun darauf achten, einen hinreichend großen Startwert zu wählen (ansonsten läuft das Verfahren wieder in die bereits bekannte Nullstelle τ_1^*), und erhalten z. B. mit $\tau_0 = 80$, wieder nach wenigen Iterationsschritten, den Wert $\tau_2^* = 143,369$, also $t_2 = 143,369 \text{ min}$.

In Abb. 5.21 sind beide Modellfunktionen mitsamt den absoluten und relativen Fehlern geplottet. Absolut beträgt der Fehler stets weniger als $0,6 \text{ ‰}$, der relative Fehler ist stets kleiner als 10 ‰ und vor allem dann groß, wenn die Konzentrationen verschwindend klein sind. Praktisch gesehen liefern also beide Modellfunktionen dieselben Vorhersagen.

Für andere Parameterwerte kann es aber durchaus problematisch sein, wenn die erste Phase, in der Alkohol noch nicht mit der maximalen Rate abgebaut wird, ignoriert wird.

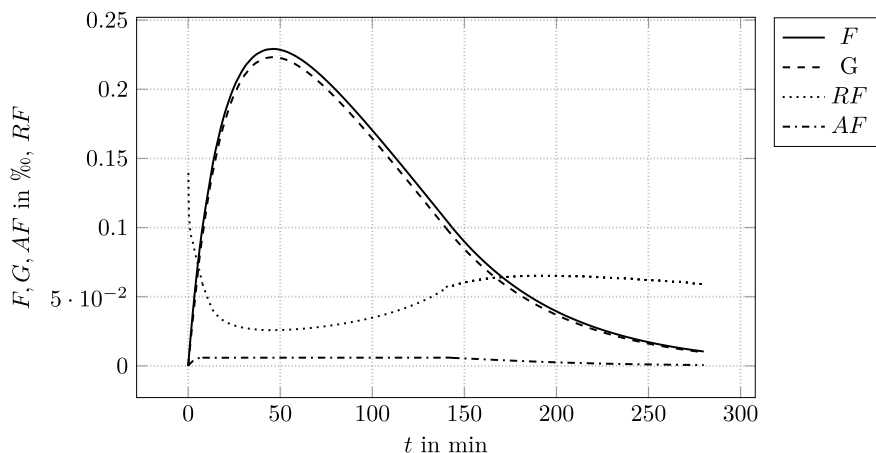


Abb. 5.21 Vergleich der Modellfunktion: F , G Modellfunktionen, $AF = F - G$ absoluter Fehler, $RF = (F - G)/F$ relativer Fehler

In Aufg. 5.21 soll numerisch mit Hilfe eines zu erstellenden GeoGebra-Blatts untersucht werden, wie sich die beiden Modelle bei unterschiedlichen Parameterwerten unterscheiden.

Zusammenfassend bleibt festzuhalten, dass die Modellfunktion der Abituraufgabe beileibe nicht „an den Haaren herbeigezogen ist“, sondern eine sinnvolle Modellierung des Kontextes liefert. Fraglich bleibt, ob für das Lösen der Abituraufgabe tatsächlich Modellierungskompetenzen erforderlich sind, weil die wesentlichen Modellierungsüberlegungen bereits durch die Aufgabenstellerin erbracht wurde und in der Aufgabe vollkommen unsichtbar bleiben.

5.8.5 Reflexionen zum mathematischen Modellieren IV: deskriptive und explikative Modellierungen

Es gibt unterschiedliche Möglichkeiten, ein mathematisches Modell für ein bestimmtes Phänomen oder eine bestimmte Problemstellung auszuwählen. Denken wir an die Modellierung des zeitlichen Verlaufs der Alkoholkonzentration, so könnte man z. B. einen oder mehrere Verläufe empirisch messen. Man würde erkennen, dass auf einen relativ schnellen Anstieg bis zu einem Hochpunkt ein langsamerer Abfall mit einem Wendepunkt erfolgt und die Kurve sich asymptotisch der Null nähert. Auf dieser Grundlage könnte man eine hinreichend große parametrisierte Familie von Funktionen auswählen, die alle diese Eigenschaften aufweisen, und anschließend, eventuell unter Zuhilfenahme statistischer Methoden, die Parameter so auswählen, dass die Modellfunktion möglichst genau zur empirisch gemessenen Funktion passt. Ein solches Vorgehen nennt man eine *deskriptive Modellierung* und hat folgende Nachteile:

- Zur Parameteranpassung sind empirische Daten nötig. Das heißt, der Vorgang muss schon stattgefunden haben. Es ist schwierig, zukünftige Vorgänge zu prognostizieren, weil dazu noch kein empirisches Material existiert.
- Die Auswahl der Funktionenfamilie ist in einem gewissen Rahmen willkürlich.
- Ändert sich die Realsituation in irgendeiner Beziehung, ist es ohne weitere empirische Daten vollkommen unklar, ob die Familie der Modellfunktionen noch brauchbar ist.

Überspitzt könnte man sagen, die mathematische Modellierung liefert lediglich eine bequeme Darstellung der empirischen Daten, sie erklärt aber wenig und erlaubt auch keine Vorhersagen.

Wir haben im Gegensatz dazu eine *explikative Modellierung* versucht: Der Ausgangspunkt besteht nicht aus empirischen Daten, sondern aus Spekulationen über die Natur der ablaufenden physiologischen Vorgänge, denen einfache Gesetzmäßigkeiten unterstellt werden. Mathematisch mündet das in Differentialgleichungen. Wir können nun a posteriori die Lösungen der Differentialgleichung mit empirischen Daten vergleichen. Wenn die Übereinstimmung gut ist, können wir hoffen, dass wesentliche Eigenschaften des Prozesses durch unsere Modellierung gut erfasst worden sind, wir die Prozesse also „verstehen“. Damit sind wir im Besitz einer Theorie der Vorgänge und können Prognosen zu noch nicht abgelaufenen Vorgängen, z. B. mit anderen Startwerten oder anderen Ratenkonstanten, abgeben. Wenn die Übereinstimmung schlecht ist, deutet dies darauf hin, dass unser Verständnis der ablaufenden Prozesse nicht ausreichend ist.

Aufgaben

5.1 Ein weiterer Weg zur logistischen Differentialgleichung: die Verbreitung von Gerüchten oder Krankheiten

Wir stellen uns die folgende Situation vor: In einer Stadt mit P Einwohnern verbreitet sich ein Gerücht (eine Krankheit). Die Anzahl der zur Zeit t über das Gerücht informierten Personen bezeichnen wir mit $I(t)$, die Anzahl der nichtinformierten mit $N(t)$. Klarerweise gilt dann zu jeder Zeit $P = I(t) + N(t)$, also $N(t) = P - I(t)$. Beim Aufeinandertreffen einer informierten und einer nichtinformierten Person wird letztere über das Gerücht in Kenntnis gesetzt, gehört also danach zur Gruppe der Informierten. Für die zeitliche Entwicklung der Informierten wird die folgende Differentialgleichung vorgeschlagen:

$$\dot{I} = kI(P - I) \quad (5.51)$$

Erläutern Sie, welche Überlegungen zu dieser Form der logistischen Differentialgleichung führen, und beschreiben Sie die Bedeutung des Parameters k . Für eine weitere Diskussion über die unterschiedlichen Formen der logistischen Differentialgleichung siehe auch [5].

5.2 Die Erstellung eines Bifurkationsdiagramms

In einem Populationsmodell wird die entdimensionalisierte Populationsgröße x durch die DGL

$$\dot{x} = x (f(x) - \hat{a}) \quad (5.52)$$

modelliert. Die Graphen von f und $\hat{a}$ sind in Abb. 5.22 skizziert.

- a) Skizzieren Sie in einem t - x -Diagramm alle Lösungen von (5.52).

Wir betrachten nun die Schar von Systemen

$$\dot{x} = x (f(x) - a), \quad (5.53)$$

wobei f weiterhin durch die Zeichnung gegeben ist. Der Parameter a sei ein Bifurkationsparameter.

- b) Skizzieren Sie das Bifurkationsdiagramm für (5.53) und $0 \leq a$. Benutzen Sie dabei die Bezeichnungen aus der obigen Abbildung. Markieren Sie auch den Einzugsbereich der stabilen stationären Punkte.
- c) Beschreiben Sie, wie sich das durch diese Gleichungen modellierte System verhält, wenn der Parameter a sehr langsam von 0 bis auf $\bar{a}$ und wieder zurück zu 0 variiert wird.

5.3 Befischung eines Fischbestands mit großen Booten

Diskutieren Sie das Befischungsmodell mit Schwellen- und Sättigungseffekt (5.12) in Abschn. 5.3 unter der Voraussetzung

- $V < \frac{K}{2} < A$
- $V = 0$ und $\frac{K}{2} < A$.

5.4 Das Um- und Zurückkippen von Seen

Bei Binnenseen unterscheidet man die Zustände *oligotroph*, *mesotroph*, *eutroph* und *hypertroph*. Grob gesprochen sind oligotrophe Seen klar, sauerstoffreich und beherbergen viele unterschiedliche Tier- und Pflanzenarten. Hypertrophe Seen sind durch Algenpopulationen trüb, sauerstoffarm und beherbergen wenige Arten mit großen Individuenzahlen. Entscheidend für den Zustand eines Sees ist der Phosphatgehalt: Seen mit niedrigem Phosphatgehalt sind oligotroph und Seen mit hohem Phosphatgehalt hypertroph. Carpenter et al. [17] schlagen zur Modellierung der Phosphatgehalts p eines Sees die folgende DGL vor:

$$\dot{p} = L - sp + r \frac{p^q}{m^q + p^q} \quad (5.54)$$

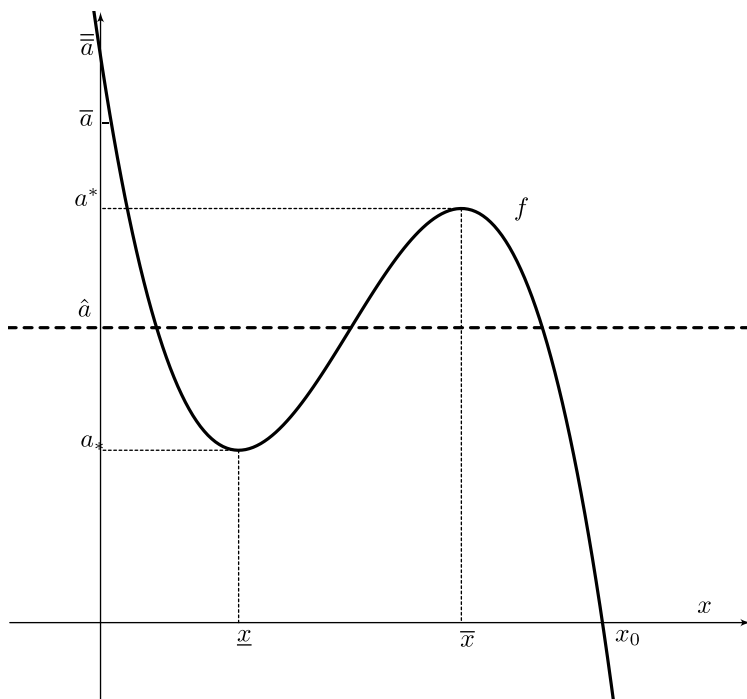


Abb. 5.22 Graph der Funktion f in (5.52)

Hierbei ist L die Zuflussrate an Phosphor, z. B. durch Abwässer oder durch Zuflüsse von Düngemitteln aus benachbarter Landwirtschaft. Der Term $-sp$ gibt den Verbrauch von Phosphat an, z. B. durch Sedimentation, Abfluss und Absorption durch Pflanzen. Der Term $+r \frac{p^q}{m^q + p^q}$ gibt an, mit welcher Rate im Seeboden gespeichertes Phosphat wieder ins offene Wasser gelangt, die sogenannte Recyclingrate. Dabei ist q ein Exponent mit Werten zwischen 2 für einen kalten und tiefen See und bis zu 20 für einen flachen und warmen See. r gibt die maximale Recyclingrate an und m den Phosphatgehalt bei halber maximaler Recyclingrate.

- a) Rechnen Sie nach, dass sich durch geeignete Skalenwahl (5.54) zu

$$\dot{x} = a - bx + \frac{x^q}{1 + x^q} \quad (5.55)$$

entdimensionalisieren lässt. Bestimmen Sie die benutzten Skalen und die dimensionslosen Parameter a, b in Abhängigkeit von L, s, r, m . Begründen Sie, warum sich der Parameter a am besten von außen kontrollieren lässt.

- b) Bestimmen Sie für fixiertes b und $1 < q < \infty$ die Anzahl und Stabilität der stationären Punkte von (5.55) in Abhängigkeit von a mit einem graphischen Verfahren und fertigen Sie ein Bifurkationsdiagramm an. Unterscheiden Sie dabei die Fälle, dass b kleiner ist

als das Maximum der Ableitung von g mit $g(x) := \frac{x^q}{1+x^q}$ und dass b größer als dieses Maximum ist.

Benutzen Sie DG-Software wie GeoGebra, um die unterschiedlichen Situationen zu veranschaulichen.

- c) Der Federsee, siehe Abb. 5.23, ist ein Endmoränensee in Oberschwaben, der ursprünglich mesotroph war. In den 50er Jahren des 20. Jahrhunderts kippte der See durch Intensivierung der benachbarten Landwirtschaft und Einleitung von Haushaltsabwässern in einen hypertrophen Zustand. 1984 wurde eine Kläranlage gebaut und die Landwirtschaft eingeschränkt. Durch diese Maßnahmen wurde die Einleiterrate des Phosphats unter die Rate zur Zeit des Umkippens gedrückt. Trotzdem verblieb der See für weitere 24 Jahre im hypertrophen Zustand und kippte erst 2008 in den mesotrophen Zustand zurück. Die Einleitung von Phosphat wurde in diesen Jahren weiter reduziert. Erklären Sie dieses Phänomen im Rahmen des aufgestellten Modells.

5.5 Ein biochemischer Schalter

Eine Aufgabe aus Murray [65]: Ein Genprodukt (also ein Protein oder RNA o.Ä.) mit der Konzentration g wird durch eine Chemikalie S mit zugehöriger Konzentration s gebildet, geht eine autokatalytische Reaktion mit sich selber ein und zerfällt spontan. Das führt auf die Differentialgleichung

$$\dot{g} = s + k_1 \frac{g^2}{1 + g^2} - k_2 g =: f(g; s),$$

wobei k_1 und k_2 gegebene positive Konstanten sind und s eine konstante Konzentration.

- a) Zeigen Sie, dass es für $s = 0$ zwei positive stationäre Punkte gibt, wenn $k_1 > 2k_2$, und bestimmen Sie ihre Stabilität.



Abb. 5.23 Der Federsee: ein Endmoränensee in Oberschwaben. (Foto: © K. Weiss)

- b) Untersuchen Sie die stationären Punkte und ihre Stabilität für $s > 0$ und $k_1 > 2k_2$.
 c) Beschreiben Sie, was passiert, wenn zu Anfang kein Genprodukt und keine Substanz S vorhanden ist und dann langsam die Konzentration s von S bis auf einen hinreichend großen Wert hochgefahren wird. Danach wird die Konzentration wieder heruntergefahren. Warum spricht man in diesem Zusammenhang von einem *biochemischen Schalter*?

5.6 Eine chemische Reaktion

Bei einer chemischen Reaktion verbinden sich zwei Stoffe A und B zu einem dritten Stoff C , der auch wieder in die beiden Ausgangsstoffe zerfallen kann:



Als Modell für die zeitliche Entwicklung der Stoffmengen u, v, w von A, B, C wird das System

$$\left. \begin{aligned} \dot{u} &= -\alpha uv + \beta w \\ \dot{v} &= -\alpha uv + \beta w \\ \dot{w} &= \alpha uv - \beta w \end{aligned} \right\} \quad (5.56)$$

vorgeschlagen, wobei $\alpha > 0$ und $\beta \geq 0$ seien.

- a) Interpretieren Sie jeden der in (5.56) auftretenden Terme.
 b) Beweisen Sie, dass

$$\varphi_1 := u - v \quad \text{und} \quad \varphi_2 := u + v + 2w$$

Erhaltungsgrößen des Modells (5.56) sind, und beschreiben Sie, was diese Erhaltungsgrößen im Sachkontext bedeuten.

Dabei heißt $\varphi : [0, \infty) \rightarrow \mathbb{R}$ eine Erhaltungsgröße, wenn $\varphi = \text{const.}$ gilt.

- c) Betrachten Sie (5.56) unter der Anfangsbedingung $u(0) = u_0, v(0) = v_0$ und $w(0) = w_0$ mit $u_0 > v_0 > w_0 = 0$. Zeigen Sie mit Hilfe von Aufgabenteil b), dass es positive Konstanten c_1 und c_2 gibt derart, dass

$$\dot{u} = c_1\beta + (c_2\alpha - \beta)u - \alpha u^2 \quad \text{für alle } t > 0 \quad (5.57)$$

gilt, und bestimmen Sie c_1 und c_2 .

- d) Bestimmen Sie im Grenzfall $\beta = 0$ die stationären Punkte von (4.12) und deren Stabilität. Beschreiben Sie damit das Langzeitverhalten der entsprechenden Version von (5.56). Bestimmen Sie für diesen Fall insbesondere die Entwicklung der Stoffmengen auf lange Sicht.

5.7 Einlagerung von Kalzium in eine Zelle

Es soll die Konzentration einer bestimmten Substanz (z. B. Kalzium) in einer Zelle modelliert werden. Die Konzentration der Substanz *innerhalb* der Zelle zur Zeit t bezeichnen wir mit $u(t)$. Ferner nehmen wir an, dass die Konzentration *außerhalb* der Zelle konstant gleich u_A

ist. Ohne Einlagerung der Substanz modellieren wir die Konzentration in der Zelle durch

$$\dot{u} = -r(u - u_A) \quad (5.58)$$

miteiner konstanten Diffusionsrate $r > 0$.

- a) Erläutern Sie, welche Modellierungsannahmen hinter dem Model (5.58) stehen. Bestimmen Sie den stationären Punkt dieser Gleichung und geben Sie seine Stabilität an.

Nun soll die Substanz zusätzlich in der Zelle eingelagert werden, z. B. um ein festes Gerüst zu bilden. Die Gl. (5.58) wird um einen Absorptionsterm erweitert:

$$\dot{u} = -r(u - u_A) - s \frac{u}{u_M + u} \quad (5.59)$$

Dabei ist $s > 0$ konstant und $u_M > 0$ die Konzentration, bei der die Absorptionsrate ihren halben maximalen Wert erreicht.

- b) Erläutern Sie, welche Modellierungsannahmen hinter dem Absorptionsterm stehen.
 c) Begründen Sie graphisch, dass das Modell immer genau einen positiven stationären Punkt hat und bestimmen Sie graphisch die Stabilität dieses Punkts.
 d) Bestimmen Sie auch einen rechnerischen Ausdruck für diesen positiven stationären Punkt und freuen Sie sich, dass Sie die Stabilität nicht rechnerisch über das Vorzeichen der Ableitung nachprüfen müssen.

5.8 Bifurkationen bei den Tannenwicklern

Erstellen Sie Bifurkationsdiagramme für das Tannenwickler-System (5.16) für

- a) $b = 10$, Bifurkationsparameter a ,
 b) $a = 0,25$, Bifurkationsparameter b ,
 c) $a = 0,55$, Bifurkationsparameter b ,
 d) $a = \frac{b}{40}$, Bifurkationsparameter a ,
 – zunächst qualitativ, indem Sie überlegen, wie die stationären Punkte auf Änderungen des Bifurkationsparameters reagieren,
 – dann quantitativ, indem Sie aus Gl. (5.17) den funktionalen Zusammenhang zwischen den positiven stationären Punkten und dem Bifurkationsparameter bestimmen. Plotten Sie das Bifurkationsdiagramm mit einem geeigneten Programm.

5.9 Die Spitze der kritischen Kurve

Berechnen Sie die Koordinaten der Spitze der kritischen Kurve der Gl. (5.16), siehe Abb. 5.14.

5.10 Bekämpfungsstrategien gegen den Tannenwickler

Stellen Sie auf Grundlage des Modells (5.15) Bekämpfungsstrategien gegen die Schädlingpopulation P auf, indem Sie untersuchen, wie Veränderungen der Parameterwerte r_P , K_P , V und S sich auf die Schädlingpopulation auswirken. Stellen Sie dabei dar, welche Maßnahmen in der Realität mit einer Änderung der Parameterwerte im Modell korrespondieren könnten.

5.11 Krebse und Bakterien

In einem Gefäß werden Krebse gezüchtet, die sich von Bakterien ernähren. Die Größe $x(t)$ der Bakterienpopulation zur Zeit t wird durch folgende Differentialgleichung modelliert:

$$\dot{x} = a - bx^2 \quad \text{mit } a, b > 0. \quad (5.60)$$

- Beschreiben Sie kurz, welche Modellierungsannahmen auf diese DGL führen können.
- Bestimmen Sie die Dimensionen der Parameter a und b . Zeigen Sie, dass sich die Gl. (5.60) in die entdimensionalisierte Form

$$\dot{u} = 1 - u^2 \quad (5.61)$$

überführen lässt, und bestimmen Sie die dabei verwendeten Skalen.

- Beschreiben Sie, wie sich die Bakterienpopulation auf lange Sicht verhält, indem Sie die stationären Punkte von (5.60) oder (5.61) und ihr Stabilitätsverhalten untersuchen.
- Berechnen Sie die explizite Lösung von (5.60) zu einem Anfangsdatum $x(0) = x_0$ mit $0 \leq x_0 < \sqrt{\frac{a}{b}}$ mit der Methode der Trennung der Variablen.

5.12 Der Fallschirmspringer

Ein Fallschirmspringer springt aus großer Höhe aus einem Flugzeug. Nach einem einfachen Modell gilt für die Geschwindigkeit v des Fallschirmspringers

$$m\dot{v} = mg - \beta v^2. \quad (5.62)$$

Dabei ist m die Masse des Fallschirmspringers und β der Luftwiderstand des Springers mitsamt Schirm, der im Wesentlichen linear von der Öffnungsfläche des Fallschirms abhängt.

- Bestimmen Sie, ohne eine explizite Lösung von (5.62) zu berechnen, mit welcher Geschwindigkeit in Abhängigkeit von m und β der Springer auf dem Boden aufkommt.
- Wie muss die Öffnungsfläche des Fallschirms verändert werden, wenn die Aufprallgeschwindigkeit halbiert werden soll?
- Bestimmen Sie die explizite Lösung von (5.62) zu einem Anfangswert $v_0 \geq 0$ mit der Methode der Trennung der Variablen.

5.13 Stetig differenzierbare Funktionen sind lokal Lipschitz-stetig

Zeigen Sie: Ist eine Funktion $f : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ stetig differenzierbar, dann ist sie lokal Lipschitz-stetig.

5.14 Explizites zur logistischen Differentialgleichung

- a) Lösen Sie das Anfangswertproblem zu $P(t_0) = P_0$ zur logistischen Differentialgleichung $\dot{P} = rP(1 - \frac{P}{K})$
- mit der Methode der Trennung der Variablen,
 - als Bernoulli'sche Differentialgleichung,
 - durch die Substitution $J = \frac{K-P}{P}$, die auf eine lineare Differentialgleichung führt, siehe [12].
- b) Geben Sie die Zeit t als Funktion von r und q an, die eine Lösung mit $P_0 = K/2$ benötigt, um auf den Wert qK mit $\frac{1}{2} \leq q < 1$ zu wachsen.

Tipp: Sie können sich das Leben erleichtern, wenn Sie die Differentialgleichung durch eine geeignete Entdimensionalisierung auf die Form $\dot{x} = x(1 - x)$ bringen.

5.15 Trennung der Variablen im nichtautonomen Fall

Die Funktionen $f, g : \mathbb{R} \rightarrow \mathbb{R}$ seien stetig. Formulieren und beweisen Sie einen Existenz- und Eindeutigkeitssatz für das Anfangswertproblem

$$\dot{x}(t) = f(x(t)) \cdot g(t), \quad x(t_0) = x_0. \quad (5.63)$$

5.16 Ein paar explizite Lösungen

- a) Bestimmen Sie die Lösungen der folgenden Anfangswertprobleme mit der Methode der Trennung der Variablen und geben Sie jeweils das maximale Intervall an, auf dem die Lösung existiert.
- i) $\dot{x} = e^{2t+3x}$ mit $x(0) = x_0 \in \mathbb{R}$.
 - ii) $x \cdot \dot{x} = e^{3t} \cdot e^{3x^2}$ mit $x(t_0) = x_0 \neq 0$.
 - iii) $(1 + t^2)\dot{x} + t(1 + x^2) = 0$ mit $x(0) = x_0 \in \mathbb{R}$.
 - iv) $t^3 \cdot \dot{x} = 2x - 5$ mit $t_0 \neq 0$ und $x(t_0) = x_0 \in \mathbb{R}$.
- b) Bestimmen Sie die allgemeine Lösung der folgenden linearen inhomogenen Gleichungen.
- i) $\dot{x} = 7x + e^t$
 - ii) $\dot{x} = 7x + te^t$
 - iii) $\dot{x} = 7x + t^2e^t$
 - iv) $\dot{x} = 7x + e^{7t}$
 - v) $\dot{x} = 7x + \sin(t)$
 - vi) $\dot{x} = \sin(t)x + \sin(t)$
 - vii) $\dot{x} = tx + t^2$

5.17 Der Vergleichssatz für die Vergangenheit

Beweisen Sie den Vergleichssatz 5.2 „in die Vergangenheit“:

Sei $f : \mathbb{R}_t \times \mathbb{R}_x \rightarrow \mathbb{R}$ stetig und lokal Lipschitz-stetig bzgl. der zweiten Variablen. Auf einem Intervall $I = [t_1, t_0]$ gelten für zwei differenzierbare Funktionen $x = x(t)$ und $y = y(t)$ die Ungleichungen

$$\dot{x}(t) \leq f(t, x(t)) \quad \text{für alle } t \in [t_1, t_0] \quad (5.64)$$

$$\dot{y}(t) \geq f(t, y(t)) \quad \text{für alle } t \in [t_1, t_0], \quad (5.65)$$

Ferner gelte $x(t_0) \geq y(t_0)$. Dann gilt

$$y(t) \leq x(t) \quad \text{für alle } t \in [t_1, t_0].$$

5.18 Ein Vergleichssatz für stetige Funktionen

Der Vergleichssatz 5.2 arbeitet mit lokal Lipschitz-stetigen Funktionen. Es gibt einen weiteren Vergleichssatz, der nur die Stetigkeit benötigt, dafür aber eine strikte Ungleichung erfordert.

Satz 5.3

Es sei g eine stetige reellwertige Funktion auf $\mathbb{R}$, $I = [t_0, t_1]$ ein (Zeit-)Intervall. Für zwei stetig differenzierbare Funktionen y, z auf I gelte:

$$\dot{y}(t) < g(y(t)) \quad \text{und} \quad \dot{z}(t) > g(y(t)) \quad \text{für alle } t \in I \quad \text{sowie} \quad y(t_0) \leq z(t_0). \quad (5.66)$$

Dann gilt

$$y(t) \geq z(t) \quad \text{für alle } t \in I.$$

Beweisen Sie diesen Satz.

Hinweis: Betrachten Sie die Differenz $d(t) = y(t) - z(t)$ und beachten Sie das Vorzeichen von $\dot{d}$ in den Nullstellen von d .

5.19 Numerische Approximation

Wir betrachten die beiden Differentialgleichungen

$$\dot{x} = 1 - 10x$$

$$\dot{x} = x(1 - 10x).$$

Lösen Sie die beiden folgenden Differentialgleichungen zum Anfangswert $x(0) = 1$

- explizit,
- numerisch mit dem expliziten Euler-Verfahren,
- numerisch mit dem impliziten Euler-Verfahren.

Nutzen Sie dazu ein TK-Programm und lassen Sie sich die Lösungen plotten.

- d) Beobachten Sie, was passiert, wenn Sie die Schrittweite Δt in den beiden numerischen Verfahren vergrößern. Versuchen Sie, Ihre Beobachtungen zu erklären.

5.20 Die logistische Differentialgleichung mit variabler Kapazität

Eine Population lebe in einer Umgebung, die sich mit der Zeit immer weiter verbessert, die Kapazitätsgröße der Population wachse linear: $K = K(t) = K_0 + kt$. Die zugehörige DGL lautet also:

$$\dot{P} = rP \left(1 - \frac{P}{K(t)} \right) \quad (5.67)$$

Eine plausible Vermutung ist, dass sich die Populationsgröße langfristig wie die Kapazität verhält, also $P(t) \approx K(t)$. Das soll jetzt untersucht werden.

- a) Zeigen Sie mit Hilfe des Vergleichsatzes 5.2: Falls $P_0 \leq K_0$, dann gilt $P_0 \leq a(t) \leq K(t)$ für alle $t \geq 0$.
b) Lösen Sie so weit wie möglich das Anfangswertproblem

$$\dot{P} = rP \left(1 - \frac{P}{K(t)} \right), \quad P(0) = P_0$$

explizit als Bernoulli-Gleichung (siehe Abschn. 5.6).

- c) Suchen Sie eine schärfere Unterlösung für P als P_0 .
d) Approximieren Sie die Lösungen des Anfangswertproblems

$$\dot{P} = rP \left(1 - \frac{P}{K_0 + kt} \right), \quad 0 < P(0) = P_0 < K_0$$

mit dem expliziten Euler-Verfahren und einem Tabellenkalkulationsprogramm. Vergleichen Sie Ihre numerische Lösung mit den gefundenen Ober- und Unterlösungen.

5.21 Die Modellierung der Alkoholkonzentration: Verallgemeinerungen

Wir verallgemeinern die Modellfunktion aus Abschn. 5.8 für allgemeine Parameter $D_0, k_1, M_F, A_{\max}, F_{\text{satt}}, k_2 = \frac{A_{\max}}{F_{\text{satt}}}$: Zunächst definieren wir die Funktion G_1

$$G_1(t) = \frac{k_1 D_0}{M_F} (1 - e^{-k_1 t}) - A_{\max} t.$$

Dann definieren wir den Zeitpunkt t^* als die größere der beiden Lösungen der Gleichung $G_1(t) = F_{\text{satt}}$. Mit dieser Zeit definieren wir

$$G_2(t) = F_{\text{satt}} e^{-k_2(t-t^*)}.$$

Die Modellfunktion setzen wir dann folgendermaßen zusammen:

$$G(t) = \begin{cases} G_1(t) & \text{für } 0 \leq t \leq t^* \\ G_2(t) & \text{für } t^* < t \end{cases}.$$

- Erstellen Sie ein GeoGebra-Blatt, in dem Sie die entdimensionalisierten Parameterwerte $\delta = D_0/g$, $\mu = M_F/g$, $\kappa_1 = k_1/\text{min}$, $\phi = F_{\text{satt}}/\text{‰}$ und $\alpha = A_{\text{max}}/(\text{‰ min})$ per Schiebereregler variieren können und das die Graphen der Modellfunktionen $f = F/\text{‰}$, $g = G/\text{‰}$ sowie ihrer absoluten und relativen Fehler zu den gerade eingestellten Parameterwerten als Funktionen der dimensionslosen Zeit $\tau = t/\text{min}$ plottet. Nutzen Sie zur numerischen Approximation der drei benötigten Zeitpunkte t_1 , t_2 und t^* das in GeoGebra vorhandene Tabellenkalkulationsprogramm.
- Variieren Sie systematisch die Parameterwerte und untersuchen Sie, unter welchen Umständen die Fehler groß werden. Interpretieren Sie Ihre Beobachtungen vor dem Hintergrund der Modellierungsannahmen.

5.22 Medikamentenkonzentration im Zentralabitur

Aus dem Zentralabitur Mathematik NRW des Jahres 2008 stammt der folgende Kontext:

„Ein Pharmaunternehmen produziert ein Medikament in unterschiedlichen Wirkstoffdosierungen, das in Tablettenform verabreicht wird. Der zeitliche Verlauf der Wirkstoffkonzentration im Blut eines Patienten kann in den ersten 24 Stunden nach Einnahme einer Tablette näherungsweise durch die Funktionenschar

$$f_a(t) = ate^{-0,25t}, \quad t \in [0, 24], \quad a > 0$$

beschrieben werden. Dabei wird die Zeit t in Stunden seit der Einnahme und die Wirkstoffkonzentration $f_a(t)$ im Blut in Milligramm pro Liter (mg/l) gemessen; die Höhe der Wirkstoffdosierung wird durch den Parameter a berücksichtigt.“

Hier betrachten wir nicht die sich anschließenden Aufgaben, sondern die aufgestellte Modellfunktion unter Modellierungsgesichtspunkten.

- Berechnen Sie, welche Differentialgleichung durch f_a gelöst wird. (Diese Aufgabe ist natürlich nicht eindeutig lösbar, finden Sie also eine „interessante“ DGL.)
- Welche Modellierungsannahmen führen auf so eine Differentialgleichung? Sind die Modellierungsannahmen vernünftig? Welche Änderungen der Modellierungsannahmen wären plausibel? Wie würden diese Änderungen die DGL und die Lösungen der DGL ändern?

Der letzte Aufgabenteil der Abituraufgabe lautet folgendermaßen:

„Untersuchen Sie das Verhalten der Funktion f_{10} für $t \rightarrow \infty$. Interpretieren Sie das Ergebnis im Hinblick auf den langfristigen Abbau des Wirkstoffs. Für $t > 24$ soll der zeitliche Verlauf

der Wirkstoffkonzentration durch eine lineare Funktion g beschrieben werden. Bestimmen Sie eine Gleichung der linearen Funktion g so, dass die zusammengesetzte Funktion h mit

$$h(t) = \begin{cases} f_{10}(t), & 0 \leq t \leq 24 \\ g(t), & t > 24 \end{cases}$$

an der Stelle $t = 24$ differenzierbar ist. Berechnen Sie für diese Modellierung den Zeitpunkt, zu dem das Medikament im Blut vollständig abgebaut ist.“

- c) Bewerten Sie diese Aufgabe unter Modellierungsaspekten.



Räuber-Beute-Modelle – Modellieren mit Differentialgleichungssystemen

6

Inhaltsverzeichnis

6.1	Das Räuber-Beute-System nach Lotka-Volterra	167
6.2	Lösungstheorie: der Existenz- und Eindeigkeitssatz von Picard-Lindelöf	169
6.3	Das Lotka-Volterra-System fortgesetzt – die Methode der Lyapunov-Funktion	171
6.4	Räuber-Beute-Systeme mit logistischem Wachstum – lineare Differentialgleichungssysteme und der Linearisierungssatz von Grobman-Hartman	184
6.5	Räuber-Beute-Systeme mit Sättigungs- und Schwelleneffekten – stabile Grenzyklen und der Satz von Poincaré-Bendixson	204
	Aufgaben	216

6.1 Das Räuber-Beute-System nach Lotka-Volterra

Während wir in unseren bisherigen Modellen die Entwicklung einer einzelnen Spezies in einem bestimmten Gebiet modelliert haben, bei der die wechselnden äußeren Einflüsse z. B. in Form von Bejagung durch von außen festgeschriebene Parameter eingebunden wurden, sollen nun die äußeren Effekte selbst eine eigenständige gesetzmäßige Entwicklung durchmachen. Wir starten mit einem klassischen Modell nach Lotka und Volterra, das unabhängig von Alfred J. Lotka 1925 [59] und Vito Volterra 1926 [89] aufgestellt wurde. Dabei untersuchte Lotka die Interaktion zwischen einer Wirtspopulation, gedacht war an Insekten, und einer Parasitenpopulation, die sich durch Eiablage in einem Wirtstier fortpflanzt. Volterra stellte das Modell ganz allgemein für zwei in einem Räuber-Beute-Verhältnis stehende Arten auf. Wir gehen ähnlich vor und betrachten die Populationen von zwei Arten, die einen gemeinsamen Lebensraum bewohnen. Dabei soll die eine Art, die wir als *Beute* bezeichnen, die Nahrungsgrundlage für die zweite Art sein, die wir althergebracht als *Räuber* bezeichnen, obwohl eine Benennung als Jäger vielleicht angebrachter wäre. Die Größe der Beutepopulation zur Zeit t sei $x(t)$ und die Größe der Räuberpopulation $y(t)$. Wir stellen nun für beide Populationen Differentialgleichungen auf.

Wir starten mit der Beutepopulation. Für die Entwicklung einer isolierten Population haben wir bereits in Kap. 5 die Modelle des exponentiellen und des logistischen Wachstums vorgestellt, auf die wir zurückgreifen werden. Wir müssen also vor allem den Interaktionsterm aufstellen, der modelliert, wie das Aufeinandertreffen von Räuber und Beute die Änderungsrate der Beutepopulation verändert. Die folgenden Annahmen sind einfach und einigermaßen plausibel (wir werden sie in späteren Abschnitten noch verfeinern):

- Der Beitrag des Interaktionsterms zur Änderungsrate der Beutepopulation hat ein negatives Vorzeichen.
- Er ist proportional zur Rate des Aufeinandertreffens von Räuber- und Beutetieren (diese Annahme werden wir später modifizieren).

Wir haben uns bereits bei der Diskussion des quadratischen Terms in der logistischen Differentialgleichung in seiner Interpretation als Konkurrenzterm überlegt, wie man die Anzahl des Aufeinandertreffens modellieren kann. Für ein einzelnes Beutetier ist es naheliegend, die Anzahl der Aufeinandertreffen mit Raubtieren als proportional zur Populationsgröße der Räuber anzusetzen. Weiterhin nehmen wir an, dass die Wahrscheinlichkeit eines Zusammentreffens für jedes Beutetier gleich groß ist und nicht von der Größe der Beutepopulation abhängt. Ferner nehmen wir an, dass die Konkurrenz innerhalb der Beutepopulation im Vergleich zur Interaktion mit der Räuberpopulation unbedeutend ist (exponentielles statt logistisches Wachstum), und erhalten die Differentialgleichung

$$\dot{x} = \alpha x - \beta xy. \quad (6.1)$$

Dabei gibt α wieder die Pro-Kopf-Reproduktionsrate (die Differenz der Pro-Kopf-Geburtenrate und der Pro-Kopf-Sterberate) an. Diese hat die Dimension $1/T$ und wir nehmen an, dass sie positiv ist. Der Parameter β hat die Dimension $\frac{1}{TA_R}$, wobei A_R die Dimension der Räuberpopulation angibt, also in der Regel die Anzahl der Räuber. Er ist also so etwas wie die „Pro-Kopf-Sterberate pro Räuber“ und gibt die Stärke der Räuber-Beute-Interaktion auf Seiten der Beutepopulation an.

Auf Seiten der Räuberpopulation modellieren wir entsprechend, nur dass diesmal die Interaktion positiv zu Buche schlägt und die Pro-Kopf-Reproduktionsrate ohne Interaktion mit der Beutepopulation negativ ist:

$$\dot{y} = -\gamma y + \delta xy. \quad (6.2)$$

Der Parameter γ ist also positiv und δ hat die Dimension $\frac{1}{TA_B}$ und gibt die Pro-Kopf-Reproduktionsrate pro Beutetier infolge der Interaktion an. Die Gl. (6.1) und (6.2) bilden zusammen das Räuber-Beute-System nach Lotka-Volterra.

Mathematisch haben wir es hier mit zwei gekoppelten Differentialgleichungen zu tun, was uns vor eine neue Situation stellt. Wir könnten zwar, wenn die Funktion y bekannt wäre, die Differentialgleichung (6.1) lösen, und wenn x bekannt wäre, (6.2) lösen. Da aber beide Funktionen unbekannt sind, brauchen wir weitere mathematische Theorien und Methoden, die in den kommenden Abschnitten bereitgestellt werden.

6.2 Lösungstheorie: der Existenz- und Eindeigkeitssatz von Picard-Lindelöf

Wir stellen zunächst einmal das Gleichungssystem aus (6.1) und (6.2) in einen größeren mathematischen Kontext: Wir wollen statt nur zwei allgemein n gekoppelte Differentialgleichungen betrachten und zudem eine explizite Zeitabhängigkeit der Differentialgleichungen zulassen. Vom Modellierungsgesichtspunkt aus bedeutet dies, dass sich die äußeren Umstände des Systems, das modelliert werden soll, auch ändern dürften. Die n zu modellierenden Größen bezeichnen wir mit x_1 bis x_n und fassen sie zu dem Vektor $x = (x_1, \dots, x_n)^t \in \mathbb{R}^n$ zusammen. Das Differentialgleichungssystem wird dann durch eine Abbildung

$$\begin{aligned} F : \mathbb{R}_t \times \mathbb{R}_x^n &\longrightarrow \mathbb{R}^n \\ (t, x) &\longmapsto F(t, x) = (F_1(t, x), \dots, F_n(t, x))^t \end{aligned} \quad (6.3)$$

gegeben. Dabei stellen wir uns unter t wieder die Zeit-Variable vor. Für den Fall, dass F nicht von t abhängt, sprechen wir von autonomen Differentialgleichungssystemen. Die Aufgabe besteht nun darin, Aussagen über die Lösbarkeit des Differentialgleichungssystems

$$\dot{x} = F(t, x) \quad (6.4)$$

zusammen mit dem Anfangswert

$$x(t_0) = x_0 \in \mathbb{R}^n \quad (6.5)$$

zu formulieren und zu beweisen. Dabei ist $t_0 \in \mathbb{R}$ der Anfangszeitpunkt und x_0 der Anfangswert der Größe x . Im Gegensatz zu Differentialgleichungen gibt es für Differentialgleichungssysteme keine elementare Lösungstheorie. Es gibt im Wesentlichen zwei fundamentale Sätze über Lösungen des Anfangswertproblems (6.4), (6.5): den Existenz- und Eindeigkeitssatz von Picard-Lindelöf, der die Existenz- und Eindeigkeit liefert, sowie den Existenzsatz von Peano, der unter geringeren Voraussetzungen an F die Existenz, aber nicht die Eindeigkeit von Lösungen garantiert. Wir werden hier die Aussage des Satzes von Picard-Lindelöf darstellen, auf einen Beweis aber verzichten, da dieser den elementaren und auf das Modellieren ausgerichteten Charakter dieses Buches sprengen würde. Zudem lässt sich der Beweis in so gut wie jedem Lehrbuch zu gewöhnlichen Differentialgleichun-

gen finden. Um die Voraussetzungen an die Funktion F zu formulieren, erinnern wir an die Begriffe *Lipschitz-stetig* und *lokal Lipschitz-stetig*:

Definition 6.1

Eine Funktion F wie in (6.3) heißt lokal Lipschitz-stetig in der Variable x , wenn es für jede kompakte Menge $K \subset \mathbb{R} \times \mathbb{R}^n$ eine Konstante L_K gibt, sodass für alle $(t, x), (t, \bar{x}) \in K$ die folgende Abschätzung erfüllt ist:

$$|F(t, x) - F(t, \bar{x})| \leq L_K |x - \bar{x}|. \quad (6.6)$$



Außerdem erinnern wir an die folgenden Sprechweisen aus Abschn. 5.5, die wir auch für Differentialgleichungssysteme übernehmen: Wir sagen, eine Lösung von (6.4) und (6.5) ist eindeutig, wenn sie auf jedem Intervall um t_0 , auf dem überhaupt eine Lösung existiert, eindeutig ist. Eine Lösung heißt maximal, wenn sie nicht fortgesetzt werden kann; eine Lösung heißt global, wenn sie für alle Zeiten $t \in \mathbb{R}$ existiert.

Satz 6.1 (Picard-Lindelöf)

Eine Funktion F wie in (6.3) sei stetig auf $\mathbb{R} \times \mathbb{R}^n$ und lokal Lipschitz-stetig bezüglich der Variablen x . Ferner sei ein Anfangswert $(t_0, x_0) \in \mathbb{R} \times \mathbb{R}^n$ gegeben. Dann hat das Anfangswertproblem (6.4), (6.5) eine eindeutige Lösung $t \mapsto x(t)$ mit einem maximalen offenen Existenzintervall $I = (t_-, t_+)$ mit $-\infty \leq t_- < t_0 < t_+ \leq \infty$. Die maximale Lösung läuft von Rand zu Rand, das heißt, es gilt entweder $t_{\pm} = \pm\infty$, oder die Lösung entwickelt einen Blow-up in endlicher Zeit: Es gilt $-\infty < t_-$ und $\lim_{t \rightarrow t_-} |x(t)| = \infty$ oder es gilt $t_+ < \infty$ und $\lim_{t \rightarrow t_+} |x(t)| = \infty$ oder auch beides.

Wir wollen diesen Satz zunächst weiter erläutern und einige Folgerungen für autonome Systeme ziehen.

Bemerkung 3

1. Die Voraussetzung der lokalen Lipschitz-Stetigkeit mag auf den ersten Blick etwas sperrig und ungewohnt klingen. Allerdings ist jede stetig differenzierbare vektorwertige Funktion automatisch lokal Lipschitz-stetig. Für reellwertige Funktionen wurde das bereits in Aufg. 5.13 gezeigt. Für vektorwertige Funktionen lässt sich das entsprechend zeigen.
2. Die Funktion F könnte auch lediglich auf einer offenen Teilmenge $D \subset \mathbb{R} \times \mathbb{R}^n$ definiert und dort stetig und lokal Lipschitz-stetig in der zweiten Variable sein. Alle entsprechenden Aussagen würden dann nach den nötigen Modifikationen ebenfalls gelten. Wir werden durch die Modellierungen aber bis auf wenige Ausnahmen immer auf Funktionen kommen, die auf ganz $\mathbb{R} \times \mathbb{R}^n$ definiert sind.
3. Der Satz von Picard-Lindelöf sagt aus, dass die Menge $\mathbb{R} \times \mathbb{R}^n$ disjunkt in Lösungsgraphen von (6.4) zerfällt: Durch jeden Punkt (t_0, x_0) verläuft ein Lösungsgraph, nämlich genau der Graph der Lösung des Anfangswertproblems (6.4) und (6.5). Durch einen

Punkt (t_0, x_0) können nicht zwei unterschiedliche Lösungsgraphen von (6.4) laufen, weil sonst die Lösung des Anfangswertproblems nicht mehr eindeutig wäre.

4. Für autonome Differentialgleichungssysteme $F(t, x) = F(x)$ definieren wir außerdem noch den Begriff einer Lösungstrajektorie, auch Orbit genannt: Es sei $t \mapsto x(t)$ eine maximale Lösung von (6.4) mit zugehörigem Existenzintervall I . Dann nennen wir die Menge

$$\{x(t) : t \in I\} \quad (6.7)$$

eine Trajektorie oder einen Orbit von (6.4). Nun zerfällt auch die Menge $\mathbb{R}^n$, das sogenannte Phasenporträt, in disjunkte Trajektorien von (6.4): Klarerweise geht durch jeden Punkt $x_0 \in \mathbb{R}^n$ eine Trajektorie, nämlich z.B. die Trajektorie, die durch die maximale Lösung des AWP's (6.4) mit Anfangswert $(t_0 = 0, x_0)$ erzeugt wird. Wir müssen nun zeigen, dass jede andere Lösung, die zu irgendeinem Zeitpunkt $t_0 \neq 0$ durch den Punkt x_0 verläuft, dieselbe Trajektorie bildet, also denselben Weg im Phasenporträt nimmt, den auch die Lösung, die zur Zeit $t = 0$ den Punkt x_0 passiert, nimmt. Das liegt aber an der Zeitautonomie der Differentialgleichung. In Aufg. 6.1 soll gezeigt werden, dass zwei maximale Lösungen von (6.4), die zu verschiedenen Zeiten durch den gleichen Punkt x_0 verlaufen, Zeittranslationen voneinander sein müssen, also dieselbe Trajektorie bilden.

5. Häufig wird es unser Ziel sein, zu den aus Modellierungen entstandenen DGLS das Phasenporträt möglichst genau zu beschreiben und daraus Rückschlüsse auf die Eigenschaften der Lösungen zu ziehen; insbesondere wird das dann der Fall sein, wenn wir nicht in der Lage sind, explizite Lösungen des DGLS zu berechnen.

6.3 Das Lotka-Volterra-System fortgesetzt – die Methode der Lyapunov-Funktion

Wir wollen nun unsere Erkenntnisse auf das System von Lotka-Volterra anwenden. Das mathematische Ziel ist eine möglichst genaue Kenntnis des Phasenporträts. Zur Erinnerung: Das Lotka-Volterra-System lautet

$$\begin{aligned} \dot{x} &= \alpha x - \beta xy \\ \dot{y} &= -\gamma y + \delta xy. \end{aligned} \quad (6.8)$$

6.3.1 Elementare Diskussion des Phasenporträts

Stationäre Punkte

Wir beginnen damit, die punktförmigen Trajektorien, also die stationären Punkte von (6.8), zu bestimmen. Dazu müssen wir das (nichtlineare) Gleichungssystem

$$\dot{x} = \alpha x - \beta xy = 0 \quad \wedge \quad \dot{y} = -\gamma y + \delta xy = 0$$

nach $x, y \in \mathbb{R}$ lösen. Wir erhalten durch Faktorisieren:

$$\begin{aligned}
 & \alpha x - \beta xy = 0 \quad \wedge \quad -\gamma y + \delta xy = 0 \\
 \Leftrightarrow & \quad x(\alpha - \beta y) = 0 \quad \wedge \quad y(-\gamma + \delta x) = 0 \\
 \Leftrightarrow & (x = 0 \quad \vee \quad \alpha - \beta y = 0) \quad \wedge \quad (y = 0 \quad \vee \quad -\gamma + \delta x = 0) \\
 \Leftrightarrow & (x = 0 \quad \wedge \quad y = 0) \quad \vee \quad \left(y = \frac{\alpha}{\beta} \quad \wedge \quad x = \frac{\gamma}{\delta} \right)
 \end{aligned}$$

Wir haben also zwei stationäre Punkte: $P_0^* = (0, 0)$, den Zustand, in dem es weder eine Räuber- noch eine Beutepopulation gibt, und den Koexistenzzustand $P_1^* = (\frac{\gamma}{\delta}, \frac{\alpha}{\beta})$. Bevor wir in der Analyse des Phasenporträts fortschreiten, wollen wir die Sache etwas übersichtlicher gestalten, indem wir entdimensionalisieren.

Entdimensionalisieren

Wir haben vier Parameter $\alpha, \beta, \gamma, \delta$ und drei Variablen t, x, y . Durch eine Entdimensionalisierung sollten wir also die Anzahl der Parameter auf eins reduzieren können. Wir führen die dimensionslosen Variablen τ, u, v und die zunächst unbestimmten Skalen $\bar{t}, \bar{x}$ und $\bar{y}$ vermöge der Gleichungen

$$\tau = \frac{t}{\bar{t}}, \quad u(\tau) = \frac{x(\bar{t}\tau)}{\bar{x}} \quad \text{und} \quad v(\tau) = \frac{y(\bar{t}\tau)}{\bar{y}}$$

ein. Einsetzen der Differentialgleichungen führt dann auf das dimensionslose Räuber-Beute-System

$$\begin{aligned}
 \dot{u} &= \bar{t}\alpha u - \bar{t}\beta\bar{y}uv \\
 \dot{v} &= -\bar{t}\gamma v + \bar{t}\delta\bar{x}uv.
 \end{aligned}$$

Wir wählen nun die Skalen so, dass die erste Gleichung parameterfrei wird und die zweite Gleichung einen globalen Parameter erhält:

$$\bar{t} = \frac{1}{\alpha}, \quad \bar{y} = \frac{\alpha}{\beta} \quad \text{und} \quad \bar{x} = \frac{\gamma}{\delta}.$$

Die typische Zeit ist also die Zeit, die ein Beutetier durchschnittlich benötigt, um einen Nachkommen zu erzeugen, und die typischen Populationsgrößen sind gerade die des Koexistenzzustands. Mit dem dimensionslosen Parameter $a := \frac{\gamma}{\alpha}$ ergibt sich damit

$$\begin{aligned}
 \dot{u} &= u - uv = u(1 - v) \\
 \dot{v} &= -av + auv = av(-1 + u).
 \end{aligned} \tag{6.9}$$

Aus den stationären Punkten P_0^* und P_1^* werden jetzt die dimensionslosen Punkte $(0, 0)$ und $(1, 1)$, die wir aber weiterhin mit P_0^* und P_1^* bezeichnen werden.

Trajektorien in den Achsen

Wir betrachten zunächst einmal den Fall, dass zu irgendeiner Zeit τ_0 keine Räuberpopulation, $v(\tau_0) = 0$, aber eine Beutepopulation, $u(\tau_0) = u_0 > 0$, vorhanden ist. Dann sind die beiden Funktionen $u(\tau) = u_0 e^{(\tau-\tau_0)}$ und $v(\tau) = 0$ für alle $\tau \in \mathbb{R}$ die eindeutige und maximale Lösung des Anfangswertproblems (6.9) zum Anfangswert $u(\tau_0) = u_0$, $v(\tau_0) = 0$. Die zugehörige Trajektorie verläuft in der u -Achse vom Ursprung nach unendlich. Genauso finden wir eine Lösungstrajektorie in der positiven v -Achse, nämlich die Trajektorie zu der Lösung $u(\tau) = 0$ und $v(\tau) = v_0 e^{-a\tau}$ mit einem $v_0 > 0$. Das zeigt uns, dass der erste Quadrant des Phasenporträts von den anderen drei Quadranten isoliert ist, denn unterschiedliche Trajektorien dürfen sich ja nicht schneiden. Es reicht also für unsere Zwecke, wenn wir uns auf den ersten abgeschlossenen Quadranten des Phasenporträts beschränken, in dem wir bis jetzt vier verschiedene Trajektorien kennen (siehe Abb. 6.1).

Nullklinen und Richtungen

In einem nächsten Schritt bestimmen wir die Punkte im Phasenporträt, in denen die Beute-Änderungsrate oder die Räuber-Änderungsrate verschwindet, man spricht von *Nullklinen*. Aus der ersten Gleichung in (6.9) erhalten wir

$$\dot{u} = 0 \Leftrightarrow u = 0 \quad \vee \quad v = 1$$

und aus der zweiten Gleichung

$$\dot{v} = 0 \Leftrightarrow v = 0 \quad \vee \quad u = 1.$$

Wir zeichnen die Nullklinen ebenfalls in das Phasenporträt ein (siehe Abb. 6.2). Dazu ein paar Bemerkungen:

- Die stationären Punkte sind die Schnittpunkte von u -Nullklinen mit v -Nullklinen.
- Die u -Nullklinen müssen von den Trajektorien rein in v -Richtung, bei unserer Wahl der Achsen also vertikal geschnitten werden; die v -Nullklinen müssen von den Trajektorien rein in u -Richtung, also hier horizontal geschnitten werden.

Abb. 6.1 Erste Skizze des Phasenporträts mit den punktförmigen Trajektorien der stationären Punkte und den Trajektorien in den Achsen

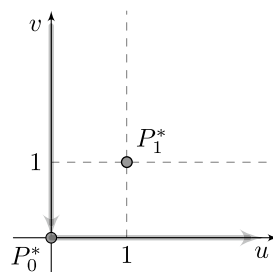
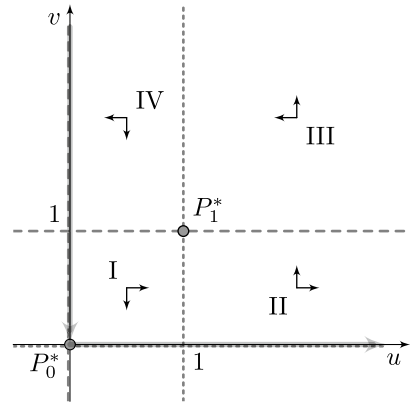


Abb. 6.2 Zweite Skizze des Phasenporträts, Nullklinen und Richtungen. Die u -Nullklinen sind gestrichelt, die v -Nullklinen sind gepunktet



- In den durch die Nullklinen definierten Sektoren herrscht eine einheitliche Richtung der Trajektorien, nämlich durch die Vorzeichen. So müssen die Trajektorien im Sektor I (siehe Abb. 6.2) sich nach rechts unten bewegen, in Sektor II nach rechts oben usw. Die Richtungen sind in Abb. 6.2 durch Pfeile angedeutet.
- Aufgrund der bis jetzt vorliegenden Informationen sind sehr viele, auch qualitativ unterschiedliche Phasenporträts denkbar. Eine Auswahl dreier möglicher Szenarien ist in Abb. 6.3 zusammengestellt: Im ersten Szenario würden die einzelnen Populationen mit wachsenden Amplituden um den Koexistenzzustand oszillieren, im zweiten mit schrumpfenden Amplituden und sich so auf lange Sicht im Koexistenzzustand einfinden; im dritten Szenario hätten wir es mit periodischen Lösungen zu tun. Um eine Entscheidung treffen zu können, reicht es nicht, lediglich die Vorzeichensituation zu nutzen, sondern wir müssen die Abbildungsvorschrift von F in Gänze berücksichtigen. Wir probieren das zunächst mit Hilfe eines numerischen Verfahrens.

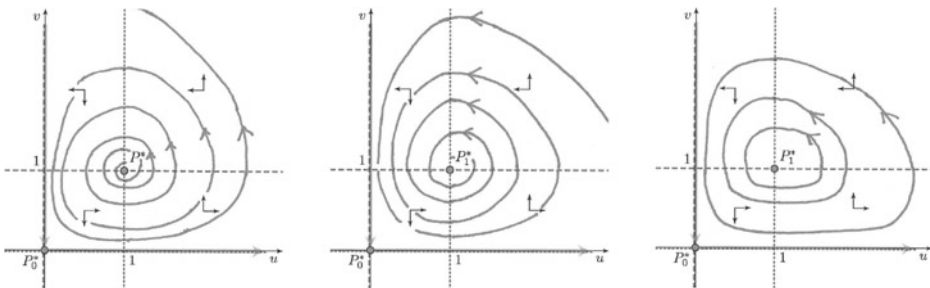


Abb. 6.3 Aufgrund der Vorzeichensituation mögliche Szenarien im Lotka-Volterra-Modell

6.3.2 Eine numerische Annäherung mit dem expliziten Euler-Verfahren

Wir benutzen wieder das explizite Euler-Verfahren. Das funktioniert im Prinzip genauso wie im eindimensionalen Fall (siehe Abschn. 5.7). Wir geben uns eine Anfangszeit τ_0 , in der Regel 0, Anfangswerte u_0 und v_0 sowie eine (kleine) Schrittweite $\Delta\tau > 0$ vor und diskretisieren das Differentialgleichungssystem

$$\tau_{k+1} := \tau_k + \Delta\tau, \quad (u_{k+1}, v_{k+1}) := (u_k, v_k) + \Delta\tau F(u_k, v_k), \quad (6.10)$$

d. h. in unserem Fall des Lotka-Volterra-Systems

$$u_{k+1} = u_k + \Delta\tau u_k(1 - v_k), \quad v_{k+1} = \Delta\tau a v_k(-1 + u_k).$$

Dieses Iterationsschema lässt sich nun einfach z. B. in einem Tabellenkalkulationsdiagramm implementieren. Mit den numerisch berechneten Lösungen können wir einzelne Trajektorien im Phasenporträt und Lösungen gegen die Zeit plotten, siehe Abb. 6.4 und 6.5.

Unsere numerischen Versuche, siehe Abb. 6.4 deuten darauf hin, dass wohl entweder Variante II oder Variante III zutreffen könnten. Die Frage kann aber auf diese Art nicht beantwortet werden: Selbst wenn Variante III zutreffend wäre, könnte das im numerischen Verfahren ja aufgrund der unumgänglichen Approximationsfehler gar nicht zum Tragen kommen. Wir stellen jetzt eine analytische Methode vor, mit der diese Frage geklärt werden kann.

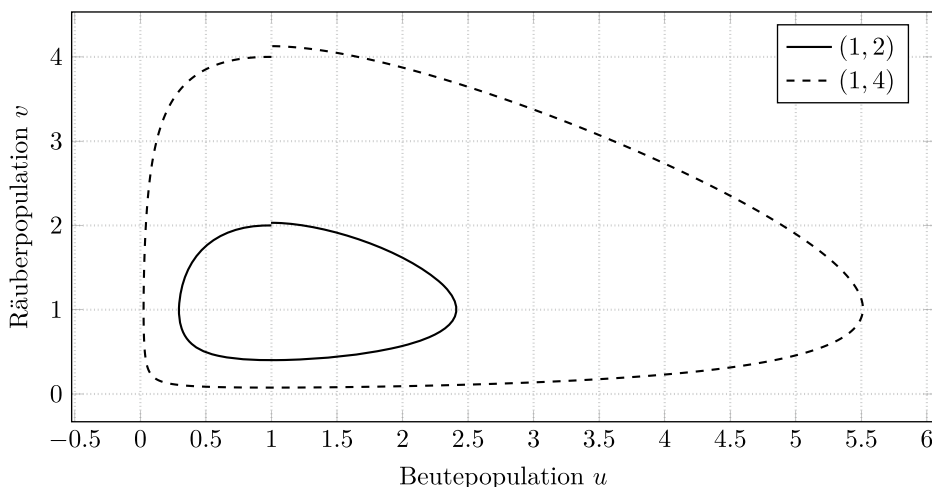


Abb. 6.4 Numerisch erstellte Skizze des Phasenporträts für $a = 0.6$, Schrittweite $\Delta\tau = 0,01$ und Anfangswerte $(1, 2)$, $(1, 4)$ und $(1, 8)$

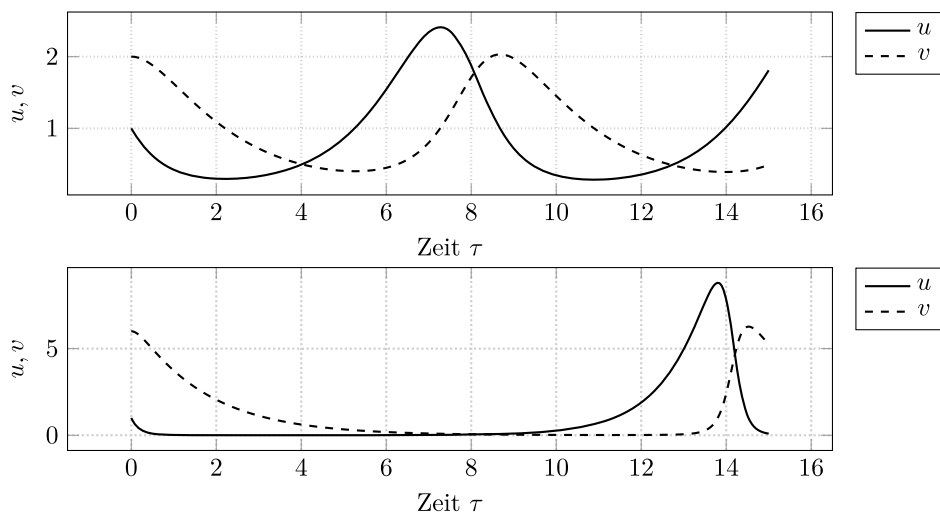


Abb. 6.5 Numerisch erstelltes τ - u , v -Diagramm für die Anfangswerte $\tau_0 = 0$ und $(u_0, v_0) = (1, 4)$ bzw. $(u_0, v_0) = (1, 8)$

6.3.3 Die Lyapunov-Funktion oder das erste Integral

Die Idee besteht darin, nicht direkt explizite Lösungen der Differentialgleichung (6.4) auszurechnen (das ist nicht möglich), sondern lediglich eine Funktion zu berechnen, welche die Trajektorien der Differentialgleichung im ersten Quadranten beschreibt: Gesucht ist eine Funktion $H : (0, \infty) \times (0, \infty) \rightarrow \mathbb{R}$ mit der Eigenschaft, dass sie konstant auf den Lösungstrajektorien von (6.4) ist. Für jede maximale Lösung $\tau \mapsto (u(\tau), v(\tau))$ mit Existenzintervall I , die im ersten Quadranten verläuft, soll es also eine Konstante c geben, so dass

$$H(u(\tau), v(\tau)) = c \quad \text{für alle } \tau \in I. \quad (6.11)$$

Wenn so eine Funktion gefunden ist, wissen wir, dass die zu dieser Lösung gehörige Trajektorie in der Niveaumenge

$$\mathcal{N}_c := \{(u, v) \in (0, \infty) \times (0, \infty) : H(u, v) = c\}$$

enthalten oder „gefangen“ sein muss. Wenn wir uns dann ein Bild über die Niveaumengen von H verschaffen, haben wir auch ein Bild vom Phasenporträt von (6.4). Man nennt so eine Funktion ein *erstes Integral* oder auch eine *Lyapunov-Funktion* nach dem russischen Mathematiker Alexander Michailowitsch Lyapunov (1857–1918). Wie findet man jetzt aber so eine Funktion? Wir schließen dazu rückwärts. Eine Funktion H hat die gewünschte Eigenschaft, wenn für alle möglichen Lösungen $\tau \mapsto (u(\tau), v(\tau))$, die im ersten Quadranten verlaufen, die erste Ableitung nach τ für alle τ aus dem Lösungsintervall verschwindet.

Außerdem wenden wir nacheinander die Kettenregel im Zweidimensionalen an und nutzen aus, dass es sich um Lösungen von (6.9) handelt:

$$\begin{aligned}
 0 &= \frac{d}{d\tau} (H(u(\tau), v(\tau))) \\
 &= \frac{\partial H}{\partial u}(u(\tau), v(\tau)) \cdot \dot{u}(\tau) + \frac{\partial H}{\partial v}(u(\tau), v(\tau)) \cdot \dot{v}(\tau) \\
 &= \frac{\partial H}{\partial u}(u(\tau), v(\tau)) \cdot u(\tau)(1 - v(\tau)) + \frac{\partial H}{\partial v}(u(\tau), v(\tau)) \cdot av(\tau)(-1 + u(\tau))
 \end{aligned} \tag{6.12}$$

Die Gl. (6.12) gilt erst recht längs aller Lösungen von (6.9), wenn sie überhaupt in allen Punkten im offenen ersten Quadranten erfüllt ist: Wir suchen also eine Funktion H mit

$$0 = \frac{\partial H}{\partial u}(u, v) \cdot u(1 - v) + \frac{\partial H}{\partial v}(u, v) \cdot av(-1 + u) \tag{6.13}$$

für alle $(u, v) \in (0, \infty) \times (0, \infty)$. Bis hierhin haben unsere Überlegungen noch keinen Gebrauch von dem konkreten Differentialgleichungssystem gemacht, das wir gerade betrachten, und hätten genauso für jedes beliebige andere DGLS durchgeführt werden können. Jetzt greifen wir aber auf eine spezielle Eigenschaft des betrachteten Systems zurück. Diese Eigenschaft wird nicht jedes DGLS haben, sodass diese Methode des ersten Integrals nur einen eingeschränkten Einsatzbereich hat: Wie wir gleich sehen werden, wäre es sehr günstig, wenn der zweite Faktor im ersten Summanden nur von v und der zweite Faktor im zweiten Summanden nur von u abhinge. Das ist zwar noch nicht der Fall, aber leicht herzustellen, indem wir die Gl. (6.13) durch uv dividieren, man spricht auch von einem *integrierenden Faktor*. Das dürfen wir, da wir ja nur den offenen ersten Quadranten betrachten. Die Aussage in (6.13) ist also äquivalent zu

$$0 = \frac{\partial H}{\partial u}(u, v) \cdot \frac{1 - v}{v} + \frac{\partial H}{\partial v}(u, v) \cdot \frac{a(-1 + u)}{u} \tag{6.14}$$

für alle $(u, v) \in (0, \infty) \times (0, \infty)$. Die Aussage (6.14) ist aber wahr, wenn wir z. B.

$$\frac{\partial H}{\partial u}(u, v) = \frac{a(-1 + u)}{u} \quad \text{und} \quad \frac{\partial H}{\partial v}(u, v) = \frac{-1 + v}{v} \tag{6.15}$$

für alle $(u, v) \in (0, \infty) \times (0, \infty)$ verlangen. Wir können die Funktion H also als Summe zweier Funktionen ansetzen, von denen die eine nur von u und die andere nur von v abhängt:

$$H(u, v) = F(u) + G(v) \quad \text{mit} \quad F'(u) = \frac{a(-1 + u)}{u} \quad \text{und} \quad G'(v) = \frac{-1 + v}{v}.$$

Eine Funktion, die unsere Wünsche erfüllt, ist also z. B.

$$H(u, v) = au - a \ln(u) + v - \ln(v) = af(u) + f(v) \tag{6.16}$$

mit $f(w) = w - \ln(w)$ für $0 < w < \infty$.

Mit einem Plotprogramm, unter GeoGebra funktioniert das mit dem Befehl `ImpliziteKurve`, können wir uns nun einzelne Niveaulinien H für einzelne Werte von a plotten lassen. In Abb. 6.6 sind die Niveaulinien für $a = 1$ zu den Niveaus 2,5, 3 und 4 geplottet. Das Ergebnis scheint darauf hinzudeuten, dass die dritte Alternative in Abb. 6.3 richtig ist.

Wir wollen diese Vermutung mit einem analytischen Argument beweisen. Dazu diskutieren wir die Funktion f und erhalten $\lim_{w \rightarrow 0} f(w) = \lim_{w \rightarrow \infty} f(w) = \infty$, f ist streng monoton fallend auf $(0, 1]$ und streng monoton wachsend auf $[1, \infty)$, hat keine Wendepunkte und das absolute Minimum 1, welches genau an der Stelle 1 angenommen wird. Damit nimmt die Funktion H ihr absolutes Minimum $1 + a$ genau in dem stationären Punkt $P_1^* = (1, 1)$ an. Wir betrachten nun das Bündel aller Geraden, die durch den Punkt $(1, 1)$ gehen. Längs jeder solchen Geraden ist die Funktion H im ersten Quadranten auf der einen Seite des Punktes $(1, 1)$ streng monoton fallend und auf der anderen Seite streng monoton wachsend; ferner wächst H gegen unendlich, wenn man sich längs der Geraden ∞ oder den Achsen nähert. Damit wird jedes Niveau $c > 1 + a$ auf jeder Geraden des Bündels auf jeder Seite von $(1, 1)$ genau einmal angenommen. Die Niveaulinie zu $c > 1 + a$ muss also eine geschlossene Kurve ergeben. (Zugegeben, diese Argumentation ist mathematisch nicht rigoros. Für einen weitergehenden Beweis müssten wir aber einen Begriffsapparat aufbauen, der den Rahmen dieses Buches sprengen würde.)

Zu überlegen bleibt, ob sich in jeder Niveaulinie von H nun auch genau eine Trajektorie befindet oder ob so eine Niveaulinie eventuell in mehrere Trajektorien zerfällt. Dazu überlegen wir uns, dass eine Lösung, die sich für ein $c > 1 + a$ in der Niveaulinie $\mathcal{N}_c$ befindet, diese auch vollständig und dann unendlich oft durchlaufen muss: Eine Lösung,

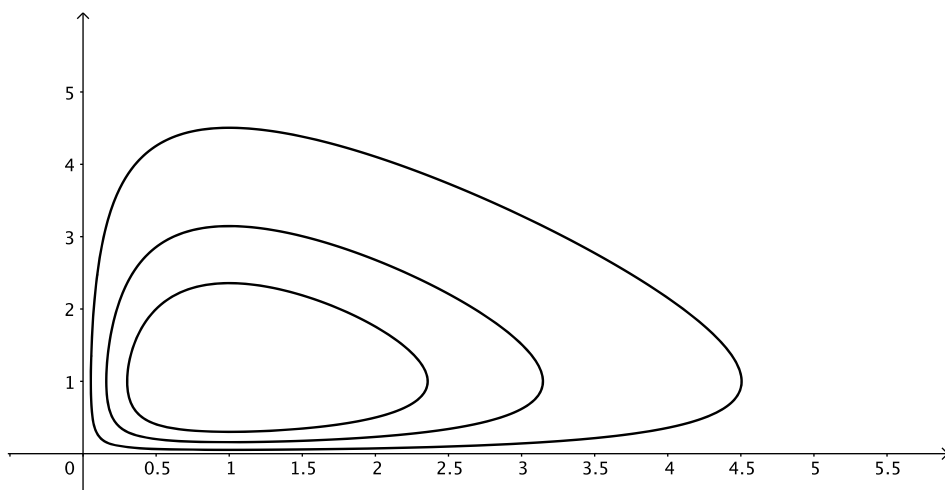


Abb. 6.6 Die Niveaulinien der Funktion H für $a = 1$ zu den Niveaus 2,5, 3 und 4. (Abbildung erstellt mit GeoGebra)

die sich zu einem Zeitpunkt τ_0 im Sektor I, siehe Abb. 6.2 in der Niveaulinie $\mathcal{N}_c$ befindet, muss den Sektor I wieder in Richtung Sektor II verlassen, Entsprechendes gilt für die anderen Sektoren. Angenommen das wäre nicht so und $\tau \mapsto (u(\tau), v(\tau))$ eine Lösung, die durch diese Niveaulinie verläuft, und τ_0 ein Zeitpunkt, zu dem sich die Lösung im Sektor I befindet. Dann wäre die Funktion $\tau_0 \leq \tau \mapsto u(\tau)$ monoton wachsend und durch 1 noch oben beschränkt, die Funktion $\tau_0 \leq \tau \mapsto v(\tau)$ monoton fallend und durch 0 nach unten beschränkt. Insbesondere würde die Lösung für alle $\tau > \tau_0$ existieren, wäre also in positiver Zeitrichtung global. Aufgrund der Beschränktheit und Monotonie gäbe es dann einen Punkt $(\hat{u}, \hat{v})$ mit $\lim_{\tau \rightarrow \infty} (u(\tau), v(\tau)) = (\hat{u}, \hat{v})$ und $0 < \hat{u} \leq 1$ sowie $1 > \hat{v} \geq 0$. Wegen des monotonen Wachstums hätten wir aber auch $\lim_{\tau \rightarrow \infty} \dot{u}(\tau) = 0$ und $\lim_{\tau \rightarrow \infty} \dot{v}(\tau) = 0$. Wegen der Stetigkeit der Funktion $F(u, v) = (u(1-v), av(-1+u))^t$ und der Differentialgleichung müsste $(\hat{u}, \hat{v})$ ein stationärer Punkt sein. Die beiden einzigen stationären Punkte haben wir aber bereits bestimmt und $(\hat{u}, \hat{v})$ gehört klarerweise nicht dazu. Die Lösung muss also unendlich oft durch die Sektoren I, II, III, IV und dann wieder I laufen und dabei den stationären Punkt P_1^* gegen den Uhrzeigersinn umrunden. Damit ist aber die gesamte Niveaulinie durch eine einzige Trajektorie ausgefüllt.

Wir haben damit die in Abschn. 6.3.1 aufgeworfene Frage beantwortet: Das Lotka-Volterra-Modell führt zu periodischen Lösungen, die im Phasenporträt den Koexistenzzustand (bei unserer Achsenwahl) gegen den Uhrzeigersinn umlaufen. Die einzelnen Populationsgrößen sind damit auch periodisch und oszillieren mit Dimensionen um die Größe γ/δ für die Beutepopulation und α/β für die Räuberpopulation. Dabei ist die Umlaufrichtung so, dass die Beutepopulation der Räuberpopulation vorweg zulaufen scheint, vergleiche mit Abb. 6.4. Dieses Verhalten lässt sich plausibel mit dem Sachkontext Räuber-Beute in Beziehung setzen.

6.3.4 Mittelwerte und Abschätzungen für die Periode

Wir wollen nun weitere Gesetzmäßigkeiten aus dem Lotka-Volterra-Modell herausziehen. Die numerischen Berechnungen, siehe Abb. 6.4, deuten darauf hin, dass sich die Populationsgrößen länger in den Bereichen aufhalten, wo die Populationsgrößen klein sind. Das lässt sich auf analytische Weise bestätigen, indem wir die Mittelwerte der Populationsgrößen über eine Periode bestimmen. Wir betrachten eine beliebige Trajektorie zu einem Niveau $c > 1 + a$, also nicht gerade den Koexistenzzustand. Eine Lösung, die in dieser Trajektorie verläuft, hat eine Periode $P = P(c)$, die wir nicht genau kennen. Bei vorgegebenem a und c könnten wir sie numerisch approximieren. Wir definieren die *mittleren Populationsgrößen* als

$$\bar{u} = \frac{1}{P} \int_{\tau_0}^{\tau_0+P} u(\tau) d\tau \quad \text{und} \quad \bar{v} = \frac{1}{P} \int_{\tau_0}^{\tau_0+P} v(\tau) d\tau. \quad (6.17)$$

Nun nutzen wir die Differentialgleichung für v , $\dot{v} = av(-1+u)$, lösen diese nach u auf, setzen den berechneten Term in das Integral ein und erhalten damit

$$\bar{u} = \frac{1}{P} \int_{\tau_0}^{\tau_0+P} \left(1 - \frac{\dot{v}(\tau)}{av(\tau)} \right) d\tau = \frac{1}{P} \left(\tau - \frac{1}{a} \ln(v(\tau)) \right) \Big|_{\tau_0}^{\tau_0+P} = 1, \quad (6.18)$$

wobei wir ausgenutzt haben, dass v periodisch mit Periode P ist. Mit einer entsprechenden Rechnung erhalten wir $\bar{v} = 1$. Damit haben wir die sogenannte zweite Lotka-Volterra-Regel gezeigt: Die Mittelwerte der Populationen hängen nicht von den Anfangsbedingungen ab (es ist egal, welche zyklische Trajektorie wir betrachten), sondern nur von den Parametern α , β , γ und δ .

Können wir außerhalb numerischer Berechnungen noch analytisch etwas über die Periode aussagen? Zumindest über Trajektorien in der Nähe des Koexistenzzustandes können wir etwas durch Linearisierung herausfinden: Dazu betrachten wir eine Trajektorie zu einem kleinen Niveau $c \approx 1 + a$, die den Koexistenzpunkt sehr eng umläuft. Das ermöglicht uns, für eine Lösung in dieser Trajektorie das nichtlineare Lotka-Volterra-System, gegeben durch die Funktion F , durch die um den Koexistenzzustand linearisierte Version zu ersetzen:

Wir haben diese Prozedur bereits bei den Differentialgleichungen in Kap. 5 kennengelernt. Es sei $\tau \mapsto (u(\tau), v(\tau))$ die betrachtete Lösung und es gelte $|u(\tau) - 1| \ll 1$ und $|v(\tau) - 1| \ll 1$ für alle τ . Dann definieren wir die Auslenkungen aus dem Koexistenzzustand $m(\tau) = u(\tau) - 1$ und $n(\tau) = v(\tau) - 1$. Diese Auslenkungen erfüllen dann exakt bzw. approximativ

$$\begin{aligned} \begin{pmatrix} \dot{m} \\ \dot{n} \end{pmatrix} &= \begin{pmatrix} \dot{u} \\ \dot{v} \end{pmatrix} = F(u, v) = F(1 + m, 1 + n) \\ &\approx F(1, 1) + J_F(1, 1) \begin{pmatrix} m \\ n \end{pmatrix} = J_F(1, 1) \begin{pmatrix} m \\ n \end{pmatrix}; \end{aligned} \quad (6.19)$$

dabei ist $J_F(1, 1)$ die Jacobi-Matrix von F an der Stelle $(1, 1)$ und wir haben benutzt, dass $(1, 1)$ ein stationärer Punkt von F ist. Wir berechnen die Jacobi-Matrix und erhalten das folgende approximative lineare Differentialgleichungssystem für m und n :

$$\begin{pmatrix} \dot{m} \\ \dot{n} \end{pmatrix} \approx \begin{pmatrix} 0 & -1 \\ a & 0 \end{pmatrix} \begin{pmatrix} m \\ n \end{pmatrix}$$

Wir werden in Abschn. 6.4 lernen, wie solche Differentialgleichungssysteme gelöst werden, können aber hier die Lösungen auch direkt erkennen: m löst approximativ $\ddot{m} \approx -am$ und n löst $\ddot{n} \approx -an$. Die exakten Lösungen dieser Gleichungen sind gegeben durch lineare Überlagerungen der Funktionen $\tau \mapsto \cos(\sqrt{a}\tau)$ und $\tau \mapsto \sin(\sqrt{a}\tau)$, die beide die Periode $P = \frac{2\pi}{\sqrt{a}}$ haben. Somit haben auch die exakten Lösungen des linearen Differentialgleichungssystems diese Periode.

Wir erwarten also, dass für Trajektorien nah am Koexistenzzustand sich die Lösungen harmonischen Funktionen annähern und wir für die Perioden die Grenzwertbeziehung

$$\lim_{c \rightarrow 1+a} P(c) = \frac{2\pi}{\sqrt{a}}$$

erhalten. Es sei aber angemerkt, dass diese Überlegung bis jetzt heuristischer Natur ist: Aus der Approximation der Funktion F durch ihre Linearisierung um $(1, 1)$ haben wir gefolgert, dass sich auch die Lösungen der unterschiedlichen Differentialgleichungen in einem genauer zu definierenden Sinne approximieren. Das klingt zwar plausibel, muss aber, um in den Rang eines mathematischen Satzes erhoben zu werden, präzisiert und bewiesen werden. Das ist gerade in dem betrachteten Fall schwierig, weil die Jacobi-Matrix $J_F(P_1^*)$ rein imaginäre Eigenwerte hat. Wir werden in Abschn. 6.4 darauf eingehen. Unseren mit heuristischen Methoden gewonnenen Grenzwert können wir aber auch durch numerische Evidenz stützen: Mit $a = 0.6$, der Schrittweite $\Delta\tau = 0.01$ und dem Anfangswert $(u_0, v_0) = (1, 1.01)$, der sich also auf der Trajektorie zum Niveau $c = H(1, 1.01) = 0.6 + 1.01 - \ln(1.01) = 1.60005 \approx 1 + a$ befindet, erhalten wir mit dem Euler-Verfahren eine Umlaufzeit von $P_{\text{num}} = 8.112$, die in guter Übereinstimmung mit dem theoretischen Wert $P = \frac{2\pi}{\sqrt{0.6}} = 8.1116$ ist. Weitere numerisch ermittelte Zusammenhänge zwischen dem Niveau c der Trajektorie und der Periode $P(c)$ sind in Abb. 6.7 dargestellt und scheint auf einen linearen Zusammenhang hinzudeuten. Dieser Zusammenhang soll in Aufg. 6.5 numerisch untersucht werden. Analytische Ergebnisse finden sich in [77] und [91].

6.3.5 Anwendung auf die Interaktion von Luchs und Schneeschuhhase

Wir wollen nun die Aussagen, die das Lotka-Volterra-System liefert, mit den Populationsentwicklungen bestimmter Räuber-Beute-Populationen vergleichen. Ein unter den Ökologen besonders prominentes Beispiel ist die Räuber-Beute-Interaktion zwischen Schneeschuhhasen- und Luchspopulationen in den unterschiedlichen Regionen Kanadas.

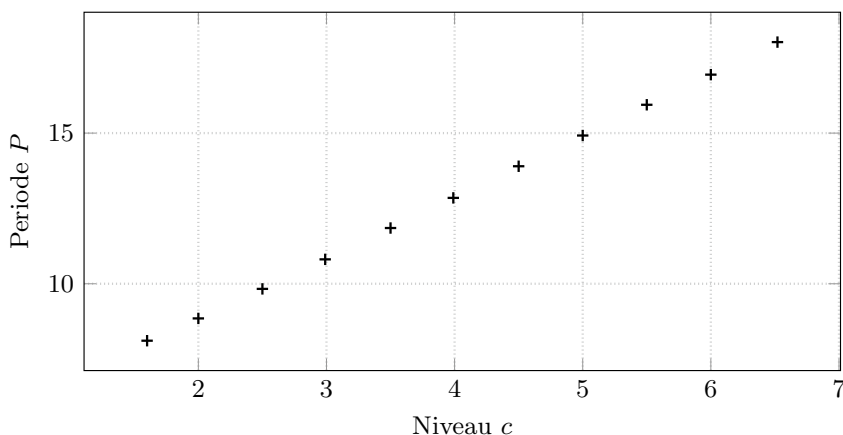


Abb. 6.7 Numerisch ermittelter Zusammenhang zwischen dem Niveau c und der Periode P für $a = 0.6$. Explizites Euler-Verfahren mit Schrittweite $\Delta\tau = 0.01$

Das hat seinen Grund darin, dass die Hudson's Bay Company Exportzahlen von Hasen- und Luchsfellen nach England über einen Zeitraum von über 200 Jahren aufgezeichnet hat, ein in der Populationsökologie konkurrenzloser Datenschatz. Zu beachten ist dabei allerdings, dass durch diese Exportzahlen die Populationsgröße nur auf eine sehr indirekte Art und Weise gemessen wird: Es kann durchaus sein, dass diese Zahlen variieren, weil sich die Anzahl der menschlichen Jäger verändert hat, weil die Nachfrage in England geringer oder größer geworden ist. Eventuell werden Exportzahlen nicht dem Jahr zugeschrieben, in denen die Tiere erlegt wurden, sondern erst späteren, und so gibt es viele Einflüsse, die auf diese Zahlen einwirken und nichts mit den Größen der Hasen- und Luchspopulationen zu tun haben. Trotzdem kann man die Zahlen, wenn auch nur mit Vorsicht, heranziehen, um etwas über die zeitliche Entwicklung der Populationen auszusagen. Wir betrachten hier die Exportzahlen für Luchsfelle aus den Jahren 1735 bis 1820 (siehe Abb. 6.8).

Dabei gibt die gestrichelte Kurve die Exportzahlen aus den originalen Verkaufsbüchern der Hudson's Bay Company an, die gepunktete Kurve die Verkaufszahlen, wie sie von Poland in [73] publiziert wurden. In den Jahren 1751 bis 1778 sowie 1794 bis 1798 liegen Zahlen aus beiden Quellen vor. Eine ausführliche Diskussion und Darstellung der Daten findet sich in [23].

Auffällig ist die über einen Zeitraum von knapp 100 Jahren regelmäßige Abfolge von Maxima, die durchgehend einen Abstand von zehn Jahren zu haben scheinen. Würden wir weitere Daten hinzuziehen, bliebe das Pattern erhalten, siehe [23]. Gleichzeitig ist die Amplitude, also die Höhe des Maximums, stark variierend, was zum einen auf unterschiedlich große Bestände, aber auch auf eine unterschiedliche Bejagungsintensität hindeuten kann.

Vergleichen wir nun diesen empirischen Befund mit den Voraussagen unseres Modell fällt erst einmal die Periodizität auf. Auf den ersten Blick scheint das gut mit dem Modell über-

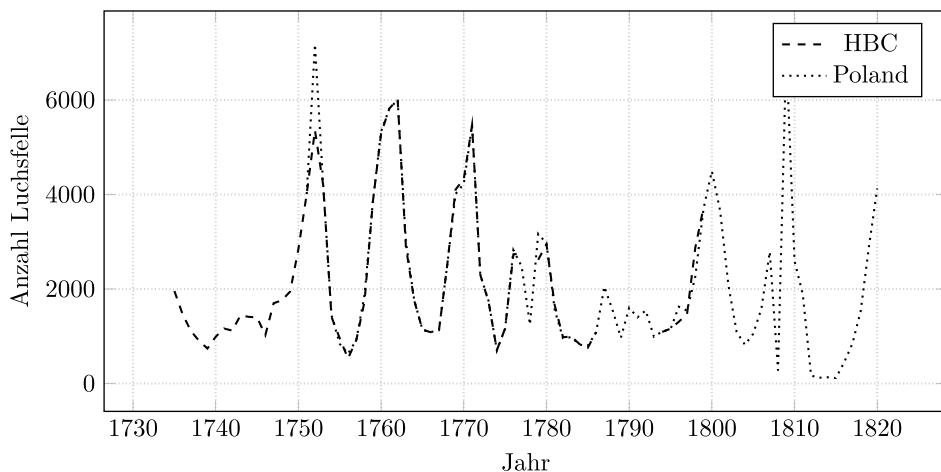


Abb. 6.8 Entwicklung der Verkaufszahlen von Luchsfellen der Hudson's Bay Company

einzustimmen, auf den zweiten Blick ist aber die stabile Periode von 10 Jahren bemerkenswert. Nach unserem Modell würde sich die Periode mit der Amplitude verändern: Kämen die beiden Populationen durch irgendwelche nichtmodellierten Einflüsse, die es natürlich immer gibt, auf eine weiter außen liegende Trajektorie (größeres c), würden wir auch eine größere Periode erwarten. Darauf haben auf jeden Fall unsere numerischen Resultate aus Abschn. 6.3.2, Abb. 6.7, hingedeutet. In den Daten ist also tatsächlich mehr Stabilität vorhanden, als wir aufgrund unseres Modells erwarten würden. In unserer Modellierung scheinen also nicht nur zufällige Einflüsse (wie unterschiedliche Wetterverhältnisse und Ähnliches) zu fehlen, sondern etwas Wesentliches. Eine gewisse Klärung können wir in Abschn. 6.5 herbeiführen.

Es gibt eine weitere berühmte Studie mit dem schönen Titel „*Do hares it lynx?*“, welche die empirischen Daten der Hasenfell-Exporte mit denen der Luchsfell-Exporte in Verbindung bringt, nämlich [34]. Diese Studie entnimmt ihre Daten der Jahre 1875 bis 1906 ebenfalls den Archiven der Hudson's Bay Company. Tragen wir diese Daten in einem Hasen-Luchs-Diagramm auf (siehe Abb. 6.9) und verbinden die Jahre der Reihenfolge nach, lassen sich drei Zyklen erkennen, die allerdings im Uhrzeigersinn anstatt gegen den Uhrzeigersinn durchlaufen werden. Im Sinne des Modells wirkt es also, als ob die Hasen die Luchse fressen würden. Es ist viel darüber spekuliert worden, wie diese erstaunlichen Zahlen zu interpretieren seien und welche Rückschlüsse sie auf das Lotka-Volterra-Modell erlauben. So wurde vermutet, dass die Rolle des Menschen als Felljäger einen so großen Einfluss habe, dass sich die Hase-Luchs-Beziehung nicht mehr sinnvoll als Interaktionsmodell zwischen zwei Spezies darstellen lasse. Eine andere Hypothese geht davon aus, dass tatsächlich die Hasen durch Übertragung von Krankheitserregern die Räuber-Rolle übernehmen könnten. Eine lapidare

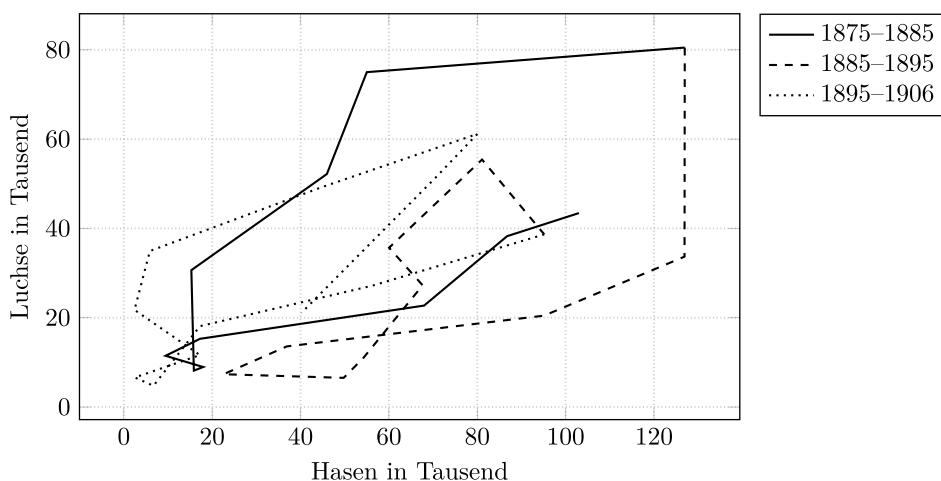


Abb. 6.9 Populationsgrößen von Schneeschuhhasen und Luchsen aus den Jahren 1875 bis 1906. (Daten aus [34])

Erklärung weist darauf hin, dass sich die in der Studie verwendeten Luchsfell-Exportzahlen auf West- und die verwendeten Hasenfell-Exportzahlen auf Ostkanada beziehen.

Wir werden in den nächsten Kapiteln das Lotka-Volterra-System modifizieren und dazu einige mathematische Hilfsmittel einführen.

6.4 Räuber-Beute-Systeme mit logistischem Wachstum – lineare Differentialgleichungssysteme und der Linearisierungssatz von Grobman-Hartman

6.4.1 Modellbildung und erste Diskussion des Modells

Wir haben gesehen, dass wir dem Lotka-Volterra-Modell mehr Stabilität geben müssen, um die empirischen Befunde halbwegs erklären zu können. Eine erste naheliegende Modifikation besteht in der Ersetzung des exponentiellen Wachstums der Beutepopulation durch ein logistisches Wachstum, also die Berücksichtigung von intraspezifischen Konkurrenzeffekten (Konkurrenz innerhalb der Spezies).

Wir erhalten das System

$$\begin{aligned}\dot{x} &= \alpha x \left(1 - \frac{x}{K}\right) - \beta xy \\ \dot{y} &= -\gamma y + \delta xy,\end{aligned}\tag{6.20}$$

hierbei ist K die Kapazität des betrachteten Gebiets bezüglich der Beutepopulation, siehe Kap. 5. Natürlich könnten wir auch noch einen intraspezifischen Konkurrenzeffekt in die Differentialgleichung für die Räuberpopulation einbauen. Das soll in einer Übungsaufgabe gemacht werden, siehe Aufg. 6.6. Dabei stellt sich allerdings heraus, dass dieses Modell gegenüber unserem (6.20) keine wesentlich anderen Aussagen liefert.

Zur Vereinfachung entdimensionalisieren wir wieder und erhalten mit den gleichen Skalen wie in Abschn. 6.1, also $\bar{t} = \frac{1}{\alpha}$, $\bar{x} = \frac{\gamma}{\delta}$ und $\bar{y} = \frac{\alpha}{\beta}$, die entdimensionalisierten Gleichungen

$$\begin{aligned}\dot{u} &= u(1 - bu) - uv \\ \dot{v} &= av(-1 + u)\end{aligned}\tag{6.21}$$

mit dem bereits bekannten dimensionslosen Parameter $a = \frac{\alpha}{\gamma}$ und dem neuen dimensionslosen Parameter $b = \frac{\gamma}{\delta K}$.

Wir gehen analog zum Lotka-Volterra-System vor, bestimmen zunächst die Nullklinen und stationären Punkte und untersuchen das System auf Trajektorien in den Achsen. In der v -Achse haben wir Lösungen der Form $u = 0$ und $v(\tau) = v_0 e^{-\tau}$. In der u -Achse befinden sich Lösungen mit $v = 0$ und u ist eine Lösung der logistischen Gleichungen $\dot{u} = u(1 - bu)$ mit der stabilen Lösung $u^* = 1/b$. Elementare Rechnungen führen uns auf die Nullklinen

$$\dot{u} = 0 \Leftrightarrow u = 0 \vee v = 1 - bu \quad \text{sowie} \quad \dot{v} = 0 \Leftrightarrow v = 0 \vee u = 1.$$

Daraus ergeben sich die stationären Punkte $P_0^* = (0, 0)$, $P_1^* = (1/b, 0)$ und $P_2^* = (1, 1 - b)$. In diesem Modell gibt es also zwei qualitativ unterschiedliche Fälle: $0 < b < 1$, entsprechend $K > \frac{\gamma}{\delta}$, und $1 < b$, also $\frac{\gamma}{\delta} > K$. Im ersten Fall liegt der Koexistenzzustand P_1^* im ersten Quadranten, im zweiten Fall nicht. Wir skizzieren die Phasenporträts, soweit das bis jetzt möglich ist (siehe Abb. 6.10). In beiden Fällen sind aufgrund der vorliegenden Informationen noch viele, auch qualitativ sehr unterschiedliche Phasenporträts möglich. Wir benötigen also noch ein paar weitere Informationen.

Im Fall des Lotka-Volterra-Modells hat uns hier das erste Integral weitergeholfen: Wir konnten eine Funktion auf dem ersten Quadranten finden, die konstant auf den Lösungen des Lotka-Volterra-Systems ist. Wir versuchen das auch in diesem Fall. Ein analoges Vorgehen zu Abschn. 6.1 führt auf die Forderung

$$0 = \frac{\partial H}{\partial u}(u, v) \cdot u(1 - bu - v) + \frac{\partial H}{\partial v}(u, v) \cdot av(-1 + u) \quad (6.22)$$

für alle $(u, v) \in (0, \infty) \times (0, \infty)$.

Auch eine Division durch uv führt nun nicht mehr zu der günstigen Situation, dass wir die Funktion H als Summe zweier Funktionen ansetzen können, von denen die eine nur von u und die andere nur von v abhängt; der kleine unscheinbare Term $-bu$ hat diese Struktur zerstört. Wir müssten jetzt eine partielle Differentialgleichung lösen, was in der Regel schwieriger ist, als mit einem gewöhnlichen Differentialgleichungssystem umzugehen.

Eine andere Idee wäre, bei unserer alten H -Funktion vom Lotka-Volterra-System zu bleiben und zu hoffen, dass diese ein monotones Verhalten auf Lösungstrajektorien von (6.21) aufweist, dass also zum Beispiel für jede Lösung $\tau \mapsto (u(\tau), v(\tau))$ von (6.21) die Abschätzung

$$\frac{d}{d\tau} (H(u(\tau), v(\tau))) < 0$$

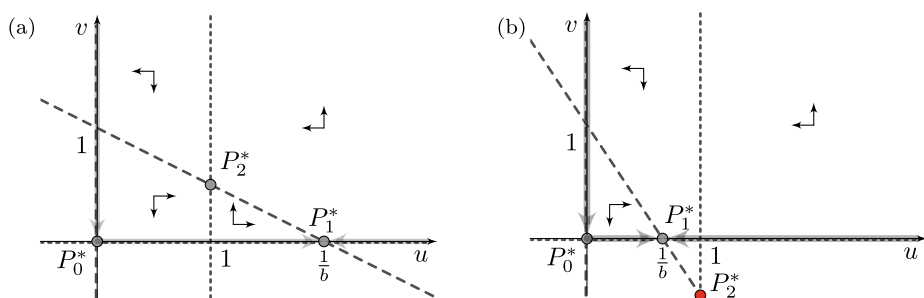


Abb. 6.10 Erste Skizze des Phasenporträts von (6.21) für (a) $0 < b < 1$ und (b) $1 < b$. u -Nullklinen gestrichelt, v -Nullklinen gepunktet

für alle τ aus dem maximalen Lösungsintervall korrekt wäre. Dann müssten sich nämlich Trajektorien zur Minimalstelle der H -Funktion hin bewegen. In einigen Anwendungen vor allem aus der Physik führt so eine Überlegung auch zum Ziel. In Aufg. 6.2 wird das mathematische Pendel mit Reibung behandelt. Dort ist diese Überlegung tatsächlich erfolgreich. Hier kann das aber nicht klappen, weil unser neuer stationärer Punkt $P_2^* = (1, 1 - b)$ nicht in der Minimalstelle der H -Funktion, nämlich in dem Punkt $(1, 1)$, dem stationären Punkt des Lotka-Volterra-Systems, liegt.

Wir benötigen also neue Methoden, die in dem kommenden Unterkapitel zusammenge stellt werden. Die zugrunde liegende Idee ist bestens bekannt: Wenn eine Funktion oder Abbildung zu schwierig ist, dann linearisiere sie um die Stelle oder den Punkt, der gerade von Interesse ist, und studiere die linearisierte Variante. Sei also $P^* = (u^*, v^*)$ ein stationärer Punkt irgendeines dynamischen Systems

$$\begin{pmatrix} \dot{u} \\ \dot{v} \end{pmatrix} = F(u, v) \quad (6.23)$$

und wir führen die Auslenkungen aus dem stationären Punkt als neue Variablen ein: $r = u - u^*, s = v - v^*$. Falls $|r|, |s| \ll 1$, können wir die Abbildung F durch ihre Linearisierung um P^* ersetzen und machen dabei nur einen kleinen Fehler:

$$\begin{aligned} \begin{pmatrix} \dot{r} \\ \dot{s} \end{pmatrix} &= \begin{pmatrix} \dot{u} \\ \dot{v} \end{pmatrix} = F(u^* + r, v^* + s) \\ &\approx F(u^*, v^*) + J_F(u^*, v^*) \begin{pmatrix} r \\ s \end{pmatrix} = J_F(u^*, v^*) \begin{pmatrix} r \\ s \end{pmatrix}; \end{aligned}$$

hierbei ist J_F die Jacobi-Matrix von F und die letzte Gleichung gilt, weil (u^*, v^*) ein stationärer Punkt ist. Wir haben dieses Vorgehen auch bereits bei der Bestimmung der Perioden im Grenzfall kleiner Amplituden in Abschn. 6.3 verfolgt.

Es sind also zwei Dinge zu tun: Wir müssen in einem ersten Schritt das Lösungsverhalten des linearen Differentialgleichungssystems

$$\begin{pmatrix} \dot{r} \\ \dot{s} \end{pmatrix} = J_F(u^*, v^*) \begin{pmatrix} r \\ s \end{pmatrix} \quad (6.24)$$

untersuchen. Das wird allgemein für den zweidimensionalen Fall in Abschn. 6.4.2 geschehen.

In einem zweiten Schritt in Abschn. 6.4.4 ist dann zu klären, in welchem Verhältnis die Lösungen des linearisierten Gleichungssystems zu den eigentlich interessierenden Lösungen des nichtlinearen Gleichungssystems stehen. Dazu werden wir den Satz von Grobman und Hartman heranziehen, den wir in diesem Buch nur zitieren und erläutern können, für einen Beweis verweisen wir auf die Darstellung in [87].

6.4.2 Lineare Differentialgleichungssysteme

Die Linearisierung des nichtlinearen Differentialgleichungssystems (6.23) um einen stationären Punkt führte uns auf ein lineares Gleichungssystem für die Auslenkungen aus dem stationären Punkt (6.24), das wir hier in einer allgemeineren Form notieren wollen:

$$\dot{\vec{x}} = A\vec{x}, \quad (6.25)$$

wobei $\vec{x} = (x_1, \dots, x_n)^t$ ein reelles Intervall I in den $\mathbb{R}^n$ abbildet und A eine $n \times n$ -Matrix mit reellen Einträgen ist. In unseren Anwendungen wird die Matrix A die Ableitungsmatrix der nichtlinearen Funktion F an einem stationären Punkt sein. Wir werden ausschließlich hier in Abschn. 6.4.2 auf eine Notation mit Pfeilen zurückgreifen, um zwischen der vektorwertigen Funktion $\vec{x}(t)$ und ihren Koordinaten $x_i(t)$ bequem unterscheiden zu können.

Wir werden uns zuerst überlegen, dass die Lösungen von (6.25) einen n -dimensionalen Vektorraum sowohl über $\mathbb{R}$ als auch über $\mathbb{C}$ bilden, und danach ausreichend viele linear unabhängige Lösungen explizit berechnen, die diesen Raum aufspannen. Den letzten Schritt werden wir allerdings nur für $n = 2$ durchführen: Die meisten folgenden Modellierungen dieses Buches werden zu zweidimensionalen Systemen führen und wir können dadurch die nötigen Notationen schlank halten. Andererseits sind die Überlegungen hinreichend allgemein, so dass sie im Wesentlichen ohne neue Ideen auf den allgemeinen Fall eines n -dimensionalen Systems verallgemeinert werden können.

Der n -dimensionale Lösungsraum

Die wichtigste Eigenschaft von linearen Gleichungen ist, dass Linearkombinationen von Lösungen ebenfalls Lösungen sind: Seien $\vec{x}_1$ und $\vec{x}_2$ zwei Lösungen von (6.25) über einem gemeinsamen Lösungsintervall I , $\alpha, \beta \in \mathbb{R}$ oder $\mathbb{C}$, dann ist auch $\alpha\vec{x}_1 + \beta\vec{x}_2$ eine Lösung von (6.25) über dem Intervall I . Wir betrachten nun Lösungen zu speziellen Anfangswertproblemen: Sei nämlich $\vec{x}_i$ die maximale Lösung von (6.25), die zur „Zeit“ $t = 0$ durch den i -ten Einheitsvektor verläuft mit Existenzintervall I_i (wir werden später sehen, dass diese Lösungen global sind, also auf ganz $\mathbb{R}$ existieren), und I der Schnitt über alle I_i . Nach dem Satz von Picard-Lindelöf ist die Lösung eindeutig (auch über jedem noch so kleinen offenen Intervall, das die Null enthält). Sei nun irgendein Anfangswert $\vec{x}_0 = (\alpha_1, \dots, \alpha_n) \in \mathbb{R}^n$ vorgegeben. Dann wird dieses Anfangswertproblem durch die Funktion $\vec{x} = \alpha_1\vec{x}_1 + \dots + \alpha_n\vec{x}_n$ gelöst, und zwar eindeutig. Damit sind die Funktionen $\vec{x}_1, \dots, \vec{x}_n$ ein Erzeugendensystem. Linear unabhängig sind sie ebenfalls: Würden sie nämlich die Nullfunktion linear kombinieren, müssten sie insbesondere für $t = 0$ den Nullvektor im $\mathbb{R}^n$ linear kombinieren. Dafür müssten aber alle Koeffizienten der Linearkombination verschwinden; die Vektoren $\vec{x}_1, \dots, \vec{x}_n$ bilden also eine Basis des Vektorraums der Lösungen von (6.25). Wir wissen also, dass es ausreichend ist, n linear unabhängige Lösungen von (6.25) zu bestimmen, daraus lassen sich dann alle Lösungen linear kombinieren.

Fundamentallösungen für $n = 2$

Haben wir eine Basis der Lösungen von (6.25) vorliegen, nennen wir diese Lösungen häufig auch ein Fundamentalsystem oder Fundamentallösungen. Inspiriert von den Lösungen der eindimensionalen linearen Gleichung $\dot{x} = \alpha x$ probieren wir einen Ansatz aus: Wie wäre es denn mit einer Funktion der Form $\vec{x}(t) = \vec{v}e^{\lambda t}$, wobei $\vec{v} \in \mathbb{R}^2$ ein konstanter Vektor ist und $\lambda \in \mathbb{R}$? Wenn diese Funktion eine Lösung von (6.25) sein soll, muss für alle $t \in \mathbb{R}$ gelten, dass

$$\dot{\vec{x}}(t) = \lambda \vec{v}e^{\lambda t} \stackrel{!}{=} A\vec{x}(t) = A\vec{v}e^{\lambda t}.$$

Das ist aber äquivalent zu der Forderung $A\vec{v} = \lambda \vec{v}$. Mit anderen Worten: Der Ansatz funktioniert genau dann, wenn $\vec{v}$ ein Eigenvektor von A zum Eigenwert λ ist. Damit ist das Problem, die Lösungen eines linearen Differentialgleichungssystems zu finden, zurück geführt worden auf das Problem, linear unabhängige Eigenvektoren für eine 2×2 -Matrix zu finden. Aber nicht jede reelle 2×2 -Matrix hat zwei linear unabhängige reelle Eigenvektoren. Wir müssen die Fälle also einzeln betrachten.

Bekanntermaßen sind die Eigenwerte einer Matrix A genau die Nullstellen des zugehörigen charakteristischen Polynoms

$$P_A(\lambda) = (a_{1,1} - \lambda)(a_{2,2} - \lambda) - a_{1,2}a_{2,1} = \lambda^2 - \text{tr}A\lambda + \det A, \quad (6.26)$$

wobei hier $a_{i,j}$ die Koeffizienten, $\text{tr}A := a_{1,1} + a_{2,2}$ die Spur und $\det A$ die Determinante der Matrix A sind. Wir gehen jetzt alle möglichen Fälle durch, die auftreten könnten.

1. Im ersten Fall hat die Matrix A zwei verschiedene reelle Eigenwerte $\lambda_1 \neq \lambda_2$. Dazu gibt es jeweils Eigenvektoren $\vec{v}_1$ und $\vec{v}_2$, die linear unabhängig sind. Damit bilden die beiden Lösungen

$$\vec{x}_1(t) = \vec{v}_1 e^{\lambda_1 t} \quad \text{und} \quad \vec{x}_2(t) = \vec{v}_2 e^{\lambda_2 t} \quad (6.27)$$

ein Fundamentalsystem. Der erste Fall tritt genau dann ein, wenn $\text{tr}^2 A - 4 \det A > 0$ gilt, wie sich aus der p - q -Formel ablesen lässt.

2. Im zweiten Fall hat P_A einen doppelten reellen Eigenwert λ_0 , also $P_A(\lambda) = (\lambda - \lambda_0)^2$. Hier müssen wir unterscheiden:

- Haben wir zwei linear unabhängige Eigenvektoren $\vec{v}_1$ und $\vec{v}_2$ zum selben Eigenwert λ_0 (geometrische Vielfachheit gleich algebraische Vielfachheit gleich zwei für alle, die die Behandlung der Jordan-Normalform kennen), dann ist jeder beliebige Vektor ein Eigenvektor, A nichts anderes als ein Vielfaches der Einheitsmatrix, und z.B. $(1, 0)^t e^{\lambda_0 t}$ und $(0, 1)^t e^{\lambda_0 t}$ bilden ein Fundamentalsystem.
- Ist der Eigenraum jedoch nur eindimensional, kommen wir mit unserem Ansatz „Eigenvektor mal e hoch Eigenwert mal t “ nicht hin, es gibt ja keine zwei linear unabhängigen Eigenvektoren. Die lineare Algebra hilft aber auch in diesem Fall weiter: Wir wählen einen Vektor $\vec{w}$, der kein Eigenvektor von A ist, es gilt also

$$\vec{v} := A\vec{w} - \lambda_0\vec{w} \neq \vec{0}. \quad (6.28)$$

Nun ist nach dem Satz von Cayleigh-Hamilton, siehe z. B. [28], die Matrix $(A - \lambda_0)^2$ die Nullmatrix. Wir haben also $(A - \lambda_0)\vec{v} = (A - \lambda_0)^2\vec{w} = \vec{0}$. Mit anderen Worten: $\vec{v}$ ist ein Eigenvektor von A zum Eigenwert λ_0 . Eine der beiden gesuchten Lösungen erhalten wir wieder durch unseren alten Ansatz: $\vec{x}_1(t) = \vec{v}e^{\lambda_0 t}$. Die zweite Lösung setzen wir jetzt mit Hilfe der Vektoren $\vec{v}$ und $\vec{w}$ zusammen: $\vec{w}e^{\lambda_0 t}$ alleine tut es nicht, aber

$$\vec{x}_2(t) := e^{\lambda_0 t} (\vec{w} + t\vec{v}) \quad (6.29)$$

ist eine Lösung von (6.25), wie sich unter Zuhilfenahme der Gl. (6.28) nachrechnen lässt. Die beiden Lösungen $\vec{x}_1, \vec{x}_2$ sind linear unabhängig: Es gilt ja $\vec{x}_1(0) = \vec{v}$ und $\vec{x}_2(0) = \vec{w}$ und die beiden Vektoren $\vec{v}$ und $\vec{w}$ sind linear unabhängig, andernfalls wäre ja auch $\vec{w}$ ein Eigenvektor von A .

3. Es verbleibt der dritte und letzte Fall: Die Matrix A hat überhaupt keinen reellen Eigenwert, sondern zwei komplexe Eigenwerte λ_1 und λ_2 , die in diesem Fall komplex konjugiert zueinander sind: $\lambda_{1/2} = \alpha \pm i\beta$ mit $\alpha = \text{tr} A/2$ und $\beta = \sqrt{\det A - \text{tr}^2 A/4} > 0$. Wir haben also $\bar{\lambda}_1 = \lambda_2$. Wir nutzen einen Umweg über die komplexen Zahlen und die komplexe Exponentialfunktion: Nach der Euler'schen Formel gilt für beliebige reelle Zahlen r_1, r_2 die Beziehung $e^{r_1 + ir_2} = e^{r_1} (\cos(r_2) + i \sin(r_2))$. Falls dem Leser die komplexe Exponentialfunktion unbekannt ist, kann er diese Gleichung auch einfach als Definition der Exponentialfunktion für komplexe Argumente ansehen. Elementares Nachrechnen liefert dann, dass $\frac{d}{dt} (e^{(r_1 + ir_2)t}) = (r_1 + ir_2)e^{(r_1 + ir_2)t}$ gilt, also auch die komplexe Exponentialfunktion die übliche Differentialgleichung löst. Zurück zu unseren Eigenwerten $\alpha \pm i\beta$: Die lineare Algebra funktioniert auch im Komplexen und wir erhalten zwei komplexe Eigenvektoren, $\vec{z}_1, \vec{z}_2 \in \mathbb{C}^2$. Da die Matrix A aber nur reelle Einträge hat, die bei komplexer Konjugation erhalten bleiben, erhalten wir durch komplexe Konjugation

$$\vec{0} = \vec{0} = \overline{(A - \lambda_1)\vec{z}_1} = (A - \lambda_2)\vec{z}_1.$$

Mit anderen Worten: $\vec{z}_1$ ist ein Eigenvektor zum Eigenwert λ_2 . Wir können also $\vec{z}_{1,2}$ so wählen, dass mit zwei reellen Vektoren $\vec{v}, \vec{w}$ die folgenden Beziehungen gültig sind:

$$\vec{z}_1 = \vec{v} + i\vec{w} \quad \text{und} \quad \vec{z}_2 = \vec{v} - i\vec{w}.$$

Wir bemerken schon einmal, dass die reellen Vektoren $\vec{v}$ und $\vec{w}$ linear unabhängig sind, andernfalls wären ja bereits $\vec{v}$ und $\vec{w}$ Eigenvektoren. Somit haben wir zwei komplexe Lösungen der Differentialgleichung, nämlich

$$\begin{aligned} \vec{z}_1(t) &= e^{\alpha t} (\cos(\beta t) + i \sin(\beta t)) (\vec{v} + i\vec{w}) \quad \text{und} \\ \vec{z}_2(t) &= e^{\alpha t} (\cos(\beta t) - i \sin(\beta t)) (\vec{v} - i\vec{w}). \end{aligned}$$

Dabei haben wir ausgenutzt, dass die Kosinusfunktion gerade und die Sinusfunktion ungerade ist. Aus diesen beiden komplexen Lösungen lassen sich durch lineare Überlagerung aber nun auch zwei reelle Lösungen gewinnen:

$$\begin{aligned}\vec{x}_1(t) &= \frac{1}{2} \left(\vec{z}_1(t) + \vec{z}_2(t) \right) = e^{\alpha t} (\cos(\beta t) \vec{v} - \sin(\beta t) \vec{w}), \\ \vec{x}_2(t) &= \frac{1}{2i} \left(\vec{z}_1(t) - \vec{z}_2(t) \right) = e^{\alpha t} (\cos(\beta t) \vec{w} + \sin(\beta t) \vec{v}).\end{aligned}$$

Auch diese beiden Lösungen sind linear unabhängig, denn wir haben $\vec{x}_1(0) = \vec{v}$ und $\vec{x}_2(0) = \vec{w}$. Damit bilden $\vec{x}_1$ und $\vec{x}_2$ ein Fundamentalsystem.

6.4.3 Die Phasenporträts der linearen Differentialgleichungssysteme für $n = 2$

Wir wollen uns nun mit Hilfe der Fundamentallösung erschließen, wie die Phasenporträts linearer Differentialgleichungssysteme in zwei Dimensionen qualitativ aussehen. Natürlich könnten wir auch mit Hilfe der expliziten Lösungen einige Lösungen plotten lassen und daraus induktiv schließen, wie alle Phasenporträts auszusehen haben. Man kann sich dann aber nie sicher sein, dass man nicht eventuell einen wichtigen Sonderfall übersehen hat. Außerdem handelt es sich um eine schöne Übung im Argumentieren.

Grundsätzlich gibt es zu jeder Lösung $\vec{x}$ von (6.25) zwei reelle Koeffizienten c_1, c_2 , sodass

$$\vec{x}(t) = c_1 \vec{x}_1(t) + c_2 \vec{x}_2(t) \quad \text{für alle } t \in \mathbb{R}. \quad (6.30)$$

Wir müssen nun die Eigenwertkonfigurationen einzeln durchgehen. Dabei lagern wir den Fall von verschwindenden Eigenwerten, also einer Matrix A , die keinen vollen Rang hat, in Aufg. 6.3 aus. Der Fall singulärer Ableitungsmatrizen wird in unseren späteren Modellierungen keine Rolle spielen.

Reelle Eigenwerte mit gleichem Vorzeichen

Zwei verschiedene Eigenwerte Wir nehmen zunächst an, dass die beiden Eigenwerte verschieden sind: Nach den bisherigen Ergebnissen existieren Eigenvektoren $\vec{v}_1$ und $\vec{v}_2$ zu den Eigenwerten λ_1 bzw. λ_2 . Wir nehmen zunächst an, dass beide Eigenwerte negativ sind. Wir bauen das Phasenporträt schrittweise auf und starten mit einer Lösung der Form $\vec{x}(t) = \alpha_1 \vec{v}_1 e^{\lambda_1 t}$ und $\alpha_1 > 0$.

Diese Lösung „läuft also die Ursprungshalbgerade $\{r \vec{v}_1 : r > 0\}$ entlang“, kommt für $t \rightarrow -\infty$ aus dem Unendlichen und konvergiert für $t \rightarrow \infty$ gegen den Ursprung. Entsprechendes gilt für Lösungen der Form $\vec{x}(t) = \alpha_1 \vec{v}_1 e^{\lambda_1 t}$ mit $\alpha_1 < 0$ und $\vec{x}(t) = \alpha_2 \vec{v}_2 e^{\lambda_2 t}$

mit $\lambda_2 \neq 0$. Wir haben also jeweils drei Trajektorien in den Eigenräumen der Matrix, der Ursprung selber ist ja auch eine punktförmige Trajektorie.

Haben wir nun eine Lösung, die nicht in den Eigenräumen verläuft ($\alpha_1 \neq 0$ und $\alpha_2 \neq 0$), so konvergiert die Lösung in jedem Fall für $t \rightarrow \infty$ gegen den Ursprung. Dabei schmiegt sich die Trajektorie dem durch $\vec{v}_1$ aufgespannten Eigenraum an, da diese Komponente der Lösung exponentiell langsamer abfällt als die Komponente der Lösung im durch $\vec{v}_2$ aufgespannten Eigenraum:

$$\frac{|\alpha_1 \vec{v}_1| e^{\lambda_1 t}}{|\alpha_2 \vec{v}_2| e^{\lambda_2 t}} = \frac{|\alpha_1 \vec{v}_1|}{|\alpha_2 \vec{v}_2|} e^{(\lambda_1 - \lambda_2)t} \quad (6.31)$$

wobei $\lambda_1 - \lambda_2 > 0$ (siehe Abb. 6.11a). Der stationäre Punkt $(0, 0)$ ist also asymptotisch stabil, man nennt ihn in diesem Fall einen **stabilen Knoten**.

Falls nun beide Eigenwerte positiv und verschieden sind, bleibt das Bild der Trajektorien bestehen, sie werden nur in der entgegengesetzten Richtung durchlaufen: Für $t \rightarrow -\infty$ kommen die Lösungen aus dem Ursprung und laufen für $t \rightarrow \infty$ gegen unendlich. Der Ursprung ist in diesem Fall instabil und wir sprechen von einem **instabilen Knoten**.

Ein doppelter Eigenwert $\lambda_0 \neq 0$ Wir diskutieren den Fall, dass λ_0 negativ ist. Wir unterscheiden zwei Fälle: Im ersten Fall sei A das λ_0 -Fache der identischen Abbildung und damit jeder Vektor ein Eigenvektor. Ein Anfangswert (u_0, v_0) entwickelt sich also einfach entlang des durch den Vektor (u_0, v_0) aufgespannten Eigenraums für $t \rightarrow \infty$ gegen den Ursprung (siehe Abb. 6.11b). Auch in diesem Fall ist der Ursprung asymptotisch stabil und man nennt ihn einen stabilen Knoten.

Im zweiten Fall ist der Eigenraum des doppelten Eigenwerts lediglich eindimensional und wir haben stattdessen noch einen Hauptvektor $\vec{w}$. Die Lösungen haben die Form

$$\vec{x}(t) = e^{\lambda_0 t} (\alpha_1 \vec{v} + \alpha_2 (\vec{w} + t\vec{v})). \quad (6.32)$$

Eine Lösung mit $\alpha_2 = 0$ läuft also wieder durch den durch $\vec{v}$ aufgespannten Eigenraum. Eine Lösung mit $\alpha_1 = 0$ dagegen befindet sich zur Zeit $t = 0$ im durch $\vec{w}$ aufgespannten Raum. Mit wachsendem t bekommt aber der $\vec{v}$ -Raum einen immer größeren Anteil, sodass die Lösung in die $\vec{v}$ -Richtung „einbiegt“ (siehe Abb. 6.11c). Auch in diesem Fall ist der Ursprung asymptotisch stabil und wir nennen ihn einen stabilen Knoten.

Für einen positiven doppelten Eigenwert werden die Trajektorien in der entgegengesetzten Richtung durchlaufen, der Ursprung ist instabil und wir sprechen von einem instabilen Knoten.

Eigenwerte mit unterschiedlichem Vorzeichen: der Sattelpunkt

Lösungen haben in diesem Fall die Form

$$\vec{x}(t) = \alpha_1 \vec{v}_1 e^{\lambda_1 t} + \alpha_2 \vec{v}_2 e^{\lambda_2 t} \quad \text{mit} \quad \alpha_1 < 0 < \alpha_2.$$

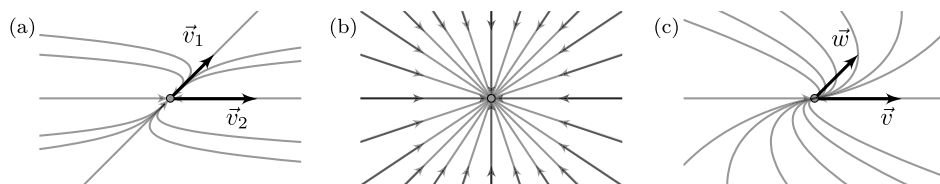


Abb. 6.11 Stabile Knoten: (a) mit zwei unterschiedlichen Eigenwerten, (b) mit einem doppelten Eigenwert mit geometrischer Vielfachheit 2; (c) mit einem doppelten Eigenwert mit geometrischer Vielfachheit 1

Falls $\alpha_2 = 0$ ist, laufen die Lösungen im $\vec{v}_1$ -Raum von unendlich kommend in den Ursprung hinein. Falls hingegen $\alpha_1 = 0$ ist, laufen die Lösungen im $\vec{v}_2$ -Raum aus dem Ursprung heraus nach unendlich. Falls $\alpha_1 \neq 0$ und $\alpha_2 \neq 0$ laufen die Lösungen für $t \rightarrow -\infty$ asymptotisch aus dem $\vec{v}_1$ -Raum heraus und laufen für $t \rightarrow \infty$ asymptotisch in den $\vec{v}_2$ -Raum (siehe Abb. 6.12). In diesem Fall ist der Ursprung instabil: Lösungen können ihm beliebig nahekommen und laufen sogar nach unendlich. Wir sprechen von einem **Sattelpunkt**. Die beiden Trajektorien, die längs des $\vec{v}_1$ -Raums in den Ursprung hineinlaufen, separieren das Phasenportrait in zwei Hälften: In der einen Hälfte nähern sich die Lösungen im Unendlichen der einen $\vec{v}_2$ -Halbgerade, in der anderen Hälfte der anderen Halbgerade. Die beiden in den stationären Punkt hineinlaufenden Trajektorien nennt man die *Separatrixen*.

Komplexe Eigenwerte: das Zentrum und der Strudel

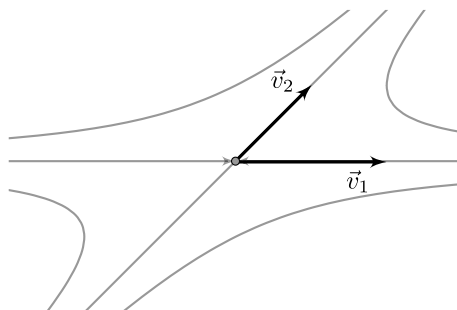
Wir unterscheiden, ob die beiden komplex konjugierten Eigenwerte einen nichtverschwindenden Realteil haben oder ob sie rein imaginär sind.

Rein imaginäre Eigenwerte: das Zentrum Die Lösungen

$$\vec{x}(t) = \alpha_1 (\cos(\beta t)\vec{v} - \sin(\beta t)\vec{w}) + \alpha_2 (\cos(\beta t)\vec{w} + \sin(\beta t)\vec{v})$$

Abb. 6.12 Der Sattelpunkt:

Die in den Ursprung hineinlaufenden Trajektorien (Separatrixen) unterteilen das Phasenportrait in zwei Hälften mit unterschiedlichem Langzeitverhalten



sind periodisch mit der Periode $P = 2\pi/\beta$, durchlaufen also geschlossene Zyklen. Man kann sich überlegen, dass die entstehenden Trajektorien tatsächlich Ellipsen sind. Der Ursprung ist in diesem Fall stabil (Lösungen bleiben nah, wenn sie einmal hinreichend nah waren), aber nicht asymptotisch stabil (die Lösungen konvergieren nicht gegen den Ursprung). Der Ursprung wird in diesem Fall als *Zentrum* bezeichnet (siehe Abb. 6.13).

Eigenwerte mit nicht-verschwindendem Realteil: der Strudel In dieser Situation hat die Lösung die Form

$$\vec{x}(t) = e^{\alpha t} (\alpha_1 (\cos(\beta t)\vec{v} - \sin(\beta t)\vec{w}) + \alpha_2 (\cos(\beta t)\vec{w} + \sin(\beta t)\vec{v}))$$

mit $\alpha \neq 0$. Die Ellipsenbahnen aus dem Zentrumsfall haben nun einen globalen Faktor, der mit zunehmender Zeit exponentiell schrumpft, falls $\alpha < 0$. Die Trajektorie strudelt gewissermaßen für $t \rightarrow \infty$ in den Ursprung. Dabei umrundet sie den Ursprung unendlich oft. Der Ursprung ist asymptotisch stabil und wir sprechen von einem **stabilen Strudel**.

Im Fall komplexer Eigenwerte mit positivem Realteil, $\alpha > 0$, drehen sich die Trajektorien für $t \rightarrow \infty$ gegen unendlich. In diesem Fall sprechen wir von einem **instabilen Strudel**. Man beachte, dass in Abb. 6.13 nur zwei Trajektorien angedeutet sind: der stationäre Punkt und eine strudelnde Trajektorie.

Im nächsten Abschnitt werden wir einen Satz erläutern, der das Verhältnis der Lösungen des nichtlinearen Systems und der Lösungen des um einen stationären Punkt linearisierten Systems in den Blick nimmt.

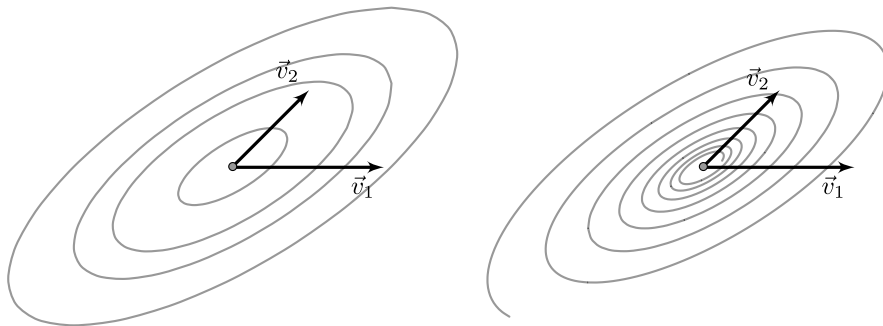


Abb. 6.13 Das Zentrum und der Strudel

6.4.4 Der Linearisierungssatz von Grobman und Hartman

Es gibt eine ganze Menge von Sätzen, die das Verhältnis von Lösungen des nichtlinearen Systems und der Lösungen des um einen stationären Punkt linearisierten Systems in den Blick nehmen. Wir werden hauptsächlich auf den Satz von Grobman-Hartman zurückgreifen, der nahezu zeitgleich und unabhängig voneinander von dem sowjetischen Mathematiker D. M. Grobman in [35] und dem amerikanischen Mathematiker P. Hartman in [38] bewiesen wurde. Sehr umgangssprachlich drückt der Satz aus, dass das Phasenporträt in der Umgebung eines *hyperbolischen* stationären Punktes qualitativ genauso aussieht wie das Phasenporträt der Linearisierung um den Ursprung. Dabei heißt ein stationärer Punkt hyperbolisch, wenn alle Eigenwerte der Ableitungsmatrix an dieser Stelle einen nichtverschwindenden Realteil haben. Keine Aussage macht der Satz also, wenn es verschwindende oder wenn es rein komplexe Eigenwerte gibt.

Anschaulich lässt sich das so vorstellen: Um vom linearen auf das nichtlineare System zu kommen, kann man die Trajektorien verbiegen und verformen, aber wenn eine Trajektorie im linearen System in den stationären Punkt reinläuft, so tut das auch ihr Gegenpart im nichtlinearen System. Wenn sie dabei den stationären Punkt unendlich oft umdreht (Strudel), so passiert das auch im nichtlinearen System. Hat aber z. B. die Ableitungsmatrix eines stationären Punktes P^* einen positiven und einen negativen Eigenwert, das linearisierte System also eine Sattelpunktstruktur, dann gibt es auch im nichtlinearen System genau zwei Trajektorien, die in P^* hineinlaufen und zwei Trajektorien, die aus P^* herauslaufen. P^* sieht also auch aus wie ein Sattelpunkt (siehe Abb. 6.14).

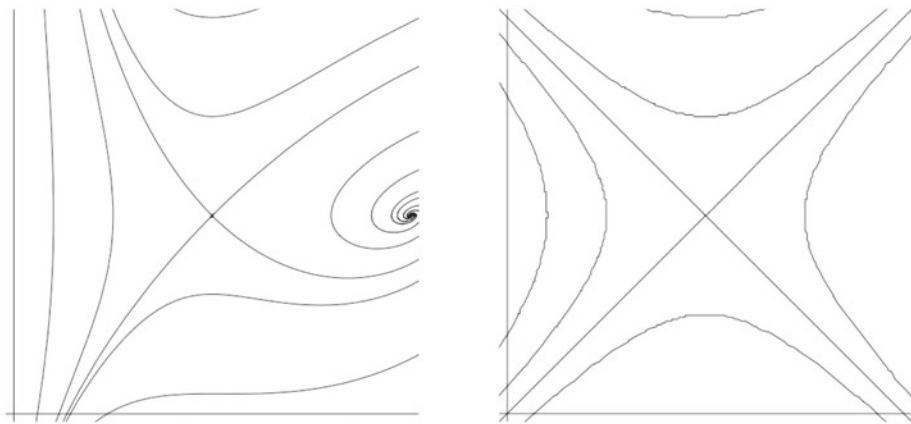


Abb. 6.14 Phasenporträt eines nichtlinearen Systems in der Umgebung eines hyperbolischen Punktes x^* im Vergleich mit der Linearisierung. (Abbildung aus [87], ©(2012) American Mathematical Society)

Um den Satz mathematisch präziser zu formulieren, müssen wir den Fluss eines dynamischen Systems $\dot{x} = F(x)$ einführen. Das machen wir für den allgemeinen Fall, die Dimension n des Systems kann also beliebig sein. Wir nehmen dazu an, dass die Abbildung $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ stetig differenzierbar ist. Es würde auch wieder reichen, dass die Abbildung auf irgendeiner offenen Menge $G \subset \mathbb{R}^n$ definiert ist, dann würde sich alles auf diese Menge G beziehen. Nach dem Satz von Picard-Lindelöf gibt es zu jedem Punkt $x \in \mathbb{R}^n$ eine eindeutige Lösung des Anfangswertproblems $\dot{\tilde{x}}(t) = f(\tilde{x}(t))$ und $\tilde{x}(0) = x$, die auf einem maximalen Existenzintervall I_x lebt. (Wir ändern hier unsere Bezeichnungsweisen, damit der zu definierende Fluss von x und nicht von x_0 abhängt.) Für globale Lösungen gilt $I_x = \mathbb{R}$, wenn es aber Blow-up-Phänomene, gibt kann I_x auch eine echte Teilmenge von $\mathbb{R}$ sein. Wir definieren dann

$$\Phi(t, x) := \tilde{x}(t)$$

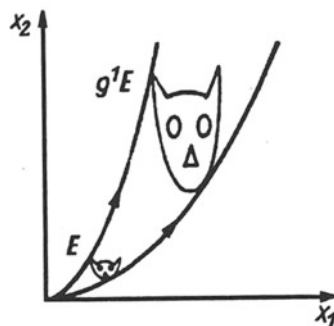
als den Fluss von F . Halten wir ein x fest und variieren t , so entsteht gerade die Trajektorie der Lösung des Differentialgleichungssystems, die sich zur Zeit $t = 0$ in x befindet. Betrachten wir dagegen eine Konstellation von Punkten zur Zeit $t = 0$ und dann das Bild dieser Konstellation zu einer späteren Zeit $\hat{t}$, sehen wir, wie sich diese Konstellationen im Intervall $[0, \hat{t}]$ geändert hat (siehe Abb. 6.15). Jetzt können wir den Satz von Grobman-Hartman wiedergeben. Dazu erinnern wir, dass ein Homöomorphismus eine stetige bijektive Abbildung mit stetiger Umkehrabbildung ist.

Satz 6.2 (Der Satz von Grobman-Hartman)

Es sei F ein stetig differenzierbares Vektorfeld, x^* ein stationärer hyperbolischer Punkt von F und $A = J_F(x^*)$ die zugehörige Ableitungsmatrix. Φ_F sei der Fluss von F und Φ_A der Fluss von A . Dann gibt es eine Umgebung U um den Ursprung, eine Umgebung V um x^* , ein offenes Intervall I um 0 sowie einen Homöomorphismus $\varphi : U \rightarrow V$, so dass

$$\varphi(\Phi_A(t, x)) = \Phi_F(t, \varphi(x)) \quad \text{für alle } x \in U, t \in I.$$

Abb. 6.15 Veränderung einer Punktekongstellation zwischen den Zeiten $t = 0$ und $t = \hat{t}$ durch einen Fluss Φ . (Abbildung aus [3], ©(1980), Springer)



Man sagt auch, dass die beiden Flüsse „lokal topologisch konjugiert“ zueinander sind. Ein Beweis dieses Satzes sprengt leider den Rahmen dieses Buches und wir verweisen z. B. auf [87].

6.4.5 Fortsetzung: Räuber-Beute-Systeme mit logistischem Wachstum

Wir kehren nun zu dem andiskutierten Modell aus Abschn. 6.4.1 zurück, siehe (6.21),

$$\begin{aligned}\dot{u} &= u(1 - bu) - uv \\ \dot{v} &= av(-1 + u),\end{aligned}$$

und erinnern an die drei stationären Punkte $P_0^* = (0, 0)$, $P_1^* = (1/b, 0)$ und $P_2^* = (1, 1-b)$. Wir bestimmen nun den Charakter dieser drei Punkte mit Hilfe des Satzes von Grobman-Hartman. Dazu berechnen wir als Erstes die Ableitungsmatrix

$$J_F(u, v) = \begin{pmatrix} 1 - 2bu - v & -u \\ av & a(-1 + u) \end{pmatrix}.$$

Wir starten mit dem Punkt P_0^* . Wir wissen bereits, dass aus diesem Punkt eine Trajektorie hinausgeht, nämlich in der u -Achse, und dass eine Trajektorie hineinläuft, nämlich in der v -Achse. Das deutet auf einen Sattelpunkt hin. Tatsächlich berechnen wir:

$$J_F(P_0^*) = \begin{pmatrix} 1 & 0 \\ 0 & -a \end{pmatrix}$$

Diese Matrix hat die beiden Eigenwerte $1 > 0$ und $-a > 0$. Nach dem Satz von Grobman-Hartman ist also P_0^* tatsächlich ein Sattelpunkt. Für P_1^* berechnen wir:

$$J_F(P_1^*) = \begin{pmatrix} -1 & -1/b \\ 0 & a(-1 + 1/b) \end{pmatrix}$$

Diese Matrix hat die Eigenwerte $\lambda_1 = -1$ und $\lambda_2 = a(-1 + 1/b)$. Der zweite Eigenwert ist genau dann positiv, wenn $0 < b < 1$ gilt. In diesem Fall ist P_1^* ein Sattelpunkt. Zwei negative Eigenwerte liegen genau dann vor, wenn $b > 1$ gilt. Dann handelt es sich bei P_1^* um einen stabilen Knoten. Falls $b = 1$ zutrifft, ist P_1^* nicht hyperbolisch und der Satz von Grobman-Hartman liefert keine Aussage. Nun zu P_2^* :

$$A := J_F(P_2^*) = \begin{pmatrix} -b & -1 \\ a(1 - b) & 0 \end{pmatrix}$$

Auch für spätere Zwecke und die Übungsaufgaben legen wir uns ein paar Rechentricks zurecht.

Determinante und Spur der Jacobi-Matrix und die Stabilität des stationären Punkts

Nach der p - q -Formel gilt für die beiden Eigenwerte $\lambda_{1/2}$ der Matrix A

$$\lambda_{1/2} = -\frac{\operatorname{tr} A}{2} \pm \frac{1}{2} \sqrt{\operatorname{tr}^2 A - 4 \det A}.$$

Aus dieser Darstellung erkennen wir Folgendes:

- Ist $\det A < 0$, so haben die beiden Eigenwerte unterschiedliche Vorzeichen. Folglich handelt es sich um einen Sattelpunkt.
- Ist $\det A > 0$ und $\operatorname{tr} A > 0$, so ist der stationäre Punkt instabil, und zwar ein instabiler Knoten, falls zusätzlich $\operatorname{tr}^2 A - 4 \det A \geq 0$, bzw. ein instabiler Strudel, falls zusätzlich $\operatorname{tr}^2 A - 4 \det A < 0$ gilt.
- Ist $\det A > 0$ und $\operatorname{tr} A < 0$, so ist der stationäre Punkt stabil, und zwar ein stabiler Knoten, falls zusätzlich $\operatorname{tr}^2 A - 4 \det A \geq 0$, bzw. ein stabiler Strudel, falls zusätzlich $\operatorname{tr}^2 A - 4 \det A < 0$ gilt.

In vielen Anwendungen ist es vor allem wichtig, ob stationäre Punkte stabil oder instabil sind, und in letzterem Fall, ob es sich um einen Sattelpunkt handelt. Die Unterscheidung in (in-)stabile Knoten und (in-)stabile Strudel ist häufig weniger relevant. Wenn wir nicht zwischen Strudeln und Knoten unterscheiden, haben wir es in unserem Fall leichter: $\operatorname{tr} A = -b$ und $\det A = a(1 - b)$. Damit ist P_2^* ein Sattelpunkt, wenn $b > 1$, und ein stabiler Knoten/Strudel, wenn $0 < b < 1$ gilt. Wollten wir auch noch zwischen Knoten und Strudeln unterscheiden, müssten wir für $0 < b < 1$ das Vorzeichen der Diskriminante $D(a, b) := \operatorname{tr}^2 A - 4 \det A = b^2 + 4a(1 - b)$ untersuchen und erhalten unter den generellen Bedingungen $a > 0$ und $0 < b < 1$, dass $D(a, b) \geq 0$ genau dann gilt, wenn $0 < a \leq \frac{b^2}{4(1-b)}$. In diesem Fall handelt es sich dann um einen stabilen Knoten, im Fall $D(a, b) < 0$, also $a > b^2/(4(1 - b))$, um einen stabilen Strudel.

Wir können die in Abschn. 6.4.1 begonnenen Phasenporträts (siehe Abb. 6.10) jetzt vervollständigen. Wir beginnen mit dem Fall $0 < b < 1$ und beschränken uns hier zunächst auf den Fall $D(a, b) \geq 0$, P_2^* ist also ein stabiler Knoten. Es erweist sich als günstig, mit den Trajektorien zu beginnen, die aus den Sattelpunkte starten, bzw. in die Sattelpunkte münden. Bezüglich P_0^* sind das die beiden bekannten Trajektorien in den Achsen sowie Trajektorien, die in die nicht interessierenden Quadranten II–IV entschwinden. Die in den Sattelpunkt P_1^* hineinlaufenden Trajektorien sind ebenfalls schon bekannt, es sind die Trajektorien in der u -Achse. Genau eine Trajektorie muss diesen Punkt in die obere Halbebene verlassen und sich nach oben links entwickeln, bis sie waagrecht die v -Nullkline $u = 1$ schneidet. Nach endlich vielen Umrundungen des stabilen Knotens P_2^* muss sie irgendwann in diesen hineinlaufen. Das kann aufgrund der vorgeschriebenen Richtungen nur aus dem Bereich I oder dem Bereich III geschehen (siehe Abb. 6.16a). Tatsächlich wäre es auch möglich, dass die Separatrix aus dem Punkt P_1^* direkt, also ohne jede Umrundung in den Punkt

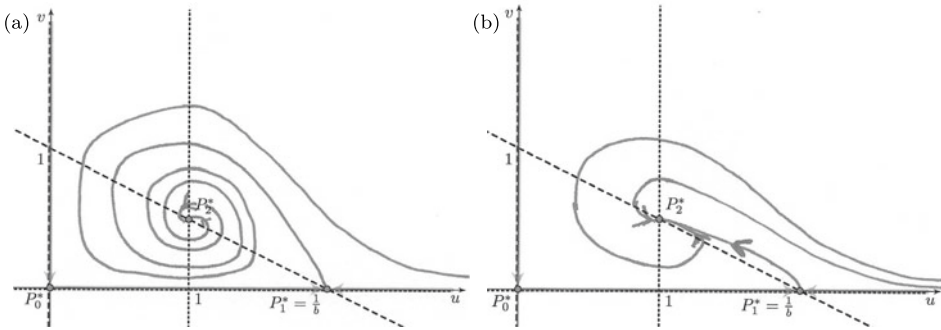


Abb. 6.16 Der Fall $0 < b < 1$ und $D(a, b) \geq 0$: P_2^* ist ein stabiler Knoten. Die beiden Populationen entwickeln sich auf lange Sicht in den Koexistenzzustand P_2^*

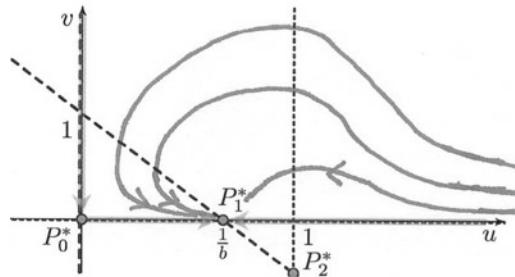
P_2^* hineinläuft (siehe Abb. 6.16). Falls nun $D(a, b) < 0$ wäre, wäre P_2^* ein stabiler Strudel und alle Trajektorien müssten den Punkt P_2^* unendlich oft umrunden. Wir verzichten auf ein Phasenporträt.

Im Fall $b > 1$ ist P_1^* ein stabiler Knoten und P_2^* ein Sattelpunkt, der allerdings nicht im ersten Quadranten liegt und uns deshalb auch nicht interessiert. Jetzt laufen alle Trajektorien in den stabilen Knoten P_1^* hinein (siehe Abb. 6.17).

6.4.6 Diskussion des Ergebnisses

Gemäß diesem Modell gibt es zwei Szenarien, der entscheidende Parameter ist $b = \frac{\gamma}{\delta K}$: Falls $0 < b < 1$, wird auf lange Sicht der Koexistenzzustand P_2^* angenommen. Wir würden also stabile Populationsgrößen erwarten, die nur durch zufällige äußere Einflüsse in zufällige Richtungen etwas schwanken. Falls $b > 1$, würden wir auf lange Sicht keine Räuberpopulation mehr sehen können, diese wäre ja ausgestorben. Im Vergleich zum Lotka-Volterra-System gibt es also keine periodischen Schwankungen mehr. Eine zu hohe intraspezifische Konkurrenz zwischen den Mitgliedern der Beutepopulation, das entspricht einer zu kleinen Kapazität K , führt zu einem Aussterben der Räuberpopulation.

Abb. 6.17 Der Fall $1 < b$, P_1^* ist ein stabiler Knoten. Die Populationen entwickeln sich auf lange Sicht in den Zustand P_1^*



Vergleichen wir diesen Befund nun mit der stabilen zehnjährigen Periodizität der Größe der Luchspopulation aus den Daten der Hudson's Bay Company (siehe Abschn. 6.3.5, Abb. 6.8), sehen wir, dass wir diese Daten auch mit diesem Modell nicht plausibel erklären können. Wir müssen also noch weitere Effekte mit berücksichtigen. Das soll im Abschn. 6.5 geschehen. Zuvor wollen wir uns aber noch einer mathematisch geprägten Frage widmen.

6.4.7 Könnte es im Räuber-Beute-Modell mit logistischem Wachstum auch unbeschränkte Lösungen geben?

Im vorigen Abschnitt haben wir aus der Natur der stationären Punkte und den vorgegebenen Richtungen der Trajektorien ein Phasenporträt so weit skizziert, dass wir das Langzeitverhalten des Systems erkennen konnten. Dabei haben wir an einigen Stellen Schlüsse gezogen, die mathematisch nicht zwingend waren. Zum Beispiel haben wir behauptet, dass die Trajektorie aus dem Sattelpunkt P_1^* entweder direkt in den stationären Punkt P_2^* mündet oder aber in den Bereich IV läuft, vergleiche Abb. 6.16a. Tatsächlich stünde aber auch eine Trajektorie, die in v -Richtung gegen unendlich verläuft, nicht im Widerspruch zu den Richtungen und den Charakteren der stationären Punkte (siehe Abb. 6.18). Im Gegensatz zu der Methode der Lyapunov-Funktion aus Abschn. 6.3, die gewissermaßen einen globalen Charakter hat, liefert der Linearisierungsansatz nur lokale Resultate, die etwas über die Umgebung eines stationären Punkts aussagen. Für die Bereiche dazwischen bleibt Raum für Spekulationen.

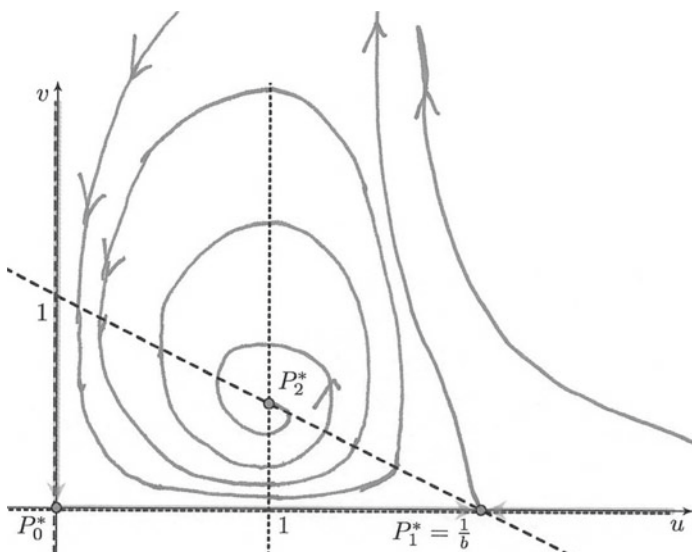


Abb. 6.18 Auch unbeschränkte Trajektorien wären mit den Vorzeichen und den Arten der stationären Punkte verträglich

Wie können wir nun damit umgehen? Die Anwenderin würde wohl einfach für ein paar unterschiedliche Parameterwerte von a und b numerische Berechnungen anstellen, sehen, dass in all diesen Fällen die Szenarien aus Abschn. 6.4.5 eintreten, und daraus mit einem plausiblen, aber nicht deduktiven Schluss folgern, dass das dann wohl immer so sein wird. Eine Mathematikerin wird nach einem mathematischen Argument suchen, warum die Szenarien aus Abb. 6.18 nicht eintreten können. Wir wollen hier den Weg der Mathematikerin beschreiten. Dazu kann es nicht reichen, nur die Vorzeichen der einzelnen Terme von F und die Linearisierung der stationären Punkte zu betrachten. Tatsächlich lassen sich nämlich Funktionen mit derselben Vorzeichensituation und derselben Art von stationären Punkten definieren, in denen die Szenarien aus Abb. 6.18 eintreten. Wir müssen also die konkreten Terme von F ins Spiel bringen: Wir betrachten den Bereich III, der durch die Ungleichungen $u > 1$ und $1 - bu - v < 0$ definiert ist (siehe Abb. 6.19). Solange eine Lösung in diesem Bereich verläuft, gilt für sie $\dot{u} < 0$ und $\dot{v} > 0$. Wir betrachten jetzt eine konkrete Lösung $\tau \mapsto (u(\tau), v(\tau))$ von (6.21) zu einem Anfangsdatum $(u(0), v(0)) = (u_0, v_0)$ aus dem Bereich III; zusätzlich nehmen wir aus technischen Gründen, die später deutlich werden, $v_0 \geq 2$ an. Diese zusätzliche Annahme ist unproblematisch: Für eine Lösung, deren v -Komponente gegen unendlich läuft, wäre diese ja irgendwann größer als 2. Ferner sei $I = [0, \hat{\tau})$ das maximale Zeitintervall, in dem diese Lösung im Bereich III verbleibt. Die Funktion $\tau \mapsto u(\tau)$ ist auf I streng monoton fallend und besitzt also eine Umkehrfunktion $u \mapsto \tau(u)$.

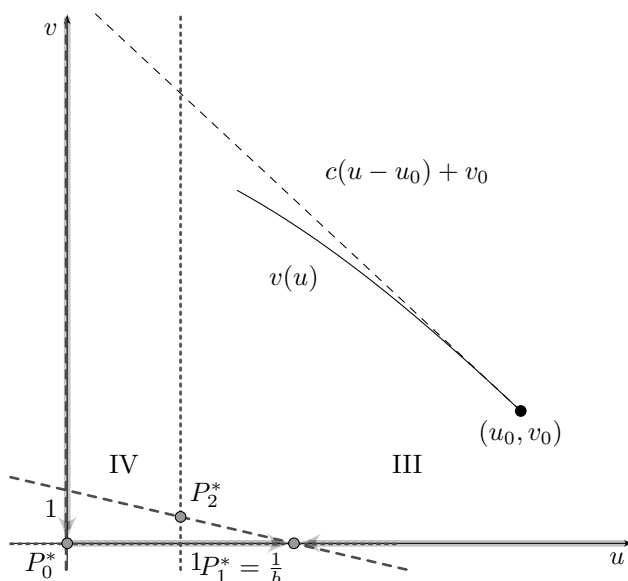


Abb. 6.19 Trajektorien können nicht nach unendlich entkommen

Wir werden im Anschluss an diese Überlegungen noch eine Bemerkung zu den Bezeichnungen machen. Hier stellt sich ja die Frage, ob man z. B. den Buchstaben τ einerseits als Bezeichnung für die unabhängige Variable der Funktion u und andererseits als Funktionsnamen für die Umkehrfunktion von u verwenden darf. Ähnliches gilt für den Buchstaben u . Auch werden wir gleich mit dem Buchstaben v unterschiedliche Funktionen bezeichnen, siehe auch die Bemerkung in Abschn. 1.2.2.

Nach dem Gesagten können wir die Funktion v als Funktion von u darstellen, mit anderen Worten: Wir stellen die Trajektorie als Graph einer Funktion von u dar, $v(\tau) = v(u) = v(\tau(u))$. Für die Ableitung der Funktion $v(u)$ gilt dann:

$$v'(u) = \frac{d}{du} (v(\tau(u))) = \dot{v}(\tau) \tau'(u) = \frac{\dot{v}(\tau)}{\dot{u}(\tau)} = \frac{av(-1+u)}{u(1-bu-v)} \quad (6.33)$$

Wir haben eine Differentialgleichung für die Trajektorie gefunden, die aber leider zu kompliziert ist, um sie explizit zu lösen. Wir können jedoch zeigen, dass die rechte Seite dieser Differentialgleichung beschränkt ist: Dabei können wir $1 < u \leq u_0$ und $v_0 \leq v$ annehmen, weil wir schon wissen, dass die betrachtete Funktion in diesem Bereich verbleibt. Wir schätzen also für solche u, v ab. Um die Probleme mit den Vorzeichen zu umgehen, nehmen wir die Beträge, müssen dann nach oben abschätzen, also immer den Zähler vergrößern und den Nenner verkleinern:

$$\begin{aligned} |v'(u)| &= \frac{|av(-1+u)|}{|u(1-bu-v)|} \leq \frac{av(-1+u)}{u(bu+v-1)} \leq \frac{avu_0}{1 \cdot (b \cdot 1 + v - 1)} \\ &= \frac{v}{v} \cdot \frac{au_0}{1 + \frac{b-1}{v}} \leq \frac{au_0}{1 + \frac{b-1}{2}} \leq 2au_0 =: c \end{aligned}$$

(An welcher Stelle haben wir $v_0 \geq 2$ benutzt?) Damit erhalten wir aber für $1 < u < u_0$

$$v(u) = v_0 + \int_{u_0}^u v'(p) dp \leq v_0 + c(u_0 - u),$$

die Lösung muss also beschränkt bleiben und in den Bereich IV übergehen. Wir schließen diesen Abschnitt mit drei Bemerkungen:

Auch eine Abschätzung der Form $|v'(u)| \leq cv$ wäre hinreichend gewesen, um zu zeigen, dass die Trajektorie nicht gegen unendlich gehen kann. In diesem Fall würde man v gegen die Lösung der Differentialgleichung $w' = -cw$, also gegen eine Exponentialfunktion abschätzen, die ebenfalls keinen Blow-up aufweist. Um eine unbeschränkte Trajektorie nachweisen zu können, wäre eine Abschätzung der Form $|v'(u)| \geq cv^\alpha$ mit einem $\alpha > 1$ geeignet: Nach Abschn. 5.5.2 entwickeln nämlich Lösungen dieser Differentialgleichungen einen Blow-up.

Unser Umgang mit Bezeichnungen ist erläuterungsbedürftig. In der Schule und in den Einführungsvorlesungen der Universität lernt man, dass eine Funktion eine eindeutige Zuordnung von Elementen eines Definitionsbereichs auf Elemente eines Wertebereichs ist (oder eine linkstotale rechtseindeutige Relation), die mit einem beliebigen bedeutungslosen

Namen versehen werden darf, der dann aber auch gerade für diese Funktion reserviert ist und nicht anderweitig Verwendung finden darf. Lässt sich eine Funktionsvorschrift in Form eines Terms angeben, so kann man der unabhängigen Variablen ebenfalls einen beliebigen bedeutungslosen Namen geben, der nur nicht schon anderweitig benutzt wird. Unterschiedliche Funktionen hören aber prinzipiell auf unterschiedliche Namen. In Modellierungskontexten ist die Situation insofern anders, als dass eine semantische Ebene ins Spiel kommt. Die Bezeichnungen u , v und τ stehen für die drei Größen „Größe der Beutepopulation“, „Größer der Räuberpopulation“ und „Zeit“, die in einem funktionalen Zusammenhang stehen, in dem mehr oder weniger jede der drei Größen in Abhängigkeit von einer der verbleibenden Größen betrachtet werden kann. Man wird aber immer dieselben Namen verwenden, $u(\tau)$, $u(v)$, $v(u)$, $\tau(v)$ und so weiter. Da aus dem Zusammenhang in der Regel klar ist, worüber man spricht, kommt es selten zu Verwechslungen und Missverständnissen. Wie sperrig (aber vielleicht für die eine oder andere Leserin auch klarer) wäre die Herleitung der Formel (6.33), wenn man die semantische Ebene nicht benutzen und stattdessen rigoros unterschiedliche Funktionen mit unterschiedlichen Namen versehen würde. Zur Illustration wollen wir das hier einmal durchführen:

Die Funktion $u : I \rightarrow (1, u_0]$ ist nach dem oben Gesagten streng monoton fallend, den Wertebereich bezeichnen wir mit $(\hat{u}, u_0]$. Die Funktion $v : I \rightarrow [v_0, \infty)$ ist streng monoton wachsend. Diese beiden Funktionen lösen das Differentialgleichungssystem (6.21). Wir führen die Umkehrfunktion $\tilde{\tau} = u^{-1} : (\hat{u}, u_0] \rightarrow I$ und die Funktion $\tilde{v} = v \circ \tilde{\tau} : (\hat{u}, u_0] \rightarrow [v_0, \infty)$ ein. Der Trajektorienabschnitt $\{(u(\tau), v(\tau)) : \tau \in I\}$ ist dann gerade der Graph der Funktion $\tilde{v}$. Mit der Kettenregel können wir nachrechnen:

$$\begin{aligned} \tilde{v}'(\tilde{u}) &= \dot{v}(\tilde{\tau}(\tilde{u}))\tilde{\tau}'(\tilde{u}) = \dot{v}(\tilde{\tau}(\tilde{u})) \cdot \frac{1}{\dot{u}(\tilde{\tau}(\tilde{u}))} \\ &= \frac{av(\tilde{\tau}(\tilde{u})) \cdot (-1 + u(\tilde{\tau}(\tilde{u})))}{u(\tilde{\tau}(\tilde{u})) \cdot (1 - bu(\tilde{\tau}(\tilde{u})) - v(\tilde{\tau}(\tilde{u})))} \\ &= \frac{a\tilde{v}(\tilde{u}) \cdot (-1 + \tilde{u})}{\tilde{u}(1 - b\tilde{u} - \tilde{v}(\tilde{u}))} \end{aligned}$$

Jede Leserin mag selber entscheiden, welche Darstellung ihr besser gefällt.

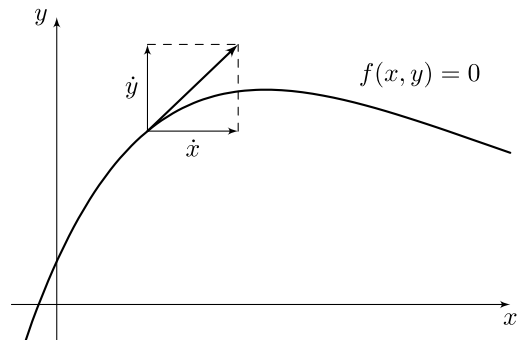
Mit der Formel

$$v'(u) = \frac{\dot{v}}{\dot{u}}$$

sind wir nah an den Ideen der *Fluxionsrechnung*, mit der Isaac Newton 1666 einen der historisch frühesten Zugänge zur Differentialrechnung eröffnete. Wir zitieren Sonar [85, S. 378]; die Variablen x , y entsprechen dabei unseren Variablen u und v , die Kurve $f(x, y) = 0$ ist unsere Trajektorie.

„Da die Kurve $f(x, y) = 0$ durch eine „fließende“ Bewegung zustande kommt, nennt Newton die Größen x und y die Fluents und $\dot{x}$ und $\dot{y}$ „Fluxionen“. Aus den Fluxionen zur Zeit t folgt sofort die Tangentensteigung an der Kurve (vgl. Abb. 6.20) zu

Abb. 6.20 Fluxionen
(Geschwindigkeitskomponenten) bei Bewegungen längs einer Kurve. (Abbildung nach Sonar, [85])



$$\frac{\dot{y}}{\dot{x}} = \frac{\frac{dy}{dt}}{\frac{dx}{dt}} = \frac{dy}{dx}.$$

6.4.8 Reflexionen über das mathematische Modellieren V: Die Art des Schließens in der mathematischen Untersuchung

Wir haben uns im letzten Abschnitt große Mühe gegeben, mathematisch zu beweisen, dass es in dem Modell (6.21) keine unbeschränkten Lösungen gibt. Manche Leserin hat sich vielleicht die Frage gestellt, ob das denn wirklich nötig sei. Man könne ja schließlich auch einfach für ein paar vernünftig ausgewählte Parameter- und Startwerte Lösungen ausrechnen, und wenn die nicht gegen unendlich laufen, warum sollten es dann irgendwelche anderen machen?

Ein solches Vorgehen genügt natürlich nicht den Ansprüchen eines mathematischen Beweises, der aus einer lückenlosen Folge korrekter Deduktionen zu bestehen hat. Wir nennen diese Vorgehensweise das *deduktive Schließen*. Tatsächlich ist man in der angewandten Mathematik auch häufig mit weniger restriktiven Schlüssen zufrieden, z. B. der Art, dass der behauptete Sachverhalt für vernünftig ausgewählte Beispiele zutrifft. Dabei ist auch schon der Ausdruck „vernünftig“ in der Regel nicht mathematisch definierbar. Wir sprechen von *plausiblen* Schlüssen.

Warum ist ein solches „vernünftiges“, aber nicht deduktives Vorgehen in der angewandten Mathematik hin und wieder nötig und zulässig? Was hätten wir denn gemacht, wenn uns der Beweis nicht gelungen wäre? Wären wir an dem System von Differentialgleichungen und der Fragestellung nach unbeschränkten Lösungen aus rein mathematischer Sicht interessiert, so hätten wir es wohl zunächst mit einem einfacheren System probiert und wären später vielleicht wieder zu dem ursprünglichen System zurückgekehrt. Aus Sicht des Anwenders haben wir aber diese Freiheit nicht. Wir benötigen Aussagen genau über dieses System, das da steht. Verhält sich nun dieses System in allen sinnvollen numerischen Tests

vernünftig und liefert gute Prognosen für das Verhalten des realen Systems, so wird der fehlende deduktive Beweis an der einen Stelle die Anwenderin kaum stören. Letztere ist ja im Wesentlichen an einem Verständnis der Realsituation und an Prognosen des zukünftigen Verhaltens interessiert.

Tatsächlich liegt z. B. für die grundlegende Gleichung der Hydrodynamik, die Navier-Stokes-Gleichung, bis heute noch kein Beweis für die globale Existenz starker Lösungen für allgemeine Anfangsdaten vor. Das hindert die Anwender aber nicht daran, mit solchen Lösungen zu rechnen.

6.5 Räuber-Beute-Systeme mit Sättigungs- und Schwelleneffekten – stabile Grenzzyklen und der Satz von Poincaré-Bendixson

6.5.1 Die Modellbildung

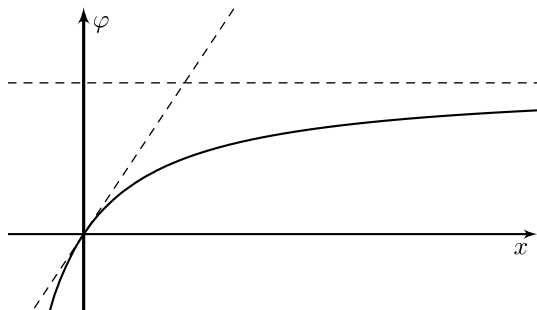
Wir sind nach wie vor auf der Suche nach einem Räuber-Beute-Modell, das die stabile Periodizität aufweist, wie sie sich in den Daten der Luchsbestände der Hudson's Bay Company findet (vergleiche Abb. 6.8). Wir nehmen jetzt den Interaktionsterm unter die Lupe, der die Interaktion zwischen der Beute- und der Räuberpopulation beschreibt. Unser bisheriges Modell beschreibt diese Interaktion als proportional zum Produkt der beiden Populationsgrößen. Aus Sicht eines Räubers bedeutet eine doppelt so große Beutepopulation einen doppelt so großen Benefit. Das ist sicherlich eine plausible Annahme, solange die Beutepopulation klein ist. Ist jedoch die Beutepopulation bereits sehr groß, sodass der Räuber ohne Probleme jederzeit auf ein Beutetier treffen, sich also zu jeder Zeit satt fressen kann, wird eine verdoppelte Beutepopulationsgröße keinen doppelten Benefit mehr hervorrufen, mehr als satt geht eben nicht.

Wir wollen also in der Räubergleichung $\dot{x} = -\gamma y + \delta xy$ (zur Modellbildung nehmen wir natürlich wieder die dimensionsbehafteten Gleichungen) den Interaktionsterm $+\delta xy$ durch einen Term der Bauart $y \cdot \varphi(x)$ ersetzen. Dabei sollte sich φ für „kleine“ Populationsgrößen annähernd proportional verhalten ($\varphi(x) \approx \beta x$ für „kleine“ x) und für „große“ Populationsgrößen nahezu konstant sein ($\varphi(x) \approx S$ für x „groß“). Der Parameter S hat dann die Dimension $\frac{1}{T}$ und gibt an, mit welcher Pro-Kopf-Rate die Räuber maximal von der Beutepopulation profitieren können. Die Funktion φ , die selber auch die Dimension $1/T$ hat, sollte also den in Abb. 6.21 dargestellten qualitativen Verlauf haben.

Nun gibt es viele Funktionsterme, die einen qualitativ ähnlichen Graphen erzeugen, und aus der Modellbildung heraus drängt sich keiner direkt auf, sodass wir nach einem mathematisch einfachen Term suchen können. Wir werden z. B. im Reich der gebrochen-rationalen Funktionen fündig (der Graph könnte ja eine Hyperbel sein) und wählen den Term

$$\varphi(x) = S \cdot \frac{x}{\frac{S}{\delta} + x} = \delta x \cdot \frac{1}{1 + \frac{\delta}{S}x}. \quad (6.34)$$

Abb. 6.21 Pro-Kopf-Reproduktionsrate aufgrund der Räuber-Beute-Interaktion mit Sättigung



Falls $x \ll S/\delta$, spielt der Summand x im Nenner keine Rolle und wir erhalten in diesem Fall $\varphi(x) \approx \delta x$. Falls $S/\delta \ll x$, spielt der erste Summand im Nenner keine Rolle und es gilt $\varphi(x) \approx S$. In Bezug auf die Beutepopulation ist es plausibel, anzunehmen, dass der Schaden, den die Beutepopulation hat, proportional zum Nutzen der Räuberpopulation ist, wir den Term $y\varphi(x)$ also nur mit einer geeigneten Konstante zu multiplizieren haben. Dahinter steht die Vermutung, dass die Räuber auch die Beutepopulation nicht mehr schädigen, wenn sie satt sind. (In Aufg. 6.8 soll eine „blutrünstige“ Räuberpopulation modelliert werden, welche die Beutepopulation ungesättigt schädigt, ohne davon zu profitieren.) Wir richten den Proportionalitätsfaktor so ein, dass für kleine Beutepopulationen approximativ der bekannte Interaktionsterm βxy auftritt. Dazu führen wir $S_B = \frac{\beta S}{\delta}$ ein und betrachten das System:

$$\begin{aligned}\dot{x} &= \alpha x \left(1 - \frac{x}{K}\right) - \beta xy \cdot \frac{1}{1 + \frac{\beta}{S_B} x} \\ \dot{y} &= -\gamma y + \delta xy \frac{1}{1 + \frac{\delta}{S} x}\end{aligned}\tag{6.35}$$

Der Parameter S_B hat die Dimension $\frac{A_B}{T A_R}$ und gibt an, wie viel Beute pro Räuber und Zeit maximal getötet wird. Wir bezeichnen das System (6.35) als Räuber-Beute-System mit Sättigung. Wir haben für den Interaktionsterm eine solche algebraische Darstellung gewählt, welche die Beziehung zum alten Interaktionsmodell $-\beta xy$ bzw. $+\delta xy$ klar heraustreten lässt. Insbesondere lässt sich im Limes $S \rightarrow \infty$ das alte Modell wiedererkennen. Wir werden dieses Modell im folgenden Abschnitt diskutieren.

In der Literatur werden auch andere Modelle diskutiert, so werden z. B. in [65] die folgenden Modelle für die Funktion φ erwähnt (siehe auch Abb. 6.22):

$$\varphi_2(x) = \frac{Sx^2}{V^2 + x^2} \quad \text{und} \quad \varphi_3(x) = S \cdot (1 - e^{-\nu x}).$$

Das auf φ_2 basierende Modell enthält zusätzlich zum Sättigungseffekt einen Schwelleneffekt, wie wir ihn bereits in Abschn. 5.3 kennengelernt haben, und soll in Auf. 6.11 diskutiert werden. Weitere Räuber-Beute-Modelle werden z. B. in [58, 68] diskutiert.

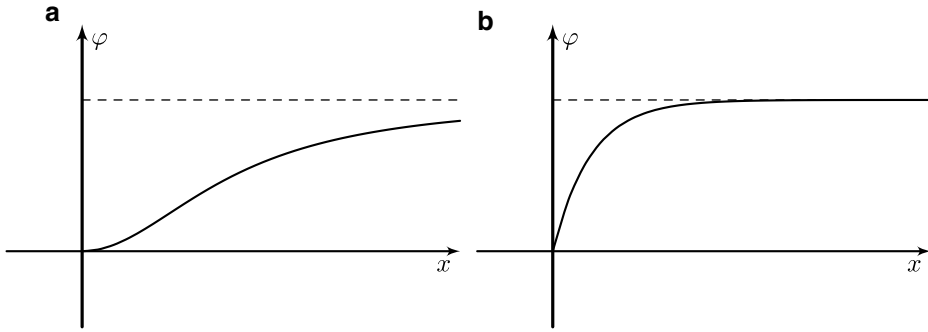


Abb. 6.22 Beispiele von alternativen Räuber-Beute-Interaktionstermen $y \cdot \varphi(x)$. a: $\varphi(x) = \frac{Sx^2}{V^2+x^2}$, b: $\varphi(x) = S \cdot (1 - \exp(-vx))$

6.5.2 Diskussion des Räuber-Beute-Modells mit Sättigung

Entdimensionalisierung

Das gewählte Modell enthält drei Variablen und sechs Parameter. S_B ist kein eigenständiger Parameter, sondern hängt von S , β und δ ab. Durch eine Entdimensionalisierung können wir die Anzahl der Parameter auf drei reduzieren. Das entdimensionalisierte System lautet mit unseren alten Bezeichnungen

$$\begin{aligned}\dot{u} &= \alpha \bar{t} u \left(1 - \frac{\bar{x}}{K} u \right) - \bar{t} \bar{y} \beta \frac{uv}{1 + \frac{\bar{x}\delta}{S} u} \\ \dot{v} &= -\delta \bar{t} v + \bar{t} \bar{x} \delta \frac{uv}{1 + \frac{\bar{x}\delta}{S} u}.\end{aligned}\tag{6.36}$$

Es gibt viele Möglichkeiten, Skalen zu wählen. Wir wollen im Sinne der besseren Vergleichbarkeit bei unseren alten Skalen bleiben und benutzen weiterhin

$$\bar{t} = \frac{1}{\alpha}, \quad \bar{x} = \frac{\gamma}{\delta} \quad \text{und} \quad \bar{y} = \frac{\alpha}{\delta},$$

das heißt die Populationsskalen, sind immer noch die stationären Populationen aus dem Lotka-Volterra-Modell. Mit dieser Wahl können wir (6.36) schreiben als

$$\begin{aligned}\dot{u} &= u \left(1 - bu - v \cdot \frac{1}{1 + cu} \right) \\ \dot{v} &= av \cdot \left(-1 + u \cdot \frac{1}{1 + cu} \right)\end{aligned}\tag{6.37}$$

mit den dimensionslosen Parametern

$$a = \frac{\gamma}{\alpha}, \quad b = \frac{\gamma}{\delta K} \quad \text{und} \quad c = \frac{\gamma}{S}.\tag{6.38}$$

Wir wollen uns zunächst überlegen, in welcher Größenordnung der Parameter c plausiblerweise liegt: Er ist das Verhältnis der Pro-Kopf-Sterberate zur maximalen Pro-Kopf-Geburtenrate infolge von Räuber-Beute-Interaktion. Es sollte also $c < 1$ gelten, damit die Wachstumsrate der Räuberpopulation nicht immer negativ ist und die Räuberpopulation aussterben muss.

Der Parameter b gibt das Verhältnis der Populationsgröße an, in der sich die Beutepopulation ohne intraspezifische Konkurrenz aber unter Bejagungsdruck der Räuberpopulation befinden würde (γ/δ), zu der Populationsgröße, in der sich die Beutepopulation unter Einfluss der intraspezifischen Konkurrenz, aber ohne Bejagungsdruck durch die Räuberpopulation befinden würde (K). Ein Parameterwert $b > 1$ sagt also aus, dass die intraspezifische Konkurrenz einen größeren Effekt hat als der Bejagungsdruck, für $0 < b < 1$ ist es umgekehrt. Wir erinnern uns, dass im ersten Fall bei einer Räuber-Beute-Interaktion ohne Sättigung ($S = \infty$ bzw. $c = 0$) die Räuberpopulation ausstirbt (siehe Abschn. 6.4.6).

Trajektorien in den Achsen, Nullklinen und stationäre Punkte

Wir bemerken zunächst, dass sich an der Situation in den Achsen im Vergleich zum Räuber-Beute-Modell mit logistischem Wachstum nichts geändert hat. Alle Änderungen betreffen ja nur den Interaktionsterm. Wir berechnen nun die Nullklinen von (6.36):

$$\begin{aligned} \dot{u} = 0 &\Leftrightarrow u = 0 \quad \vee \quad v = (1 + cu)(1 - bu) \\ \dot{v} = 0 &\Leftrightarrow v = 0 \quad \vee \quad u = \frac{1}{1 - c} \end{aligned}$$

Die u -Nullkline ist der Graph einer quadratischen Funktion mit den Nullstellen $-1/c$ und $1/b$ und der Scheitelpunktstelle $u_S = \frac{1}{2} \left(\frac{1}{b} - \frac{1}{c} \right)$. Wir müssen also eine ganze Menge qualitativ unterschiedlicher Fälle betrachten (siehe Abb. 6.23) mit $u_{\dot{v}=0} := \frac{1}{1-c}$, nämlich

- I $1/b < u_{\dot{v}=0}$,
- II $u_S < u_{\dot{v}=0} < 1/b$ und $0 < u_{\dot{v}=0}$,
- III $0 < u_{\dot{v}=0} < u_S$ sowie
- IV $u_{\dot{v}=0} < 0$.

Wir stellen diese vier Bereiche übersichtlich in einem b - c -Diagramm dar, siehe Abb. 6.24.

Der dritte Fall kann natürlich nur auftreten, wenn $u_S > 0$. Im vierten Fall kann der Schnittpunkt der v -Nullkline $u = u_{\dot{v}=0}$ mit der u -Nullkline $v = (1 + cu)(1 - bu)$ auch im vierten Quadranten liegen; in beiden Fällen ist der stationäre Punkt P_2^* für den Kontext des Räuber-Beute-Modells irrelevant.

Wir beschränken uns in den weiteren Analysen auf die interessanten Fälle II und III mit einem stationärer Koexistenzzustand im ersten Quadranten. Man kann relativ einfach zeigen, dass in den Fällen I und IV die Räuberpopulation auf lange Sicht ausstirbt.

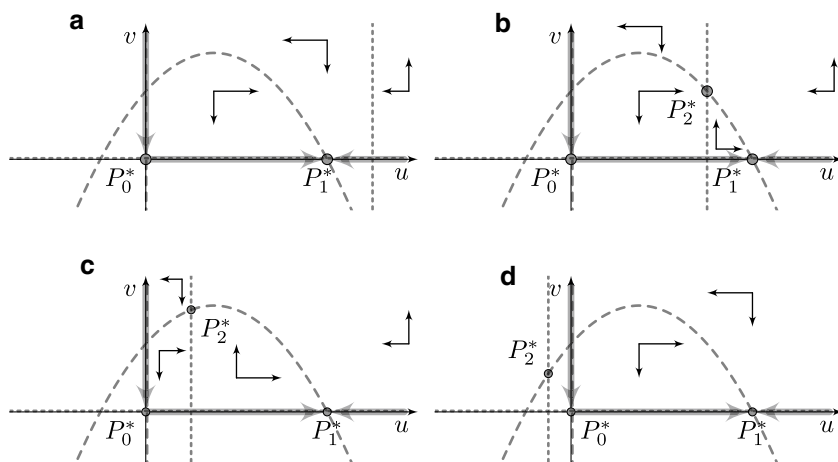
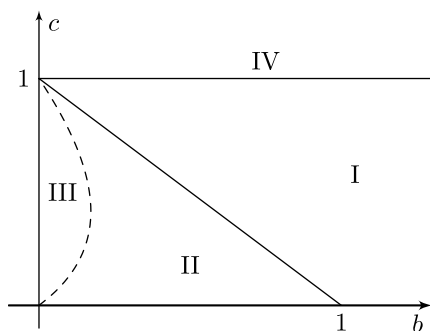


Abb. 6.23 Nullklinen und stationäre Punkte für das Räuber-Beute-System mit Sättigung (6.36), a Fall I, b Fall II, c Fall III, d Fall IV

Der nächste Schritt in unserer Analyse ist die Bestimmung der Art der stationären Punkte. Für P_0^* und P_1^* wird das schon aus den Trajektorien in den Achsen und den Richtungen deutlich: Beide müssen wohl Sattelpunkte sein (wir werden das gleich bestätigen). Der stationäre Punkt P_2^* hat die Koordinaten $u_2^* = \frac{1}{1-c}$ und $v_2^* = \frac{1-b-c}{(1-c)^2}$. Es sieht unangenehm aus, diesen Ausdruck in die Jacobi-Matrix einzusetzen und die Eigenwerte oder zumindest deren Vorzeichen zu berechnen. Wir wenden stattdessen ein grafisches Verfahren an, das uns lediglich die Vorzeichen der Einträge der Jacobi-Matrix liefert, die sogenannte Signaturmatrix.

Der Fall II: $u_S < u_{i=0} < 1/b$ und $0 < u_{i=0}$ Wir werden die Methode der Bestimmung der Signaturmatrix mehrmals benutzen, deshalb kommt die Erklärung in einen Kasten.

Abb. 6.24 b - c -Diagramm zur Darstellung der unterschiedlichen Parameterbereiche für das System (6.37)



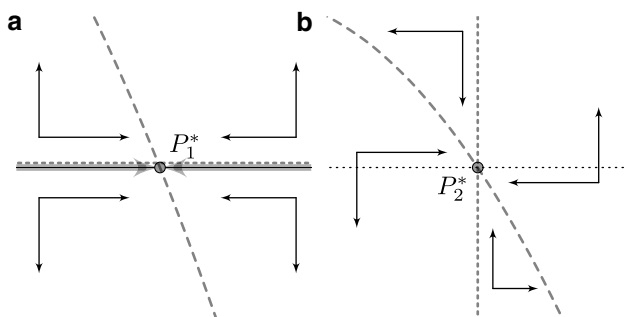


Abb. 6.25 Die Vorzeichensituation an den stationären Punkten P_1^* (a) und P_2^* (b) im Fall II

Die Bestimmung der Signaturmatrix aus den Vorzeichen der Ableitungen

Wir starten mit dem Punkt P_2^* . Die Vorzeichensituation ist in Abb. 6.25(b) in größerem Format dargestellt. Die unser System (6.36) definierende Abbildung bezeichnen wir wieder mit $F = (F_1, F_2)^t$. Im obere linken Eintrag der Jacobi-Matrix steht die partielle Ableitung von F_1 nach u . Wir müssen also die Ableitung der Funktion $s \mapsto F_1(s, v_2^*)$ an der Stelle $s = u_2^*$ bilden. Diese Funktion hat an der Stelle $s = u_2^*$ einen Vorzeichenwechsel von $+$ nach $-$, wir laufen nämlich in Abb. 6.25(b) längs der gepunkteten Linie und am Punkt P_2^* wechselt das u -Wachstum von $+$ nach $-$. Dann kann die Ableitung dieser Funktion an dieser Stelle aber nicht positiv sein. Wir wenden jetzt ein rationales und nichtdeduktives Argument an: Es gibt keinen plausiblen Grund dafür, dass die angesprochene Funktion genau hier eine Wendestelle haben sollte. Also folgern wir (nur rational, nicht deduktiv), dass die Ableitung negativ und damit auch der Eintrag links oben in der Jacobi-Matrix negativ ist.

Für den Eintrag rechts oben in der Jacobi-Matrix benötigen wir die partielle Ableitung von F_1 nach v am Punkt P_2^* , also die Ableitung der Funktion $s \mapsto F_1(u_2^*, s)$ an der Stelle $s = v_2^*$. Diese Funktion hat an dieser Stelle einen Vorzeichenwechsel von $+$ nach $-$, wir gehen nämlich jetzt senkrecht durch den Punkt P_2^* und das u -Wachstumsverhalten ändert sich dabei gerade von $+$ nach $-$. Damit kann diese Ableitung nicht positiv sein. Wir schließen rational, sie ist negativ.

Für den Eintrag links unten in der Jacobi-Matrix benötigen wir die partielle Ableitung von F_2 nach u am Punkt P_2^* , also die Ableitung der Funktion $s \mapsto F_2(s, v_2^*)$ an der Stelle $s = u_2^*$. Diese Funktion hat an dieser Stelle einen Vorzeichenwechsel von $-$ nach $+$, wir gehen nämlich längs der gestrichelten Linie durch den Punkt P_2^* und das v -Wachstumsverhalten ändert sich dabei gerade von $-$ nach $+$. Damit kann diese Ableitung nicht negativ sein. Wir denken rational, sie ist positiv.

Für den Eintrag rechts unten in der Jacobi-Matrix benötigen wir die partielle Ableitung von F_2 nach v am Punkt P_2^* , also die Ableitung der Funktion $s \mapsto F_2(u_2^*, s)$ an der Stelle $s = v_2^*$. Diese Funktion ist aber konstant null, wir laufen nämlich die v -Nullkline entlang durch P_2^* und hier verschwindet das v -Wachstum, also F_2 . Damit ist diese Ableitung null. Die Jacobi-Matrix hat also die folgende Signatur:

$$J(P_2^*) = \begin{pmatrix} - & - \\ + & 0 \end{pmatrix}$$

Daraus lässt sich $\text{tr} J(P_2^*) < 0$ und $\det J(P_2^*) > 0$ schließen, also ist P_2^* ein stabiler Knoten oder ein stabiler Strudel, auf jeden Fall aber asymptotisch stabil. Die Unterscheidung zwischen Strudeln und Knoten kann dieses Verfahren nicht leisten.

Aus der Vorzeichensituation bei P_1^* (siehe Abb. 6.25(a), diesmal benötigen wir tatsächlich auch die Vorzeichen im vierten Quadranten) erhalten wir die Signaturmatrix

$$J(P_1^*) = \begin{pmatrix} - & 0 \\ 0 & + \end{pmatrix},$$

P_1^* ist also ein Sattelpunkt. Analoge Betrachtungen bestätigen ebenfalls, dass P_0^* ein Sattelpunkt ist. Damit können wir das Phasenporträt skizzieren, siehe Abb. 6.26.

Dabei gehen wir wieder stillschweigend davon aus, dass es im Bereich rechts außen, also $v > 0$ und $v > (1 - bu)(1 + cu)$ und $u > u_{\dot{v}=0}$, keine in Durchlaufrichtung unbeschränkten Halbtrajektorien gibt. In Aufg. 6.9 soll das, analog zum Vorgehen in Abschn. 6.4.7, ausgeschlossen werden. Wir erwarten also, dass sich das ungestörte System auf lange Sicht im Koexistenzzustand P_2^* befindet und bei nicht mitmodellierten Störungen aus diesem Gleichgewichtszustand wieder in diesen zurückkehrt. Wir werfen noch einmal einen Blick auf die Bedingungen: Die Bedingungen sind äquivalent zu $1 - c > b$, also insbesondere $c < 1$ und $b < 1$ (b und c sind immanent positiv). In Abschn. 6.4.6 haben wir gesehen, dass auch im ungesättigten Fall $c = 0$ für $b < 1$ eine stabile Koexistenz eintritt. Der Sättigungseffekt spielt also in dieser Parameterkonstellation keine große Rolle und kann insbesondere nicht die stabile 10-Jahres-Periodizität der Luchspopulation Kanadas erklären.

Der Fall III: $0 < u_{\dot{v}=0} < u_S$ Die Vorzeichensituation an den stationären Punkten P_1^* und P_2^* ist in Abb. 6.27 dargestellt.

Damit ergeben sich die folgenden Signaturmatrizen

$$J(P_1^*) = \begin{pmatrix} - & 0 \\ 0 & + \end{pmatrix} \quad \text{und} \quad J(P_2^*) = \begin{pmatrix} + & - \\ + & + \end{pmatrix}.$$

Damit ist in diesem Fall P_2^* ein instabiler Knoten oder Strudel, während sich an der Situation bei P_1^* nichts geändert hat, P_1^* ist ebenso wie P_0^* ein Sattelpunkt. Wollen wir nun das Phasenporträt skizzieren, erscheint ein Problem (siehe Abb. 6.28): Einerseits drehen sich die Trajektorien von außen kommend in den Bereich um P_2^* hinein, andererseits ist P_2^*

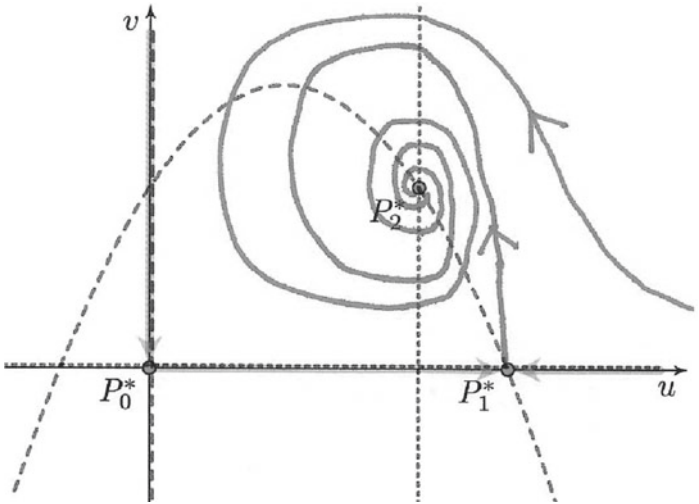


Abb. 6.26 Phasenportrait im Fall II: Auf lange Sicht befindet sich das System im Koexistenzzustand P_2^*

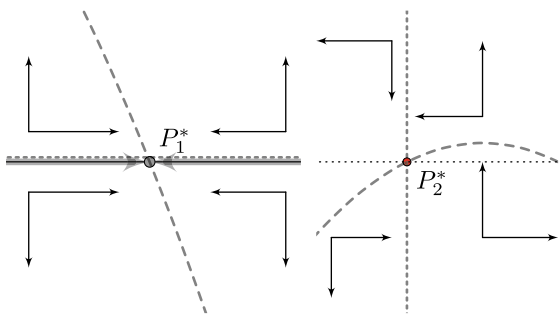
instabil und die Trajektorien müssen den direkten Bereich um P_2^* verlassen. Wie lässt sich das lösen? Das soll im nächsten Abschnitt diskutiert werden.

Die noch fehlenden Fälle I und IV lagern wir in Aufg. 6.10 aus.

6.5.3 Stabile Grenzyklen und der Satz von Poincaré-Bendixson

Um eine Idee zu erhalten, wie sich diese Situation klärt, lösen wir für geeignete Parameterwerte a , b und c das Differentialgleichungssystem (6.36) numerisch mit dem expliziten Euler-Verfahren für einige Anfangsdaten. Mit einer Schrittweite h implementieren wir das folgende Verfahren in ein Tabellenkalkulationsprogramm:

Abb. 6.27 Die Vorzeichensituation an den stationären Punkten P_1^* und P_2^* im Fall III



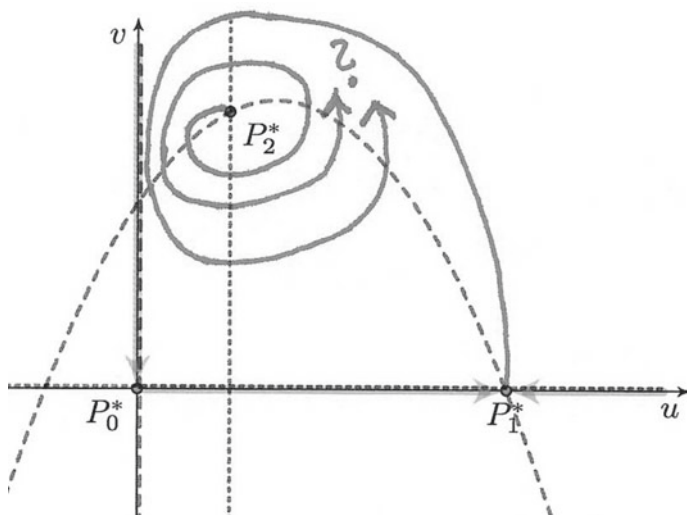


Abb. 6.28 Phasenportrait im Fall III: P_2^* ist instabil

$$u_{n+1} = u_n + hu_n \left(1 - bu_n - \frac{v_n}{1 + cu_n} \right)$$

$$v_{n+1} = v_n + hv_n \left(-1 + \frac{u_n}{1 + cu_n} \right)$$

Konkret wählen wir die Parameterwerte $a = 1$, $b = 0.125$ und $c = 0.7$: Der stationäre Punkt P_2^* hat dann die Koordinaten $(\frac{10}{3}, \frac{35}{18}) = (3.333, 1.944)$. Auf diese konkreten Werte sind wir durch Ausprobieren gekommen, in diesem Fall entstehen gut darstellbare Plots. Wir plotten die numerischen Lösungen zu den Anfangsdaten $(3, 2)$, $(1, 2)$ sowie $(4.2038058516, 1.8242244313)$ der Schrittweite $h = 0.1$ und 2000 Iterationsschritten in ein u - v -Diagramm, siehe Abb. 6.29. Die Trajektorie zum ersten Anfangsdatum dreht sich ein, die zum zweiten dreht sich aus, die zum dritten scheint eine geschlossene Kurve, einen sogenannten Zyklus zu bilden, welcher die beiden ersten Trajektorien voneinander trennt. (Auf das Anfangsdatum der dritten Lösung sind wir gekommen, indem wir das Anfangsdatum der zweiten Lösung ca. 8000-mal iteriert und das letzte Wertepaar als neues Anfangsdatum benutzt haben.)

Das gerechnete Beispiel deutet darauf hin, dass es eine geschlossene Kurve, einen sogenannten Grenzzzyklus gibt, der die sich aus dem instabilen stationären Punkt P_2^* herausdrehenden Trajektorien von den von außen kommenden, sich eindrehenden Trajektorien trennt. Ferner scheint dieser Grenzzzyklus anziehend zu sein, denn alle Trajektorien nähern sich diesem Grenzzzyklus von innen oder von außen an. Wir sprechen von einem *stabilen Grenzzzyklus*. Dieser Grenzzzyklus ist die Trajektorie periodischer Lösungen des Differentialgleichungssystems, ähnlich wie die Lösungen zum Lotka-Volterra-System in Abschn. 6.3. Im Gegensatz zur Situation dort gibt es hier aber nur einen Zyklus. Die u - und v -Komponenten

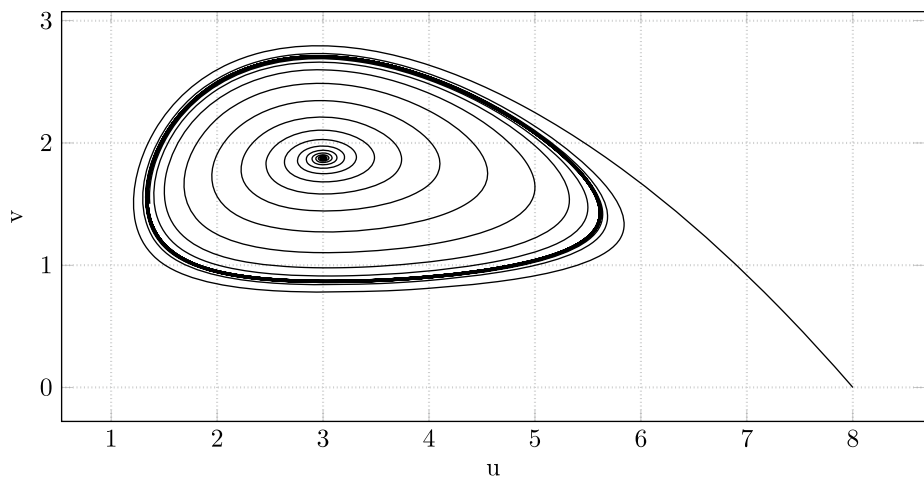


Abb. 6.29 Numerische Lösung von (6.36) zu den Parametern $a = 1$, $b = 0.125$ und $c = 0.7$ und verschiedenen Anfangswerten

der numerisch berechneten Lösung, die in dem Grenzyklus verläuft, sind in Abb. 6.30 geplotet. Es ist plausibel, dass auf lange Sicht jede beliebige Lösung im ersten Quadranten eine passende zeitliche Verschiebung dieser Lösung approximiert.

Numerische Versuche mit anderen Parameterkombinationen b, c aus dem Bereich III zeigen qualitativ dasselbe Verhalten: einen stabilen Grenzyklus, dem sich die Trajektorien von

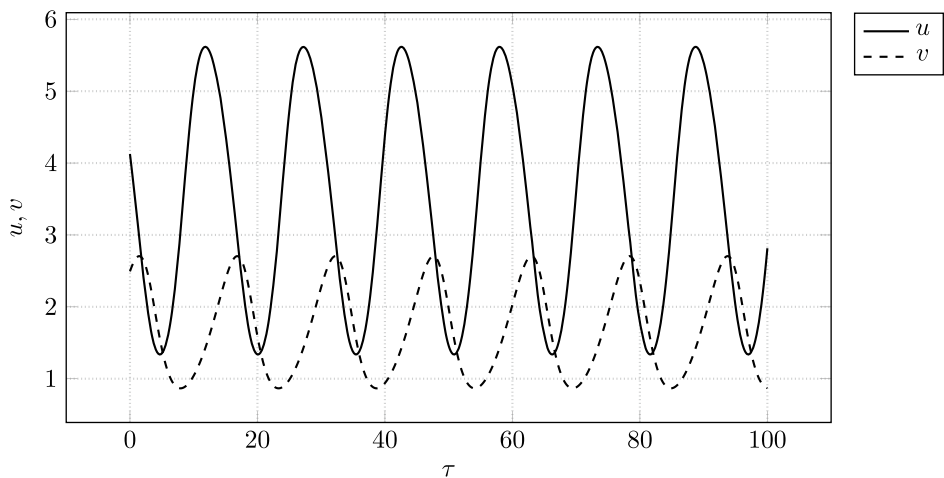


Abb. 6.30 Numerisch ermittelte periodische Lösung im Grenzyklus für $b = 0.125$ und $c = 0.7$

innen und außen annähern. Wir schließen rational, dass das für alle Parameterkombinationen aus dem Bereich III so ist.

Das bedeutet, dass alle Lösungen von (6.36) auf lange Sicht in dem Grenzyklus verlaufen und somit ein periodisches Verhalten mit einer festen Periode aufweisen. Im Gegensatz zum Lotka-Volterra-Modell gibt es hier nur einen Zyklus und nur eine Periode. Wir haben also endlich ein Modell gefunden, in dessen Rahmen sich die stabile 10-Jahres-Periode der Schwankungen der Luchspopulationen erklären lässt (siehe Abb. 6.31).

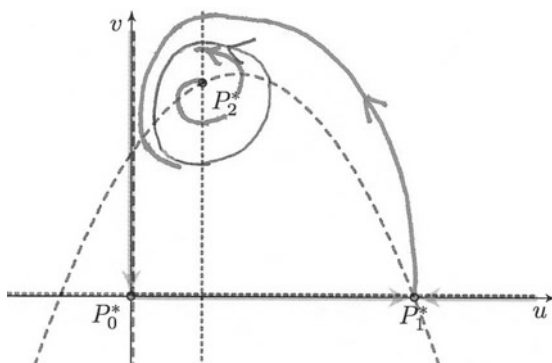
Erfreulicherweise sind wir in dieser Frage nicht nur auf rationale Schlüsse angewiesen, sondern es gibt tatsächlich einen mathematischen Satz, der die Wahrheit der aufgestellten Vermutungen deduktiv beweist, nämlich der Satz von Poincaré-Bendixson, der 1882 von Poincaré vermutet (siehe [72]) und 1901 von dem schwedischen Mathematiker Ivar Bendixson bewiesen wurde (siehe [7]). Wir wollen die Aussage des Satzes im folgenden Abschnitt darstellen und begründen, dass unser System (6.36) die Voraussetzungen des Satzes erfüllt.

Der Satz von Poincaré-Bendixson Um den Satz vernünftig zu formulieren, müssen wir zunächst präzisieren, was es bedeuten soll, dass sich eine Trajektorie an eine andere anschmiegt. Dazu führen wir die ω -Limesmenge einer Trajektorie γ ein. Sei $t \mapsto x(t)$ die (maximale) Lösung eines Differentialgleichungssystems $\dot{x} = F(x)$ und γ die zugehörige Trajektorie. Ein positiver Grenzpunkt x von γ ist ein Punkt, für den es eine monoton wachsende Folge (t_n) gibt mit $\lim_{n \rightarrow \infty} t_n = \infty$, so dass $\lim_{n \rightarrow \infty} x(t_n) = x$ gilt. Das heißt, die Lösung muss dem Punkt x in der Zukunft immer wieder beliebig nahekommen, darf sich zwischenzeitlich aber auch wieder entfernen. Wir definieren die ω -Limesmenge von γ als die Menge dieser Punkte,

$$\omega(\gamma) := \{x \in \mathbb{R}^n : x \text{ ist positiver Grenzpunkt von } \gamma\}.$$

Weil alle Lösungen $t \mapsto x(t)$, die in der betrachteten Trajektorie γ verlaufen, Translationen voneinander sind, ist die ω -Limesmenge wohldefiniert, also unabhängig von der Wahl der konkreten Lösung. Mit dieser Definition können wir den Begriff *stabiler Grenz-*

Abb. 6.31 Im Fall III gibt es einen stabilen Grenzyklus, dem sich alle nicht-stationären Trajektorien von innen oder außen anschmiegen



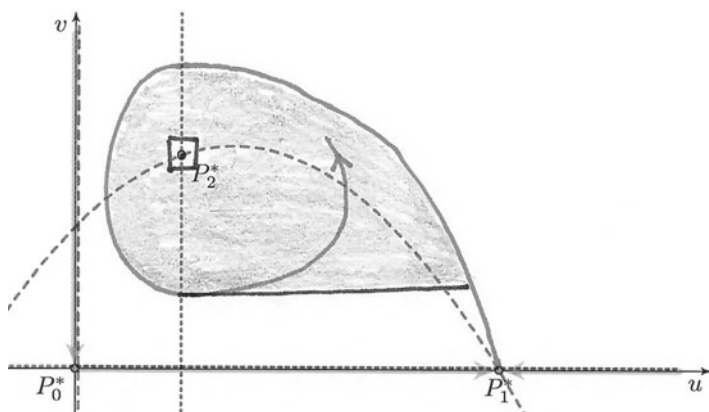


Abb. 6.32 Die Voraussetzungen des Satzes von Poincaré-Bendixson sind erfüllt: die Halbtrajektorie und die kompakte Menge K schraffiert

zyklus präzisieren: Eine Trajektorie γ heißt stabiler Grenzyklus, wenn sie die Trajektorie einer periodischen Lösung und die ω -Limesmenge einer anderen Trajektorie γ' ist. Die zweite Bedingung ist wesentlich, sonst wären ja auch alle geschlossenen Trajektorien im Lotka-Volterra-System bereits stabile Grenzyklen. So muss sich aber auch eine „fremde“ Trajektorie „anschießen“.

Nun können wir den Satz von Poincaré-Bendixson formulieren, so wie er z. B. in [71] zu finden ist. Für den nicht elementaren Beweis verweisen wir ebenfalls auf diese Literaturangabe.

Satz 6.3 (Der Satz von Poincaré-Bendixson)

Gegeben sei eine stetig differenzierbare Abbildung $F : \mathbb{R}^2 \rightarrow \mathbb{R}^2$ zusammen mit dem zugehörigen Differentialgleichungssystem

$$\dot{x} = F(x). \quad (6.39)$$

Enthält eine kompakte Menge $K \subset \mathbb{R}^2$ keinen stationären Punkt von (6.39) und gibt es eine Lösung $t \mapsto x(t)$ von (6.39) und ein $t_0 \in \mathbb{R}$, sodass $x(t) \in K$ für alle $t \geq t_0$, dann ist diese Lösung entweder selber periodisch oder K enthält einen stabilen Grenzyklus, dem sich die Lösung anschießt.

Es ist bemerkenswert, dass dieser Satz nur in der Ebene gilt. Dies liegt daran, dass im Höherdimensionalen eine geschlossene Kurve den Raum nicht in eine beschränkte innere und eine unbeschränkte äußere Menge zerlegt, wie das in der Ebene der Fall ist.

Erfüllt unser System nun die Voraussetzungen von diesem Satz? Dafür müssen wir eine kompakte Menge K , die keinen stationären Punkt enthält, und eine „Zukunfts-Halbtrajektorie“ angeben, die ganz in K verläuft. Als Trajektorie wählen wir die aus dem Sattelpunkt P_1^* hinauslaufende Trajektorie. Für die Menge K wählen wir das in Abb. 6.32

schräffierte Gebiet, aus dem wir den stationären Punkt P_2^* mit einer hinreichend kleinen offenen Umgebung ausgeschnitten haben. Nun ist das Gebiet kompakt, enthält keinen stationären Punkt, wohl aber eine Halbtrajektorie: Weil P_2^* ein instabiler Knoten oder Strudel ist, kann nämlich keine Lösung in das ausgeschnittene Loch hineinlaufen.

Aufgaben

6.1 Das Phasenprotrait zerfällt in disjunkte Trajektorien

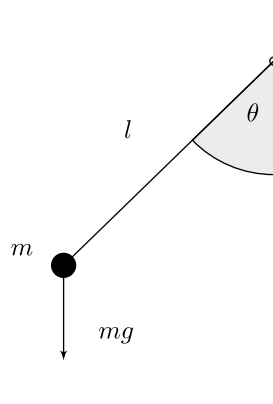
- Die Funktionen y und z seien zwei maximale Lösungen des zeitautonomen DGLS $\dot{x} = F(x)$ mit Existenzintervall I_y bzw. I_z , deren Trajektorien sich schneiden, wobei $F : \mathbb{R}^n \rightarrow \mathbb{R}^n$ lokal Lipschitz-stetig ist. Zeigen Sie: Es gibt ein $\hat{t} \in \mathbb{R}$, sodass $I_y = I_z + \hat{t} := \{t + \hat{t} : t \in I_z\}$ und $y(t) = z(t + \hat{t})$ für alle $t \in I_y$, die beiden Lösungen sind also Zeittranslationen voneinander.
- Wir können auch für ein nichtzeitautonomes DGLS die Trajektorie einer Lösung in gleicher Art und Weise definieren wie in (6.7). Zeigen Sie, dass in diesem Fall im Allgemeinen zwei Trajektorien nicht entweder disjunkt oder gleich sein müssen.

6.2 Das mathematische Pendel ohne und mit Reibung

Wir betrachten einen masselosen Stab der Länge l , der an einem Ende aufgehängt ist und am anderen Ende einen Massenpunkt m trägt, siehe Abb. 6.33. Mit θ bezeichnen wir den Auslenkungswinkel des Pendels aus der Vertikalen. Nach den Gesetzen (oder den Modellen) der Mechanik ist die Winkelbeschleunigung $\ddot{\theta}$ proportional zum Kraftmoment des Gewichts:

$$I\ddot{\theta} = -mgl \sin \theta, \quad (6.40)$$

Abb. 6.33 Das mathematische Pendel mit der Stablänge l , einem Pendelkörper der Masse m , dem Auslenkungswinkel θ und der Gravitationskraft mg



wobei $I = ml^2$ das Trägheitsmoment ist.

- a) Entdimensionalisieren Sie die Gleichung und überführen Sie diese Differentialgleichung zweiter Ordnung durch eine geeignete Substitution in ein Differentialgleichungssystem erster Ordnung.

Tipp: Führen Sie als zweite Variable die Winkelgeschwindigkeit $\mu = \dot{\theta}$ ein.

- b) Bestimmen Sie eine Funktion $H : \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$, die konstant auf den Trajektorien der Lösungen des Differentialgleichungssystems aus a) ist. Diese Funktion nennt man die Energie des Systems.
- c) Plotten Sie z. B. mit GeoGebra eine geeignete Auswahl an Niveaulinien der Energie und erläutern Sie, wie die Bewegung des Pendels auf diesen Niveaulinien aussieht.
- d) In der Realität kann sich ein Pendel nicht reibungsfrei bewegen, in die Gl. (6.40) muss also noch ein Reibungsterm eingebaut werden. Eine plausible Modellierung ist, dass die Reibungskraft proportional zur Winkelgeschwindigkeit und dieser entgegengesetzt ist. Das führt auf die DGL

$$I\ddot{\theta} = -mgl \sin \theta - r\dot{\theta}. \quad (6.41)$$

Entdimensionalisieren Sie auch diese Gleichung, überführen Sie sie in ein DGLS erster Ordnung. Wie entwickelt sich die Energie H aus Aufgabenteil b) längs den Lösungen dieses DGLS? Was bedeutet das für die Bewegung des Pendels?

6.3 Phasenporträts für lineare Differentialgleichungssysteme bei nicht vollem Rang

Skizzieren Sie die möglichen Phasenporträts des linearen DGLS $\dot{\vec{x}} = A\vec{x}$, wenn die 2×2 -Matrix A

- a) den einfachen Eigenwert 0 hat.
b) den doppelten Eigenwert 0 hat.

Welches Stabilitätsverhalten haben die stationären Punkte?

6.4 Lineare Geichungssysteme

- a) Geben Sie, falls möglich, eine 2×2 -Matrix A an, die nur positive Einträge hat und
- zwei unterschiedliche positive Eigenwerte hat,
 - einen positiven und einen negativen Eigenwert hat,
 - zwei negative Eigenwerte hat,
 - zwei komplexe Eigenwerte hat,
 - einen doppelten positiven Eigenwert hat,

und begründen Sie, warum das ggf. nicht geht.

- b) Berechnen Sie zu den in (a) aufgestellten Matrizen ein Fundamentalsystem des DGLS $\dot{\vec{x}} = A\vec{x}$. Plotten Sie mit geeigneter Software eine Skizze des Phasenporträts.

6.5 Der funktionale Zusammenhang zwischen dem Niveau der H -Funktion und der Periode im Lotka-Volterra-Modell

Untersuchen Sie numerisch den Zusammenhang zwischen der Periode P von Lösungen des Lotka-Volterra-Systems (6.9) und dem zugehörigen Niveau der Lyapunov-Funktion H , (6.16), auch in Abhängigkeit von dem Parameter a .

6.6 Ein Räuber-Beute-System mit konkurrierenden Räubern

Gegeben ist das folgende Räuber-Beute-Modell (alle auftretenden Parameter sind positiv):

$$\begin{aligned}\dot{x} &= \alpha x \cdot \left(1 - \frac{x}{K_x}\right) - \beta xy \\ \dot{y} &= -\gamma y - \varepsilon y^2 + \delta xy\end{aligned}\tag{6.42}$$

- a) Erläutern Sie die Modellierungsannahmen.
b) Zeigen Sie, dass sich (6.42) durch eine geeignete Entdimensionalisierung auf die Form

$$\begin{aligned}\dot{u} &= u(1 - bu - v) \\ \dot{v} &= av(-1 + u - cv)\end{aligned}\tag{6.43}$$

bringen lässt. Geben Sie die dabei benutzten Skalen an und drücken Sie die dimensionslosen Parameter a, b, c durch die dimensionsbehafteten Parameter $\alpha, \beta, \gamma, \delta, K_x$ und ε aus.

- c) Bestimmen Sie **rechnerisch** die stationären Punkte des Systems.
d) Bestimmen Sie **rechnerisch** die Art der stationären Punkte (in Abhängigkeit von den dimensionslosen Parametern). Es reicht, die stationären Punkte im abgeschlossenen ersten Quadranten zu betrachten.
e) Geben Sie an, für welche Parameter es einen stabilen Koexistenzzustand im ersten Quadranten gibt, und zeichnen Sie für diesen Fall das Phasenporträt.
f) Interpretieren Sie für diesen Fall das Ergebnis.

Tipp für Aufgabenteil d): Der Koexistenzzustand $P_{ko} = (u_{ko}^*, v_{ko}^*)$ berechnet sich als Lösung des linearen Gleichungssystems $1 - bu - v = 0$ und $-1 + u - cv = 0$. Dadurch lässt sich die Jacobi-Matrix am Punkt P_{ko}^* folgendermaßen ausdrücken:

$$J(P_{ko}^*) = \begin{pmatrix} -bu_{ko}^* & u_{ko}^* \\ av_{ko}^* & -acv_{ko}^* \end{pmatrix}$$

Mit Hilfe dieser Darstellung lässt sich der Charakter des stationären Punktes einfach bestimmen, wenn man nicht zwischen Knoten und Strudel unterscheiden will.

6.7 Signaturmatrix und Stabilität

Gegeben sei ein Differentialgleichungssystem

$$\begin{aligned}\dot{u} &= f(u, v) \\ \dot{v} &= g(u, v)\end{aligned}\tag{6.44}$$

und (u^*, v^*) sei ein stationärer Punkt von (6.44). Untersuchen Sie systematisch, in welchen Fällen Sie etwas aus der Signatur der Jacobi-Matrix

$$J(u^*, v^*) =: A$$

über den Charakter des stationären Punkts aussagen können.

6.8 Ein Räuber-Beute-System mit blutrünstigen Räubern

Stellen Sie ein Räuber-Beute-Modell auf, in dem die Räuber-Beute-Interaktion einen Sättigungseffekt in der Räubergleichung aufweist, die Räuber aber trotzdem (blutrünstig) weiter jagen, obwohl sie keinen Nutzen mehr davon haben. Diskutieren Sie Ihr Modell möglichst umfassend. Beurteilen Sie, ob es nach den Aussagen dieses Modells eine solche Räuber-Beute-Interaktion in der Realität geben kann.

6.9 Unbeschränkte Lösungen im Räuber-Beute-System mit Sättigung?

Zeigen Sie, dass es im Räuber-Beute-Modell mit Sättigung (6.36) keine (in positiver Zeitrichtung) unbeschränkten Lösungen, also auch keine in Durchlaufrichtung unbeschränkten Halbtrajektorien gibt.

6.10 Das Räuber-Beute-System mit Sättigung – die Fälle I und IV

Diskutieren Sie für das System (6.36) in Abschn. 6.5.2 die Fälle I und IV.

6.11 Räuber-Beute-Modelle mit Schwelle und Sättigung

In dieser Aufgabe soll ein Räuber-Beute-System mit Sättigungs- und Schwelleneffekten diskutiert werden. Dazu betrachten wir den Interaktionsterm $y \cdot \varphi_2(x)$ mit

$$\varphi_2(x) = \frac{Sx^2}{V^2 + x^2}.\tag{6.45}$$

- Diskutieren Sie die beiden Funktionen φ aus (6.34) und φ_2 und vergleichen Sie sie.
- Beschreiben Sie die Unterschiede der beiden Wechselwirkungsterme $y\varphi(x)$ und $y\varphi_2(x)$ im Sachzusammenhang.

Wir betrachten das folgende Räuber-Beute-Modell:

$$\left. \begin{aligned}\dot{x} &= \alpha x \left(1 - \frac{x}{K}\right) - y\varphi_2(x) \\ \dot{y} &= -\gamma y + \delta y\varphi_2(x)\end{aligned}\right\}\tag{6.46}$$

- c) Zeichnen Sie die Nullklinen, bestimmen Sie grafisch die stationären Punkte sowie ihren Charakter und skizzieren Sie die möglichen Phasenporträts (Fallunterscheidung).
- d) Zeigen Sie, dass jede Trajektorie, die im ersten Quadranten beginnt, für $t \geq 0$ beschränkt bleibt.
- e) Geben Sie an, für welche Parameter es einen stabilen stationären Punkt und für welche es einen Grenzzyklus gibt. (Gefragt ist nach möglichst einfachen Ungleichungen zwischen den Parametern. Sie dürfen dazu alle technischen Hilfsmittel heranziehen, die Ihnen einfallen.)
- f) Wir betrachten zwei Arten B (Beute) und R (Räuber) mit folgenden Parametern:

$\alpha :$	0.7/Jahr	$V :$	80 Beutetiere
$K :$	15 000 Beutetiere	$S :$	300 Beutetiere pro Jahr und Räuber
$\gamma :$	2 /Jahr	$\delta :$	0.1/Raubtier pro Beutetier

Befindet sich das System gemäß der Modellierung auf lange Zeit in Koexistenz oder stirbt eine Spezies aus? Gibt es ein periodisches Verhalten?

6.12 Symbiosemodelle

Viele Tier- und Pflanzenarten leben mit anderen Arten in einer symbiotischen Beziehungen, das heißt, beide Populationen profitieren voneinander. In der mathematischen Modellierung unterscheiden sich Modelle für symbiotische Beziehungen von Räuber-Beute-Beziehungen durch die Vorzeichen in den Interaktionstermen, welche in symbiotischen Beziehungen für beide Partner positiv sind. Stellen Sie möglichst viele mathematische Modelle für symbiotische Beziehungen auf und diskutieren Sie diese mit den vorgestellten Methoden. Denken Sie dabei an

- intraspezifische Konkurrenz,
- Sättigungseffekte und
- Schwelleneffekte.

Finden Sie ein Modell mit periodischen Lösungen?

Bemerkung: Eine umfassende Bearbeitung dieser Aufgabe könnte sehr umfangreich werden. Ganz erhellend ist vielleicht auch ein Blick in die Literatur, z. B. [65].

6.13 Konkurrenzmodelle

Viele Tier- und Pflanzenarten leben mit anderen Arten in einer Konkurrenzbeziehungen, das heißt, beide Populationen schädigen sich gegenseitig. Man spricht in diesem Zusammenhang von *interspezifischer Konkurrenz* im Gegensatz zur intraspezifischen Konkurrenz. In der mathematischen Modellierung unterscheiden sich Modelle für Konkurrenzbeziehungen von Räuber-Beute-Beziehungen durch die Vorzeichen in den Interaktionstermen, welche in Konkurrenzbeziehungen für beide Populationen negativ sind.

Stellen Sie möglichst viele mathematische Modelle für Konkurrenzbeziehungen auf und diskutieren Sie diese mit den vorgestellten Methoden. Denken Sie dabei

- an intraspezifische Konkurrenz und
- Schwelleneffekte.

Sind Sättigungseffekte in Konkurrenzbeziehungen zu erwarten? Finden Sie ein Modell mit periodischen Lösungen?

Bemerkung: Eine umfassende Bearbeitung dieser Aufgabe könnte sehr umfangreich werden. Ganz erhellend ist vielleicht auch ein Blick in die Literatur, z. B. [65].

6.14 Schädlingsbekämpfung mit sterilen Insekten

Die *Steril Insect Release Method* (SIRM) zur Bekämpfung etwa der Pest besteht darin, in eine Insektenpopulation sterile Insekten einzubringen. Wir haben das bereits in Verbindung mit Differenzengleichungen kennengelernt.

Bringt man nun in eine Population X mit Größe x von Insekten y sterile Insekten ein (die Anzahl der sterilen Insekten wird konstant auf dem Wert y gehalten), so wird für die Entwicklung der Populationsgröße der fertilen Insekten die folgende Differentialgleichung vorgeschlagen:

$$\dot{x} = x \cdot \left(\frac{\alpha x}{x + y} - \beta \right) - \gamma x(x + y)$$

mit $\alpha, \beta, \gamma > 0$. Dabei sind α die Pro-Kopf-Geburtenrate und β die Pro-Kopf-Sterberate.

- Beschreiben Sie die Modellannahmen, die auf diese Differentialgleichung führen.
- Bestimmen Sie die kritische Zahl an sterilen Insekten y_c , die zu einem Aussterben der Insekten führt, und zeigen Sie, dass diese Zahl kleiner als ein Viertel der zugrunde liegenden Kapazität ist.
- Nehmen Sie an, dass die Insekten durch das einmalige Aussetzen von sterilen Insekten bekämpft werden sollen und die sterilen Insekten dieselbe Sterberate wie die fertilen Insekten haben. Stellen Sie ein geeignetes Differentialgleichungssystem für x und y auf und zeigen Sie, dass es nicht möglich ist, die Insekten durch ein einmaliges Aussetzen von sterilen Insekten zum Aussterben zu bringen.
- Wird ein Anteil δ an Insekten bereits steril geboren, wird das folgende Modell vorgeschlagen:

$$\begin{aligned}\dot{x} &= x \cdot \left(\frac{\alpha x}{x + y} - \beta \right) - \gamma x(x + y) \\ \dot{y} &= \delta x - \beta y\end{aligned}$$

Bestimmen Sie, wie groß δ sein muss, damit die Insektenpopulation ausstirbt.



Würfe, Enzymreaktionen und der Tannenwickler – asymptotische Methoden

7

Inhaltsverzeichnis

7.1	Der senkrechte Wurf im Galilei’schen und im Newton’schen Modell	224
7.2	Reaktionskinetik, Enzyme und die Methode des Asymptotic Matching	235
7.3	Der Tannenwickler und die Balsamtanne – zweiter Teil	245
	Aufgaben	264

In diesem Abschnitt soll eine wichtige und grundlegende Methode der Modellbildung vorgestellt werden. Die Idee ist sehr einfach: Um ein kompliziertes Modell einfacher zu machen, lasse man einfach die kleinen Effekte unberücksichtigt und konzentriere sich auf die wichtigen Elemente. Aber kann man immer erkennen, wann ein Effekt klein ist? Und was ist, wenn ein Effekt zwar klein, aber doch nicht völlig unwichtig ist? Kann man ihn dann auch nur halb berücksichtigen? Diese Fragen werden im kommenden Kapitel beantwortet. Wir beginnen in Abschn. 7.1 mit einem historischen Beispiel aus der Physik, den Bewegungen im Gravitationsfeld. In der Darstellung orientieren wir uns an [21].

In den folgenden Abschnitten in Abschn. 7.2 sollen die kennengelernten mathematischen Methoden genutzt werden, um einfache mathematische Modelle für chemische und biochemische Vorgänge aufzustellen. Wir stellen die Grundlagen der Reaktionskinetik dar und besprechen die Methode des *Asymptotic Matching* am Beispiel einer Enzymreaktion. Zum Abschluss dieses Kapitels greifen wir in Abschn. 7.3 die in Abschn. 5.4 begonnene Modellierung der Interaktion zwischen dem Tannenwickler und der Balsamtanne wieder auf und erklären die periodischen Wicklerepidemien mit anschließendem Zusammenbruch des Tannenbestands.

7.1 Der senkrechte Wurf im Galilei'schen und im Newton'schen Modell

7.1.1 Das Galilei'sche Modell des senkrechten Wurfs

Eine der ersten mathematischen Modellierungen von physikalischen Vorgängen auf der Erde stammt von Galileo Galilei: Er untersuchte, wie sich die Geschwindigkeiten von frei fallenden Kugeln verhalten (genauer gesagt untersuchte Galileo relativ schwere Kugeln, die sehr reibungsarm eine sogenannte Fallrinne herunterrollten).

Aus den Beobachtungen heraus war klar ersichtlich, dass die Geschwindigkeiten der Kugeln mit der Zeit zunehmen, unklar war aber, ob sich für diese Zunahme eine Gesetzmäßigkeit formulieren lässt. Galileo proklamiert nun, dass der Geschwindigkeitszuwachs nach der einfachsten möglichen Gesetzmäßigkeit vonstattengeht, die man sich nur vorstellen kann: In den *Discorsi* führt er aus:

„Wenn ich daher bemerke, dass ein aus der Ruhelage von bedeutender Höhe herabfallender Stein nach und nach neue Zuwächse an Geschwindigkeit erlangt, warum sollte ich nicht glauben, dass solche Zuwächse in allereinfachster, Jedermann plausibler Weise zu Stande kommen? Wenn wir genau aufmerken, werden wir keinen Zuwachs einfacher finden, als denjenigen, der in immer gleicher Weise hinzutritt.“ [33]

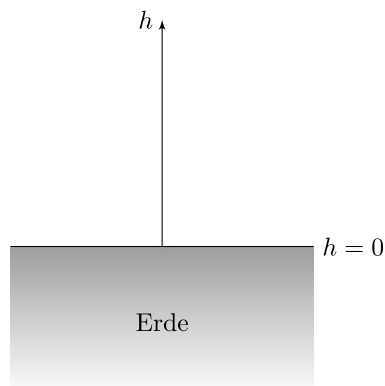
In der heutigen Sprache postuliert Galileo, dass die Änderung der Geschwindigkeit proportional zur Länge des Zeitintervalls ist, in dem sich die Änderung vollzieht. Mit anderen Worten: $\dot{v} = g$ mit einer Konstanten g . Interessanterweise stehen am Anfang also keine Messungen (außer der Beobachtung, dass die Geschwindigkeiten zunehmen), sondern eine vermutete Gesetzmäßigkeit, die sich nun vor dem Experiment bewähren muss, in dem nach Möglichkeit alle diese Gesetzmäßigkeit störenden Einflüsse wie Reibung und Luftwiderstand klein gehalten werden können.

Unternehmen wir experimentelle Untersuchungen des freien Falls, wie sie z.B. im Physikunterricht stattfinden, lässt sich diese Gleichung im Rahmen der experimentellen Möglichkeiten bestätigen und für die Proportionalitätskonstante ergibt sich der Wert $g = 9.81 \text{ m s}^{-2}$. Dieser Wert ist als die sogenannte *Erdbeschleunigung* bekannt. Wir führen nun ein Koordinatensystem wie in Abb. 7.1 ein. Da die Zunahme der Geschwindigkeit nach unten zeigt, erhalten wir für die Geschwindigkeit die Gleichung $\dot{v} = -g$. Integration liefert $v(t) = -gt + v_0$, wobei v_0 die Geschwindigkeit zur Zeit $t = 0$ ist. Wollen wir nun wissen, in welcher Höhe h sich die Kugel zur Zeit t befindet, müssen wir noch einmal integrieren und erhalten für die Höhe

$$h(t) = -\frac{1}{2}gt^2 + v_0t + h_0, \quad (7.1)$$

wobei h_0 die Höhe der Kugel zur Zeit $t = 0$ ist. Mit dieser Gleichung kann jetzt nicht nur der *freie Fall* behandelt werden, sondern es können alle möglichen Fragen beantwortet werden,

Abb. 7.1 Koordinatensystem für den senkrechten Wurf



z. B.: Wie hoch fliegt eine nach oben geworfene Kugel maximal? Wann schlägt die Kugel auf dem Boden auf? Zu welcher Zeit befindet sich die Kugel in dieser oder jenen Höhe?

Wir wollen dieses Modell in denselben mathematischen Kontext bringen wie die Räuber-Beute-Modelle, schreiben es also als ein Differentialgleichungssystem:

$$\begin{aligned}\dot{h} &= v \\ \dot{v} &= -g\end{aligned}\tag{7.2}$$

Wir könnten jetzt also auch wieder anfangen, Phasenporträts zu diesem System zu zeichnen. In Aufg. 7.1 sollen sie zeigen, dass auch dieses System ein erstes Integral besitzt. Die zugehörige Funktion ist proportional zur sogenannten Energie der Kugel. Das Modell ist gültig, solange $h > 0$ gilt, denn beim Aufprall auf dem Boden ist es mit dem senkrechten Wurf und der freien Fallbewegung vorbei. Aus der Bewegungsgleichung (7.1) oder auch aus der Energiefunktion (siehe Aufg. 7.1) sehen wir, dass, egal wie groß die Abwurfgeschwindigkeit auch sein mag, die geworfene Kugel immer wieder auf den Boden zurückfällt und nie in die Weiten des Weltalls entkommen kann.

7.1.2 Das Newton'sche Gravitationsmodell

Newton beschreibt die Wurfbewegung als Teil eines viel umfassenderen Bildes, insbesondere erklärt er auch, wodurch die Beschleunigung des Körpers bewirkt wird: Auf die Kugel wirkt die Gravitationskraft zwischen der Kugel und der Erde. Für diese Gravitationskraft postuliert Newton die folgende Gesetzmäßigkeit: Die Gravitationskraft ist proportional zur Masse der Erde, zur Masse der Kugel und zum Inversen des Abstandsquadrats, wobei der Abstand von der Kugel zum Mittelpunkt der Erde zu nehmen ist (siehe Abb. 7.2).

Das Newton'sche Modell ist dann gegeben durch die Differentialgleichungen

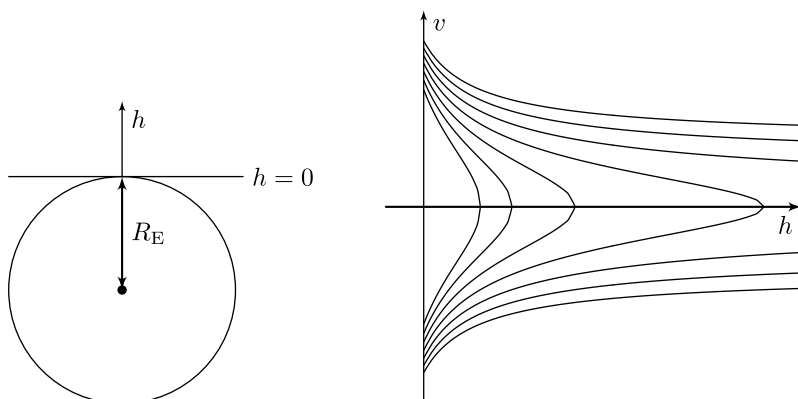


Abb. 7.2 Das Newton'sche Gravitationsmodell: Koordinatisierung und Phasenporträt

$$\begin{aligned}\dot{h} &= v \\ \dot{v} &= -\frac{Gm_E}{(R_E + h)^2};\end{aligned}\tag{7.3}$$

dabei ist m_E die Masse der Erde, R_E der Radius der Erdkugel und G die Gravitationskonstante, eine der fundamentalen Naturkonstanten, mit den Werten

$$\begin{aligned}m_E &= 5.9723 \times 10^{24} \text{ kg}, \\ R_E &= 6.37 \times 10^6 \text{ m}, \\ G &= 6.674 \times 10^{-11} \text{ m}^3 \text{ kg}^{-1} \text{ s}^2.\end{aligned}\tag{7.4}$$

Das entstandene Gleichungssystem ist nichtlinear, explizite Lösungen liegen nicht vor. In Aufg. 7.2 sollen Sie zeigen, dass die Funktion

$$E(h, v) = \frac{1}{2}v^2 - \frac{Gm_E}{R_E + h}\tag{7.5}$$

ein erstes Integral von (7.3) ist. Uns interessiert nur die rechte Halbebene. Verschiedene Niveaus sind in Abb. 7.2 skizziert. Die positiven Niveaus erzeugen dabei unbeschränkte Niveaulinien und die negativen Niveaus beschränkte Niveaulinien. Die Lösungen in den unbeschränkten Ästen der Niveaulinien stehen dabei für Würfe, in denen die Kugel nach unendlich entkommt ($v > 0$), bzw. für Kugeln, die aus dem Unendlichen auf die Erde fallen ($v < 0$). Das Nullniveau $\frac{1}{2}v^2 = Gm_E/(R_E + h)$ trennt dabei die beiden Bereiche. Aus dieser Gleichung können wir ablesen, mit welcher Geschwindigkeit wir eine Kugel (oder Rakete) von der Erdoberfläche abwerfen (abschießen) müssten, damit sie gerade in den Weltraum entkommt. Diese Geschwindigkeit nennt man *Fluchtgeschwindigkeit*:

$$v_{\text{Flucht}} = \sqrt{2 \frac{Gm_E}{R_E + 0}} = 1.2 \times 10^4 \text{ m s}^{-1},$$

das sind ca. $40\,000 \text{ km h}^{-1}$. Die beiden Modelle (7.2) und (7.3) sind also sowohl mathematisch als auch in Bezug auf ihre Aussagen über die Realität grundverschieden, aber trotzdem werden beide benutzt, um senkrechte Würfe zu beschreiben. Wie kann das sein?

Einen ersten Hinweis erhalten wir, wenn wir uns die Werte der Parameter g , m_E , R_E und G anschauen. Dazu schreiben wir die Newton'sche Geschwindigkeitsgleichung etwas um,

$$\dot{v} = -\frac{Gm_E}{R_E^2} \frac{1}{\left(1 + \frac{h}{R_E}\right)^2},$$

und berechnen den Vorfaktor

$$\frac{Gm_E}{R_E^2} = 9.82 \text{ m s}^{-2}.$$

Das ist offenbar gerade die Erdbeschleunigung g . (Die Erdbeschleunigung variiert über die Erdoberfläche zwischen 9.832 m s^{-2} an den Polen und 9.780 m s^{-2} am Äquator. Diese Unterschiede rühren hauptsächlich daher, dass die Erde keine perfekte Kugel, sondern zu den Polen hin abgeplattet ist. Weitere lokale Unterschiede in kleineren Größenordnungen stammen von unterschiedlichen Gesteinsdichten in der Erde.) Ist die Höhe, die die Kugel erreicht, klein im Vergleich zum Erdradius ($h/R_E \ll 1$), haben wir

$$\dot{v} = -\frac{g}{\left(1 + \frac{h}{R_E}\right)^2} \approx -g,$$

das heißt, in diesem Fall unterscheiden sich die Differentialgleichungen im Newton'schen und im Galilei'schen Modell nur geringfügig. Die Hoffnung besteht dann, dass sich auch die Lösungen nur geringfügig unterscheiden. Die Situation stellt sich also folgendermaßen dar: Das Newton'sche Modell ist das Modell mit dem größeren Gültigkeitsbereich. Es ist gültig für alle Höhen. Das Galilei-Modell hat hingegen einen beschränkten Gültigkeitsbereich. Es darf nur angewendet werden, wenn die erreichten Höhen klein sind im Vergleich zum Erdradius. Das Ziel der kommenden Abschnitte ist es,

- die Abweichung der Aussagen der beiden Modelle zu quantifizieren und
- Zwischenmodelle zwischen Galilei und Newton zu entwickeln, die aus einfacheren Gleichungen bestehen als das Newton-Modell, aber auch für größere Höhen h richtige Vorhersagen liefern.

7.1.3 Entdimensionalisierung und Wahl der geeigneten Skalen

In den folgenden Betrachtungen geht es darum, kleine Effekte unter den Tisch fallen zu lassen. Aber was sind kleine Effekte? Eine Größe allein kann nie groß oder klein sein, sondern nur im Vergleich zu einer anderen Größe der gleichen Dimension. Hier wird also jetzt die Entdimensionalisierung wesentlich (bis hierhin haben wir sie hauptsächlich benutzt, um die Anzahl der Parameter zu reduzieren). Im Gegensatz zu den vorherigen Betrachtungen, in denen die Wahl der benutzten Skalen mehr oder weniger willkürlich war, wird sie nun entscheidend sein, wie wir gleich sehen werden. Wir nutzen die Bezeichnungen $x = h/\bar{h}$, $y = v/\bar{v}$ und $\tau = t/\bar{t}$ und erhalten dann die folgenden entdimensionalisierten Systeme:

$$\begin{aligned} \dot{x} &= \frac{\bar{t}\bar{v}}{\bar{h}} y \\ \dot{y} &= \begin{cases} -\frac{\bar{t}g}{\bar{v}} & \text{Galilei} \\ -\frac{\bar{t}Gm_E}{\bar{v}R_E^2} \frac{1}{\left(1 + \frac{\bar{h}}{R_E}x\right)^2} & \text{Newton.} \end{cases} \end{aligned} \quad (7.6)$$

1. Möglichkeit: Möglichst einfache Gleichungen Wir folgen zunächst einmal der Maxime, möglichst die Parameter aus den „schwierigen“ Gleichungen zu beseitigen. Dazu wählen wir $\bar{h} = R_E$ und stellen die Forderungen

$$\frac{\bar{t}\bar{v}}{\bar{h}} = 1 \quad \text{und} \quad \frac{\bar{t}Gm_E}{\bar{v}R_E^2} = 1.$$

Das führt auf die beiden Skalen

$$\bar{v} = \sqrt{\frac{Gm_E}{R_E}} \quad \text{und} \quad \bar{t} = \frac{\bar{h}}{\bar{v}} = \sqrt{\frac{R_E^3}{Gm_E}}.$$

Diese Skalen sind gut zu interpretieren: Die Längenskala ist der Erdradius, die Geschwindigkeitsskala ist bis auf einen Faktor $\sqrt{2}$ die Fluchtgeschwindigkeit und die Zeitskala gibt die Zeit an, die benötigt wird, um mit dieser Geschwindigkeit die Strecke eines Erdradius' zurückzulegen. Das Newton-System ist nun parameterfrei:

$$\begin{aligned} \dot{x} &= y \\ \dot{y} &= -\frac{1}{(1+x)^2}. \end{aligned} \quad (7.7)$$

Wie kommt jetzt die Kleinheit ins Spiel? Über die Anfangswerte! Wir betrachten einen Körper, der von der Erdoberfläche senkrecht nach oben mit einer Geschwindigkeit v_0 abgeworfen oder -geschossen wird. Das heißt, wir fügen die Anfangswerte $x_0 = 0$ und $y_0 = v_0/\bar{v}$ zu (7.7) hinzu. Stellen wir uns tatsächlich einen von Menschenhand geworfenen Körper vor, ist eine Anfangsgeschwindigkeit von 10 m s^{-1} realistisch und wir erhalten als dimensionsloses

Verhältnis

$$\varepsilon = \frac{v_0}{\bar{v}} = 8,9 \times 10^{-4}, \quad (7.8)$$

also circa ein Tausendstel. Folgen wir der Strategie, kleine Einflüsse unter den Tisch fallen zu lassen, würden wir also die Anfangsgeschwindigkeit vernachlässigen und die Anfangswerte $x_0 = 0$ und $y_0 = 0$ betrachten. Die Änderungsrate der Geschwindigkeit ist aber negativ. Damit werden für positive Zeiten τ sowohl die Geschwindigkeiten als auch die Höhen negativ, der Körper würde also direkt unter die Erdoberfläche verschwinden (dafür müssten wir ihm wohl einen Schacht graben). Von der nach oben gerichteten Wurfbewegung ist in dieser Annäherung überhaupt nichts mehr zu sehen. Wie kann das sein? Betrachten wir die konkreten Werte der benutzten Höhen und Zeitskalen,

$$\begin{aligned} \bar{h} &= R_E = 6370 \text{ km} \\ \bar{t} &= \sqrt{\frac{R_E^3}{Gm_E}} = 805 \text{ s}, \end{aligned}$$

so sehen wir, dass das im Verhältnis zur erwarteten Bewegung zu große Skalen sind. Aus der Galilei'schen Bewegungsgleichung (7.1) erwarten wir ja, dass zu den Anfangsdaten $h_0 = 0$ und $v_0 = 10 \text{ m s}^{-1}$ der Körper eine maximale Höhe von ca. 30 m erreicht und ca. 5 s in der Luft bleibt. Diese Höhe und diese Zeitspanne sind aber im Vergleich zur gewählten Höhen- bzw. Zeitskala verschwindend klein. Wir müssen also Skalen wählen, die mit der erwarteten maximalen Höhe, der erwarteten Flugzeit und der gewählten Anfangsgeschwindigkeit vergleichbar sind.

2. Möglichkeit: Zur betrachteten Situation passende Skalen Wir fixieren jetzt von Beginn an die Geschwindigkeitsskala $\bar{v} = v_0$ und wünschen uns wieder

$$\frac{\bar{t}\bar{v}}{\bar{h}} = 1 \quad \text{und} \quad \frac{\bar{t}Gm_E}{\bar{v}R_E^2} = 1$$

(dass der Wunsch $\bar{h}/R_E = 1$ in der betrachteten Situation nicht sinnvoll ist, haben wir eben schon gesehen), was auf die Skalen

$$\bar{t} = \frac{v_0 R_E^2}{Gm_E} = 1.0 \text{ s} \quad \text{und} \quad \bar{h} = \frac{v_0^2 R_E^2}{Gm_E} = 10 \text{ m} \quad (7.9)$$

führt. Mit dem dimensionslosen Parameter

$$\varepsilon = \frac{\bar{h}}{R_E} = \frac{v_0^2 R_E}{Gm_E} = 1.6 \times 10^{-6} \quad (7.10)$$

erhalten wir das dimensionslose System

$$\begin{aligned}\dot{x} &= y \\ \dot{y} &= -\frac{1}{(1 + \varepsilon x)^2}\end{aligned}\tag{7.11}$$

zusammen mit den Anfangsbedingungen $x_0 = 0$ und $y_0 = 1$. Wenn wir jetzt kleine Effekte vernachlässigen, also $\varepsilon = 0$ setzen, erhalten wir das mit denselben Skalen entdimensionalisierte Galilei-System $\dot{x} = y$, $\dot{y} = -1$.

7.1.4 Die Idee der asymptotischen Entwicklung

Wenn wir in (7.10) die Anfangsgeschwindigkeit v_0 vergrößern, wächst auch der Parameter ε und die Abweichung zwischen dem Newton- und dem Galilei-Modell wird größer. Unser Ziel ist es nun, ein Modell aufzustellen, das gewissermaßen zwischen Galilei und Newton steht, also nicht nur für sehr kleine $\varepsilon \approx 1/1000$, sondern auch noch für kleine Parameter $\varepsilon \approx 1/100$ oder $\varepsilon \approx 1/10$ eine gute Approximation an das Newton-System liefert, aber immer noch einfach zu lösen ist. Wir lassen uns dabei von der Idee der Taylor-Entwicklung einer Funktion um eine Stelle x leiten: Für hinreichend häufig differenzierbare Funktionen f und für $|\varepsilon| \ll 1$ haben wir

$$\begin{aligned}f(x + \varepsilon) &\approx f(x) \\ f(x + \varepsilon) &\approx f(x) + \varepsilon f'(x) \\ f(x + \varepsilon) &\approx f(x) + \varepsilon f'(x) + \frac{\varepsilon^2}{2} f''(x),\end{aligned}$$

wobei der Fehler der nullten Approximation von der Ordnung ε , der Fehler der ersten Approximation von der Ordnung ε^2 und der der zweiten Approximation von der Ordnung ε^3 ist. Wenn nun ε zwar klein, aber nicht sehr klein ist, ist doch dann ε^2 und ε^3 wiederum sehr klein.

Die Idee der asymptotischen Entwicklung besteht nun darin, dass wir nicht nur ein System mit einem festen Parameter ε und einer Lösung betrachten, sondern direkt eine ganze, durch ε parametrisierte Schar von Systemen mit der zugehörigen Lösungsschar. Ausführlicher: Für $0 \leq \varepsilon < 1$ und eine Schar von Anfangsdaten $(x^{\varepsilon,0}, y^{\varepsilon,0})$ betrachten wir die zugehörigen Lösungen $(x^\varepsilon, y^\varepsilon)$ der Systemschar (7.11). Die durchaus forsche Grundannahme der asymptotischen Entwicklung besteht nun darin, dass diese Lösungsschar nach ε entwickelbar ist, also

$$\begin{aligned}x^\varepsilon &= x_0 + \varepsilon x_1 + \varepsilon^2 x_2 + \dots \\ y^\varepsilon &= y_0 + \varepsilon y_1 + \varepsilon^2 y_2 + \dots\end{aligned}\tag{7.12}$$

Dabei sind die Koeffizienten x_k, y_k Funktionen, die unabhängig von ε sind. Die Erwartung ist dann, dass der Fehler der Approximation k -ter Ordnung von der Ordnung ε^{k+1} ist. Das sind natürlich viele Erwartungen und Hoffnungen, die allesamt beweisbedürftig sind. Wir lassen

uns aber erst einmal auf das Spiel ein und schauen, wohin uns das führt. Wir setzen diesen Ansatz (7.12) in das Differentialgleichungssystem (7.11) ein. Auf der linken Seite gehen wir dabei davon aus, dass die Differentiation nach der Zeit summandenweise ausgeführt werden kann (wir sind momentan in der Welt der Heuristik, nicht der rigorosen Mathematik). Für die rechte Seite der zweiten Differentialgleichung nutzen wir die Entwicklung

$$-1/(1+s)^2 = -1 + 2s - 3s^2 + 4s^3 \mp \dots = \sum_{n=0}^{\infty} (-1)^{n+1} (n+1) s^n \quad \text{für } |s| < 1 \quad (7.13)$$

(siehe Aufg. 7.3). Damit erhalten wir:

$$\begin{aligned} \dot{x}_0 + \varepsilon \dot{x}_1 + \varepsilon^2 \dot{x}_2 + \dots &= y_0 + \varepsilon y_1 + \varepsilon^2 y_2 + \dots \\ \dot{y}_0 + \varepsilon \dot{y}_1 + \varepsilon^2 \dot{y}_2 + \dots &= -1 + 2\varepsilon(x_0 + \varepsilon x_1 + \varepsilon^2 x_2 + \dots) - 3\varepsilon^2(x_0 + \varepsilon x_1 + \dots)^2 \pm \dots \end{aligned}$$

Diese Gleichungen sollen für alle $0 \leq \varepsilon < 1$ gelten. Das ist nur möglich, wenn sie in jeder Ordnung von ε einzeln stimmen. Wir sortieren also nach Ordnungen von ε . Für das Beispiel bleiben wir bei unseren alten Anfangswerten, die unabhängig von ε sind:

$$x^{\varepsilon,0} = x^{\varepsilon}(0) = 0 \quad \text{und} \quad y^{\varepsilon,0} = y^{\varepsilon}(0) = 1.$$

In nullter Ordnung ergibt sich gerade das bekannte Galileo-System

$$\dot{x}_0 = y_0, \quad \dot{y}_0 = -1 \quad (7.14)$$

mit den Anfangswerten $x_0(0) = 0$, $y_0(0) = 1$. Die Lösung ist gegeben durch

$$x_0(\tau) = -\frac{1}{2}\tau^2 + \tau, \quad y_0(\tau) = -\tau + 1. \quad (7.15)$$

Im asymptotischen Jargon kann man das jetzt so bezeichnen: Das Galileo-System ist der asymptotische Grenzwert des Newton-Systems im Limes v_0/v_{grenz} gegen 0. Jetzt kommen wir zu den Termen in der Ordnung ε^1 :

$$\dot{x}_1 = y_1, \quad \dot{y}_1 = 2x_0 \quad (7.16)$$

mit den Anfangswerten $x_1(0) = y_1(0) = 0$. Auch dieses Problem ist einfach zu lösen, die Funktion x_0 kennen wir ja bereits. Wir erhalten:

$$\begin{aligned} y_1(\tau) &= \int_0^{\tau} 2x_0(s) ds = -\frac{1}{3}\tau^3 + \tau^2 \\ x_1(\tau) &= \int_0^{\tau} y_1(s) ds = -\frac{1}{12}\tau^4 + \frac{1}{3}\tau^3 \end{aligned}$$

In der Ordnung ε^2 erhalten wir die Gleichungen

$$\dot{x}_2 = y_2, \quad \dot{y}_2 = 2x_1 - 3x_0$$

mit verschwindenden Anfangsdaten $x_2(0) = y_2(0) = 0$. Auch diese Gleichungen sind durch simple Integration zu lösen:

$$y_2(\tau) = \int_0^\tau 2x_1(s) - 3x_0^2(s) ds = -\tau^3 + \frac{11}{12}\tau^4 - \frac{11}{60}\tau^5$$

$$x_2(\tau) = \int_0^\tau y_2(s) ds = -\frac{1}{4}\tau^4 + \frac{11}{60}\tau^5 - \frac{11}{360}\tau^6$$

Es besteht kein Zweifel, dass man dieses Verfahren sukzessive bis in beliebige Ordnungen fortsetzen kann. Die Rechnungen werden dann zwar umfangreich, sind aber immer elementar. Die asymptotische Entwicklung bis zur Ordnung zwei ist dann durch

$$x_\varepsilon^{\text{II}}(\tau) = x_0(\tau) + \varepsilon x_1(\tau) + \varepsilon^2 x_2(\tau), \quad y_\varepsilon^{\text{II}}(\tau) = y_0(\tau) + \varepsilon y_1(\tau) + \varepsilon^2 y_2(\tau) \quad (7.17)$$

gegeben. Wir wollen zunächst einmal numerisch untersuchen, wie sich die einzelnen Approximationen und die vollen Lösungen bei verschiedenen Werten von ε zueinander verhalten. Da wir keine expliziten Lösungen des Newton-Systems kennen, müssen wir diese wiederum mit dem Euler-Verfahren approximieren. In Abb. 7.3 sind die Approximationen und

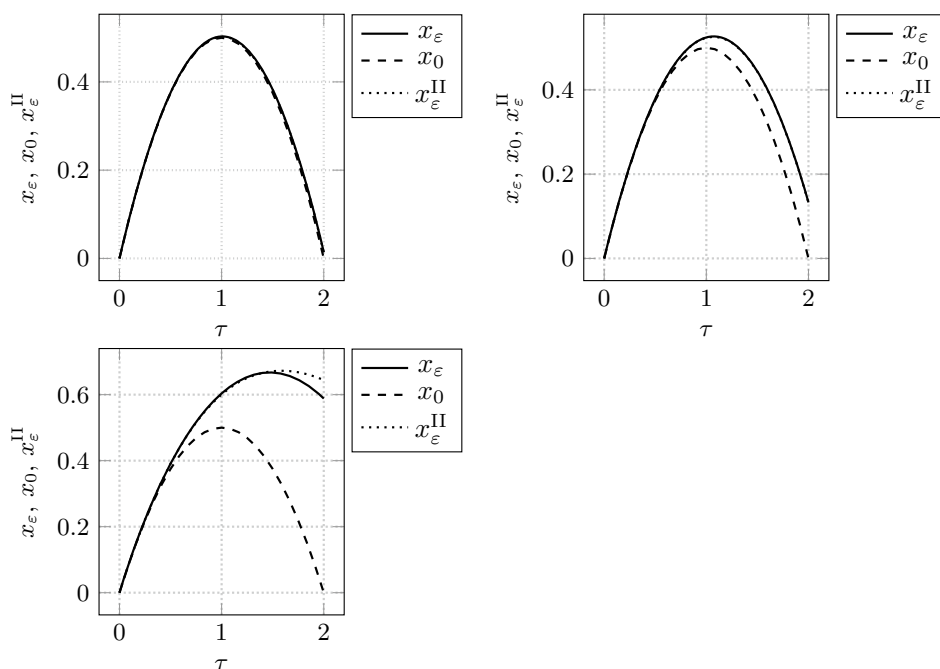


Abb. 7.3 Vergleich der numerisch berechneten Lösung des Newton-Systems (Schrittweite 0.001) mit den Approximationen nullter und zweiter Ordnung für $\varepsilon = 0.01, 0.1$ und 0.5

die volle Lösung für $\varepsilon = 0.01$, $\varepsilon = 0.1$ und $\varepsilon = 0.5$ dargestellt. Während für $\varepsilon = 0.01$ $x_\varepsilon^{\text{II}}$ und x_ε nicht unterscheidbar sind, sind für $\varepsilon = 0.5$ deutliche Unterschiede zu erkennen. Aus unseren numerischen Versuchen sehen wir, dass für $\varepsilon \ll 1$ kein erkennbarer Unterschied zwischen den Approximationen nullter und zweiter Ordnung und der vollen Lösung besteht. Für größer werdendes ε wird zunächst die nullte und später auch die zweite Approximation ungenauer. In allen Fällen nimmt der Approximationsfehler mit der Zeit zu.

7.1.5 Beweis der Approximationseigenschaft

Als Mathematiker möchte man gerne einen rigorosen Beweis haben, dass die asymptotischen Verfahren funktionieren, nicht nur numerische Evidenz. Zunächst müssen wir uns darüber klar werden, wie so eine asymptotische Aussage überhaupt aussehen soll.

Umgangssprachlich haben wir gesagt, dass die Entwicklung bis zur Ordnung k die Lösung auch bis zur Ordnung k approximieren soll, der Fehler also von der Ordnung $k+1$ sein soll. Was soll das präzise heißen? Es soll das Folgende bedeuten: Es gibt Konstanten $\hat{\varepsilon}_k$, c_k und T_k , sodass für alle $0 \leq \varepsilon < \hat{\varepsilon}_k$ und $0 \leq \tau \leq T_k$ die folgende Abschätzung gilt:

$$|x_\varepsilon(\tau) - x_\varepsilon^k(\tau)| \leq c_k \varepsilon^{k+1}$$

In der Regel wird die Konstante c_k nicht nur von k abhängen, sondern auch von der Konstante T_k . Je länger die Zeit läuft, desto mehr häufen sich die Fehler an und desto schlechter wird die Abschätzung. Geht man also von einer Konstanten T_k zu einer auch erlaubten Konstante $\tilde{T}_k < T_k$ über, so kann man häufig auch eine kleinere Konstante $\tilde{c}_k$ finden, mit der die Abschätzung gültig ist. Wir werden die Abschätzung für den einfachsten Fall $k = 0$ beweisen. Eine Übertragung auf den Fall $k = 1$ wird in Aufg. 7.7 gefordert. Die Schwierigkeit des Beweises besteht darin, dass von der Lösungsschar x_ε lediglich bekannt ist, dass sie die Differentialgleichungen und die Anfangsbedingungen erfüllt, aber keine explizite Lösung vorliegt. Allein aus den genannten Tatsachen muss die Abschätzung gefolgert werden.

Wir führen als neue Größe die Differenz von x_ε und x_0 ein: $w_\varepsilon = x_\varepsilon - x_0$ und $z_\varepsilon = \dot{w}_\varepsilon = \dot{x}_\varepsilon - \dot{x}_0$. Für die Zeit $\tau = 0$ gilt dann $w_\varepsilon(0) = z_\varepsilon(0) = 0$. Wir rechnen nach, welche Differentialgleichung z_ε löst, und versuchen mit Hilfe der Differentialgleichungen für w_ε und z_ε den Fehler zwischen den Höhen w_ε abzuschätzen. Aus den Systemen (7.11) und (7.14) folgt:

$$\dot{z}_\varepsilon = -\frac{1}{(1 + \varepsilon x_\varepsilon)^2} + 1 \quad (7.18)$$

Wir wollen zeigen, dass dieser Ausdruck auf einem festen Zeitintervall von der Ordnung ε ist. Dazu benutzen wir, wie die Approximation der nullten Ordnung entstanden ist. Wir nutzen den Mittelwertsatz der Differentialrechnung: Es gibt eine Stelle $\theta = \theta(s) \in (0, 1)$, sodass

$$f(s) := -\frac{1}{(1+s)^2} = f(0) + sf'(\theta s) = -1 + \frac{2}{(1+\theta s)^3} \cdot s.$$

Wir setzen jetzt für s den Ausdruck $\varepsilon x_\varepsilon(\tau)$ ein und erhalten

$$\dot{z}_\varepsilon(\tau) = \frac{2}{(1 + \theta \varepsilon x_\varepsilon(\tau))^3} \cdot \varepsilon x_\varepsilon(\tau)$$

mit $0 < \theta = \theta(\varepsilon, \tau) < 1$. (Die Notation $\theta(\varepsilon, \tau)$ soll ausdrücken, dass die Zwischenstelle θ sowohl von ε als auch von τ abhängt.) Allerdings benötigen wir erst eine grobe Abschätzung für x_ε : Es gilt $x_\varepsilon(0) = 0$ und $\dot{x}_\varepsilon(0) = y_\varepsilon(0) = 1$. Damit gibt es ein maximales Intervall $[0, \tau_\varepsilon)$, auf dem die Abschätzung $0 \leq x_\varepsilon(\tau)$ gültig ist. Für dieses Intervall können wir abschätzen:

$$-1 \leq \dot{y}_\varepsilon(\tau) = -\frac{1}{(1 + \varepsilon x_\varepsilon(\tau))^2} < 0$$

Durch Integration und wegen $y_\varepsilon(0) = 1$ schließen wir aus dieser Ungleichung für das Intervall $[0, \tau_\varepsilon)$

$$1 - \tau \leq y_\varepsilon(\tau) = y_\varepsilon(0) + \int_0^\tau \dot{y}_\varepsilon(s) ds \leq 1.$$

Und mit einer weiteren Integration unter Berücksichtigung des Anfangswerts $x_\varepsilon(0) = 0$ erhalten wir

$$\tau - \frac{1}{2}\tau^2 \leq x_\varepsilon(\tau) = x_\varepsilon(0) + \int_0^\tau y_\varepsilon(s) ds \leq \tau \quad (7.19)$$

für alle $\tau \in [0, \tau_\varepsilon)$.

Diese Abschätzung nutzen wir nun, um eine gute Abschätzung für den Fehler w_ε zu gewinnen: Für alle $\tau \in [0, \tau_\varepsilon)$ gilt

$$\begin{aligned} 0 \leq w_\varepsilon(\tau) &= \int_0^\tau z_\varepsilon(s) ds = \int_0^\tau \int_0^s \dot{z}_\varepsilon(\sigma) d\sigma ds \\ &= 2\varepsilon \int_0^\tau \int_0^s \frac{x_\varepsilon(\sigma)}{(1 + \theta \varepsilon x_\varepsilon(\sigma))^3} d\sigma ds \\ &\leq 2\varepsilon \int_0^\tau \int_0^s \frac{\sigma}{1} d\sigma ds = 2\varepsilon \frac{1}{6} \tau^3 \\ &\leq \frac{\varepsilon \tau_\varepsilon^3}{3}. \end{aligned}$$

Bis jetzt haben wir nur die rechte Ungleichung aus (7.19) benutzt. Um nun einen gleichmäßigen Gültigkeitsbereich der Abschätzung zu gewinnen, benötigen wir die linke Ungleichung aus (7.19): Weil $0 \leq \tau - \frac{1}{2}\tau^2$ für alle $0 \leq \tau \leq 2$, folgern wir $\tau_\varepsilon \geq 2$ für alle $\varepsilon \geq 0$. Damit haben wir, was wir wollten: Für alle $\varepsilon \geq 0$ und alle $0 \leq \tau \leq 2$ gilt

$$|x_\varepsilon(\tau) - x_0(\tau)| \leq \frac{8}{3} \varepsilon. \quad (7.20)$$

In Aufg. 7.5 soll gezeigt werden, dass sich der Fehler zwischen x_ε und der Approximation erster Ordnung tatsächlich gegen ε^2 abschätzen lässt.

Anwendung auf die Situation In einem letzten Schritt beziehen wir diese Abschätzung nun wieder auf unsere Ausgangssituation: Wie gut approximiert die Lösung des Galilei-Modells die Lösung des Newton-Modells bei einer Abwurfgeschwindigkeit von $v_0 = 10 \text{ m s}^{-1}$? Wir erinnern an die konkreten Skalen und den konkreten Wert von ε in dieser Situation, siehe (7.9) und (7.10):

$$\bar{t} = 1 \text{ s}, \quad \bar{h} = 10 \text{ m} \quad \text{und} \quad \hat{\varepsilon} = 1.6 \times 10^{-6}.$$

Aus (7.20) erhalten wir die Fehlerschranke

$$|h_N(t) - h_G(t)| \leq \bar{h} |x_{\hat{\varepsilon}}(t/\bar{t}) - x_0(t/\bar{t})| \leq 10 \text{ m} \cdot \frac{8}{3} \cdot 1.6 \times 10^{-6} = 4.3 \times 10^{-5} \text{ m}$$

für alle Zeiten $t \leq 2 \text{ s}$. Nach zwei Sekunden trifft der Ball aber nach dem Galilei-Modell gerade wieder auf den Boden und beide Modelle gelten danach nicht mehr. Das Galilei-Modell liefert also unter diesen Umständen praktisch dieselbe Vorhersage wie das Newton-Modell.

7.2 Reaktionskinetik, Enzyme und die Methode des Asymptotic Matching

7.2.1 Das grundlegende Modell für chemische Reaktionen: Reaktionskinetik

Chemische Stoffe bestehen aus kleinen Teilen, den Molekülen, die wiederum aus Atomen zusammengesetzt sind. Bei chemischen Reaktionen reagieren die Moleküle von in der Regel unterschiedlichen chemischen Stoffen miteinander und bilden dabei Moleküle irgendwelcher anderer Stoffe. Wir notieren ein paar abstrakte Beispiele:

1. Beispiel: $A \longrightarrow B$ Der Stoff A zerfällt in den Stoff B, genauer: Moleküle des Stoffes A zerfallen in Moleküle des Stoffes B.
2. Beispiel: $A + B \longrightarrow C$ Moleküle des Stoffes A und des Stoffes B verbinden sich zu Molekülen des Stoffes C.
3. Beispiel: $A + A \longrightarrow A_2$ Zwei Moleküle des Stoffes A schließen sich zu einem Molekül zusammen, der resultierende Stoff wird A_2 genannt.

Den Prozess im ersten Beispiel nennt man einen Prozess erster Ordnung: Es reicht, dass ein Molekül vorliegt, damit der Prozess ablaufen kann. Die Prozesse im zweiten und dritten Beispiel nennt man Prozesse zweiter Ordnung: hier müssen sich zwei Moleküle hinreichend nah kommen, damit der Prozess ablaufen kann. Prozesse dritter Ordnung spielen in der Regel keine Rolle, denn hierfür müssten drei Moleküle sich zusammen in einem hinreichend kleinen Raumbereich aufhalten, was so gut wie nie vorkommt. Wir stellen uns nun vor,

dass sich die Moleküle in einer Lösung zufällig hin und her bewegen und in der gesamten Lösung die gleiche Konzentration aufweisen. Das grundlegende mathematische Modell für diese Situation ist sehr einfach. Es sei $a = a(t)$ die Konzentration des Stoffes A zur Zeit t in der Lösung,

$$a(t) = \frac{\text{Anzahl der Moleküle des Stoffes A zur Zeit } t}{\text{Volumen der Lösung}}, \quad (7.21)$$

$b = b(t)$ die Konzentration des Stoffes B usw. Die zeitlichen Entwicklungen werden durch Differentialgleichungssysteme beschrieben.

Prozesse erster Ordnung Ein Prozess erster Ordnung $A \longrightarrow B + C$ führt auf das System

$$\begin{aligned}\dot{a} &= -ka \\ \dot{b} &= +ka \\ \dot{c} &= +ka.\end{aligned}$$

Dabei gibt der Koeffizient k die Reaktionsgeschwindigkeit an und hat die Dimension $1/T$. Da sich Moleküle ineinander umwandeln können, haben alle Konzentrationen dieselbe Dimension, nämlich Anzahl pro Volumen. Da die Anzahl der Moleküle in der Regel sehr hoch ist, wird die Anzahl meist in der Einheit mol angegeben, wobei 1 mol eines Stoffes $6.02214076 \times 10^{23}$ Teilchen also ca. 602 Trilliarden Teilchen, enthält.

Prozesse zweiter Ordnung Prozesse zweiter Ordnung führen zu quadratischen Termen in den Differentialgleichungen. Dahinter steckt wie schon bei den Konkurrenz- und Räuber-Beute-Termen die Überlegung, dass die Anzahl der Annäherungen zwischen zwei Teilchen, welche dann zu einer Reaktion führen, proportional zu beiden Stoffkonzentrationen und damit proportional zum Produkt der beiden Stoffkonzentrationen ist. Die Reaktion $A + B \longrightarrow C$ führt auf das System

$$\begin{aligned}\dot{a} &= -kab \\ \dot{b} &= -kab \\ \dot{c} &= +kab.\end{aligned}$$

Wir können dieses Bildungsprinzip auch auf komplexe Reaktionsketten anwenden. Der Prozess $E + S \xrightleftharpoons[k_{-1}]{k_1} X \xrightarrow{k_2} Y \xrightarrow{k_3} E + P$ führt auf das Differentialgleichungssystem

$$\begin{aligned}\dot{e} &= -k_1es + k_{-1}x + k_3y \\ \dot{s} &= -k_1es + k_{-1}x \\ \dot{x} &= +k_1es - (k_{-1} + k_2)x \\ \dot{y} &= k_2x - k_3y \\ \dot{p} &= k_3y.\end{aligned}$$

Ein Doppelpfeil $\rightleftharpoons$ deutet dabei an, dass der Prozess auch in der Rückrichtung ablaufen kann. Die Geschwindigkeit der Rückrichtung steht dann unter dem Pfeil.

7.2.2 Enzymatische Reaktionen – zwei unterschiedliche Zeitskalen und das Verfahren des Asymptotic Matching

Eine enzymatische Reaktion ist eine Reaktion, in der ein Stoff, das Enzym, eine Reaktionskette in Gang setzt, aber am Ende der Reaktionskette unverändert heraustritt. Eine der einfachsten enzymatischen Reaktionen ist die folgende, die von Michaelis und Menten als erstes aufgestellt wurde, siehe [63]. In die Reaktion gehen das Substrat S und das Enzym E ein. Diese bilden den Substratenzymkomplex SE , der wiederum in das Produkt P und das Enzym E zerfällt: $S + E \xrightleftharpoons[k_{-1}]{k_1} SE \xrightarrow{k_2} P + E$. Wir bezeichnen den Substrat-Enzym-Komplex mit C und die zugehörigen Konzentrationen mit s, e, c und p . Das zugehörige Differentialgleichungssystem lautet:

$$\begin{aligned}\dot{s} &= -k_1 s e + k_{-1} c \\ \dot{e} &= -k_1 s e + (k_{-1} + k_2) c \\ \dot{c} &= k_1 s e - (k_{-1} + k_2) c \\ \dot{p} &= k_2 c\end{aligned}\tag{7.22}$$

Das Ganze versehen wir mit Anfangswerten

$$s(0) = s_0 > 0, \quad e(0) = e_0 > 0 \quad \text{und} \quad c(0) = p(0) = 0.\tag{7.23}$$

Zu Beginn liegen also nur das Substrat und das Enzym vor. Wir starten mit ein paar einfachen Beobachtungen:

- 1) Die Produktkonzentration p ist in dem Differentialgleichungssystem entkoppelt und kann alleine aus der Komplexkonzentration berechnet werden:

$$p(t) = \int_0^t k_2 c(s) ds$$

- 2) Die Gesamtzahl der Enzyme (frei oder im Komplex gebunden) bleibt erhalten:

$$(e + c) = -k_1 s e + (k_{-1} + k_2) c + k_1 s e - (k_{-1} + k_2) c = 0$$

Wir können also eine Variable eliminieren:

$$e(t) + c(t) = e(0) + c(0) = e_0$$

für alle Zeiten t , also $e(t) = e_0 - c(t)$.

Wir können also das Gleichungssystem (7.22) auf zwei Gleichungen reduzieren:

$$\begin{aligned}\dot{s} &= -k_1 s (e_0 - c) + k_{-1} c \\ \dot{c} &= k_1 s (e_0 - c) - (k_{-1} + k_2) c\end{aligned}\quad (7.24)$$

mit den Anfangswerten $s(0) = s_0$ und $c(0) = 0$. Wir versuchen das qualitative Verhalten durch eine Skizze des Phasenporträts zu klären. Wir erhalten die Nullklinen

$$\begin{aligned}\dot{s} = 0 &\Leftrightarrow c = \frac{k_1 s e_0}{k_{-1} + k_1 s} = e_0 - \frac{k_{-1} e_0}{k_{-1} + k_1 s} \\ \dot{c} = 0 &\Leftrightarrow c = \frac{k_1 s e_0}{k_{-1} + k_2 + k_1 s} = e_0 - \frac{(k_{-1} + k_2) e_0}{k_{-1} + k_2 + k_1 s}.\end{aligned}$$

Das qualitative Phasenporträt mit einer Skizze der Trajektorie durch den Punkt $(s_0, 0)$ ist in Abb. 7.4 dargestellt. Die Substratkonzentration geht streng monoton auf null zurück, der Komplex baut sich zunächst auf und fällt anschließend wieder auf null. Das Substrat wird also komplett in das Produkt umgewandelt. Für den Anwender ist diese Aussage über das qualitative Langzeitverhalten oft nicht ausreichend, sondern Wissen über die genauen zeitlichen Abläufe gewünscht. Wir betrachten nun ein konkretes Beispiel für solch eine einfache enzymatische Reaktion, nämlich die Hydrolyse von Benzol-L-Argine-Ethyl-Ester

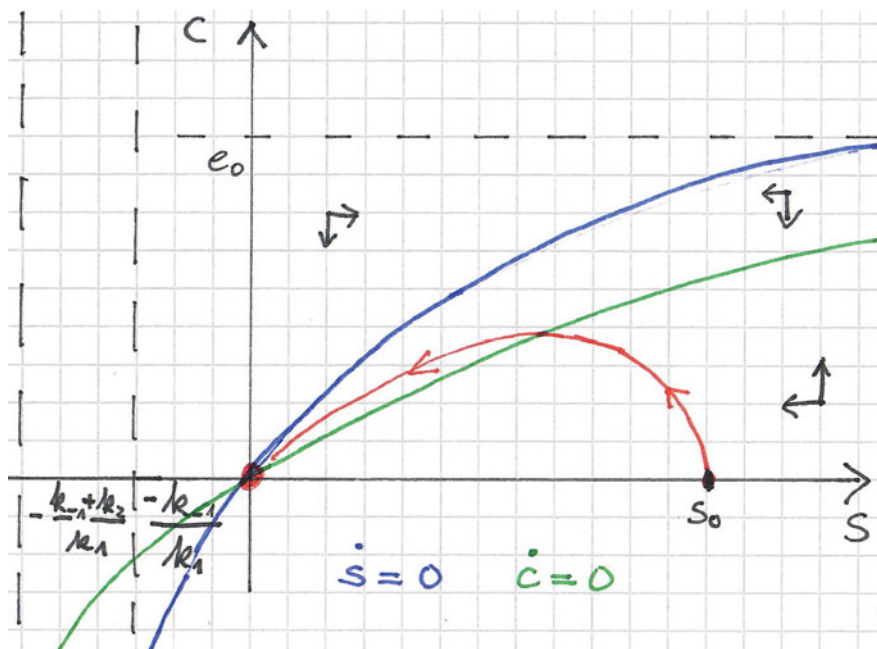


Abb. 7.4 Phasenporträt einer einfachen enzymatischen Reaktion

zu Tryptin. Nach [14] wird diese Reaktion durch folgende Parameterwerte beschrieben:

$$k_1 = 4 \times 10^6 \text{ mol}^{-1} \text{ s}^{-1}, \quad k_{-1} = 25 \text{ s}^{-1} \quad \text{und} \quad k_2 = 15 \text{ s}^{-1}. \quad (7.25)$$

Wir lösen das resultierende System mit den Anfangswerten

$$s_0 = 10 \times 10^{-5} \text{ mol} \quad \text{und} \quad e_0 = 10 \times 10^{-8} \text{ mol} \quad (7.26)$$

numerisch. Dabei lässt sich erkennen, dass in diesem Prozess zwei sehr unterschiedliche Geschwindigkeiten auftreten: Zunächst baut sich sehr schnell, innerhalb von 0.05 s, der Komplex auf, anschließend werden Substrat und Komplex über einen Zeitraum von Hunderten von Sekunden wieder abgebaut (siehe Abb. 7.5). Der Anwender hätte nun gerne eine genauere Beschreibung dieser beiden einzelnen Vorgänge, und zwar nicht nur durch numerische Verfahren. Die Chemiker machen in dieser Situation die Approximation $\dot{c} \approx 0$, die sogenannte *Quasi-Steady-State-Hypothese*. Wir werden uns gleich Gedanken machen, unter welchen Voraussetzungen diese Annahme gerechtfertigt ist und wie man überhaupt auf sie kommen kann, wenn man nicht mit der begnadeten Intuition eines Biochemikers gesegnet ist. Zunächst einmal spielen wir aber durch, wohin uns diese Annahme führt. Interpretieren wir sie im Phasenporträt, bedeutet sie, dass wir uns nur auf einer c -Nullkline bewegen dürfen. Rechnerisch wird die Differentialgleichung für c durch eine algebraische Gleichung ersetzt:

$$0 = k_1 s (e_0 - c) - (k_{-1} + k_2) c \quad (7.27)$$

Wir lösen diese Gleichung nach c auf und erhalten

$$c = \frac{k_1 s e_0}{k_{-1} + k_2 + k_1 s} = \frac{e_0 s}{\frac{k_{-1} + k_2}{k_1} + s} = \frac{e_0 s}{k_m + s}$$

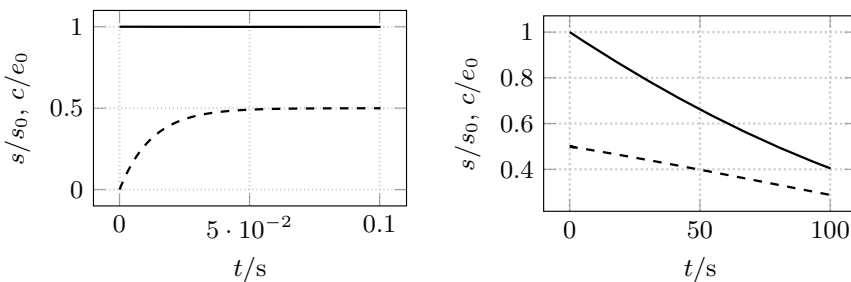


Abb. 7.5 Numerik zur enzymatischen Reaktion mit den Parameterwerten aus (7.25), Schrittweite 1×10^{-4} und den Anfangswerten aus (7.26) bzw. Schrittweite 0.25 und den Anfangswerten $s(0) = 1 \times 10^{-5} \text{ mol}$ und $c(0) = 0.5 \times 10^{-8} \text{ mol}$

mit $k_m = \frac{k_{-1} + k_2}{k_1}$, der sogenannten Michaelis-Konstanten. Wir setzen nun diesen Ausdruck für c in die Differentialgleichung für s ein und erhalten mit etwas Rechnerei eine Differentialgleichung mit getrennten Variablen:

$$\begin{aligned}\dot{s} &= -k_1 s (e_0 - c) + k_{-1} c \\ &= e_0 s \cdot \frac{k_{-1} - k_1 k_m}{k_m + s} \\ &= -s \frac{l}{k_m + s}\end{aligned}\tag{7.28}$$

mit $l = -e_0(k_{-1} - k_1 k_m)$ (das Vorzeichen wurde so gewählt, damit l positiv ist). Mit dem Ansatz von der Trennung der Variablen, dem Anfangswert $s(0) = s_0$ und Integration erhalten wir

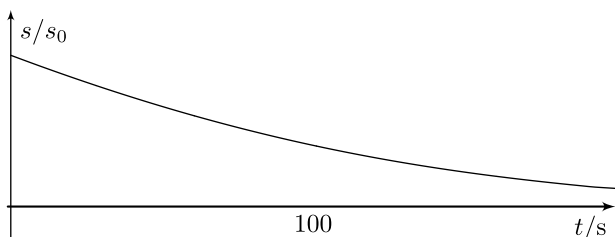
$$s(t) + k_m \ln(s(t)) - s_0 - k_m \ln(s_0) = -lt.$$

Diese Gleichung können wir zwar jetzt nicht explizit nach s auflösen, aber haben immerhin die Zeit in Abhängigkeit von der Konzentration und können uns die Lösungen plotten lassen. Ein Plot mit den Parameterwerten aus (7.25) und den Anfangswerten aus (7.26) ist in Abb. 7.6 dargestellt. Diese Funktion stimmt hervorragend mit der numerisch berechneten Lösung in Abb. 7.5 überein. Die Quasi-Steady-State-Hypothese scheint sich also für diese Werte zu bewähren. Es bleiben die folgenden Fragen:

- 1) Für welche Parameter- und Anfangswerte ist die Approximation $\dot{c} \approx 0$ sinnvoll?
- 2) Durch die Approximation $\dot{c} = 0$ gilt $c(0) = \frac{s_0 e_0}{k_m + s_0} \neq c_0$, das heißt, diese Approximation liefert für den Komplex den falschen Anfangswert und kann somit in der Nähe von $t = 0$ den Komplex nicht richtig beschreiben. Lässt sich dieser Fehler irgendwie reparieren?
- 3) Wieso dürfen wir trotzdem für $s(0)$ den alten Anfangswert s_0 verwenden? Wenn der Anfangswert für c geändert wurde, warum sollte der für s gleich bleiben?

Der Weg der Quasi-Steady-State-Approximation im Phasenporträt ist in Abb. 7.7 veranschaulicht. Die Lösungstrajektorie „klebt“ gewissermaßen an der c -Nullkline fest, vergleiche auch mit Abb. 7.4.

Abb. 7.6 Die Lösung der quasi-steady-state Approximation mit den Parametern und Anfangswerten aus (7.25) bzw. (7.26) und den resultierenden Konstanten $k_m = 10^{-5}$ mol und $l = 1.5 \times 10^{-7}$ mol s⁻¹



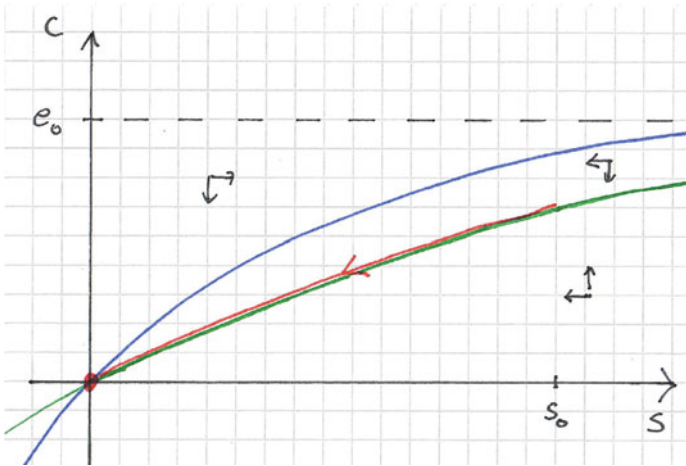


Abb. 7.7 Die Lösungstrajektorie (rot) der Quasi-Steady-State-Approximation im Phasenporträt, c -Nullkline in grün, s -Nullkline in blau

7.2.3 Entdimensionalisierung und asymptotische Betrachtungen

Wir versuchen asymptotischen Methoden anzuwenden. Dazu müssen wir zunächst einmal entdimensionalisieren. Mit den Bezeichnungen

$$\tau = \frac{t}{\bar{t}}, \quad u = \frac{s}{\bar{s}} \quad \text{und} \quad v = \frac{c}{\bar{c}}$$

erhalten wir das entdimensionalisierte System

$$\begin{aligned} \dot{u} &= -\bar{t}k_1\bar{c}u\left(\frac{e_0}{\bar{c}} - v\right) + \frac{\bar{t}k_{-1}\bar{c}}{\bar{s}}v \\ \dot{v} &= \bar{t}k_1\bar{s}u\left(\frac{e_0}{\bar{c}} - v\right) - \bar{t}(k_2 + k_{-1})v \end{aligned} \quad (7.29)$$

mit den Anfangswerten

$$u(0) = \frac{s_0}{\bar{s}} \quad \text{und} \quad v(0) = 0.$$

Die Skalen für die Stoffmengen sind kanonisch $\bar{s} = s_0$ und $\bar{c} = e_0$. Wir setzen diese Wahl in (7.29) ein:

$$\begin{aligned} \dot{u} &= -\bar{t}k_1e_0u(1-v) + \frac{\bar{t}k_{-1}e_0}{s_0}v \\ \dot{v} &= \bar{t}k_1s_0u(1-v) - \bar{t}(k_2 + k_{-1})v. \end{aligned}$$

Welche Zeitskala ergibt jetzt Sinn? Es gibt zunächst einmal vier Möglichkeiten:

$$\bar{t} = \frac{1}{k_1 e_0}, \quad \bar{t} = \frac{s_0}{k_{-1} e_0}, \quad \bar{t} = \frac{1}{k_1 s_0} \quad \text{und} \quad \bar{t} = \frac{1}{k_2 + k_{-1}}.$$

Die große Zeitskala Wir beginnen, ein bisschen auf gut Glück, mit dem ersten Vorschlag $\bar{t} = \frac{1}{k_1 e_0}$. Für unsere Beispielreaktion mit den Parameter- und Anfangswerten (7.25) und (7.26) hat diese Skala den Wert $\bar{t} = 25$ s. Das ist eine plausible Zeiteinheit, wir erwarten aus unseren numerischen Erfahrungen ja einen Prozess, der ca. 200 s dauert. Damit ergibt sich

$$\begin{aligned} \dot{u} &= -u(1-v) + \frac{k_{-1}}{k_1 s_0} v \\ \dot{v} &= \frac{s_0}{e_0} u(1-v) - \frac{k_2 + k_{-1}}{k_1 e_0} (k_2 + k_{-1}) v \\ &= \frac{s_0}{e_0} \left(u(1-v) - \frac{k_m}{s_0} v \right). \end{aligned}$$

Mit den gegebenen Werten erhalten wir für die dimensionslosen Parameterkombinationen

$$\frac{s_0}{e_0} = 10^3 \gg 1, \quad \frac{k_{-1}}{k_1 s_0} =: k_l = \frac{5}{8} \quad \text{sowie} \quad \frac{k_m}{s_0} = 1.$$

Die erste Kombination ist also groß, die anderen sind von der Größenordnung 1. Wir definieren $\varepsilon = \frac{e_0}{s_0} \ll 1$, schreiben die zweite Gleichung um und erhalten schlussendlich

$$\begin{aligned} \dot{u} &= -u(1-v) + k_l v \\ \varepsilon \dot{v} &= u(1-v) - \frac{k_m}{s_0} v. \end{aligned} \tag{7.30}$$

Wir setzen jetzt eine asymptotische Entwicklung für $\varepsilon \ll 1$ an und erhalten in der nullten Ordnung

$$\begin{aligned} \dot{u}_0 &= -u_0(1-v_0) + k_l v_0 \\ 0 &= u_0(1-v_0) - \frac{k_m}{s_0} v_0. \end{aligned}$$

Damit können wir die Quasi-Steady-State-Approximation als asymptotischen Limes für $\varepsilon = e_0/s_0 \ll 1$ identifizieren. Der Erfolg hängt also davon ab, dass zu Beginn mehrere Größenordnungen mehr Substrat als Enzym vorhanden sind. Es ist auch klar, warum das Anfangsdatum für den Komplex nicht realisiert werden kann: Aus einer Differentialgleichung ist eine algebraische Gleichung geworden. In der Sprache der Physiker ist durch den Limes der Freiheitsgrad für den Komplex verloren gegangen. Mathematisch liegt das daran, dass der kleine Parameter vor der Ableitung aufgetaucht ist. Der Mathematiker spricht von einem *singulären asymptotischen Limes*.

Die schnelle Zeitskala Gibt es denn eine Möglichkeit, auch den Aufbau des Komplexes in einer asymptotischen Betrachtung sichtbar zu machen? Die gewählte Zeitskala $\bar{t} = 25$ s

scheint dafür zu groß zu sein. Wir erwarten ja, dass der Aufbau nur Bruchteile von Sekunden benötigt. Was passiert nun, wenn wir stattdessen eine kurze Zeitskala benutzen, z.B. $\bar{t}_2 = \varepsilon \bar{t} = 0.025 \text{ s}$? Das läuft auf neue entdimensionalisierte Größen heraus, die wir mit Großbuchstaben bezeichnen: U , V und T . Es gelten die Beziehungen

$$\begin{aligned} T &= \frac{t}{\bar{t}_2} = \frac{t}{\varepsilon \bar{t}} = \frac{\tau}{\varepsilon}, \\ U(T) &= \frac{s(\bar{t}_2 T)}{s_0} = \frac{s(\varepsilon \bar{t} T)}{s_0} = u(\varepsilon T) \quad \text{und} \\ V(T) &= \frac{c(\bar{t}_2 T)}{c_0} = \frac{c(\varepsilon \bar{t} T)}{c_0} = v(\varepsilon T), \end{aligned} \quad (7.31)$$

denn es gilt ja z.B. $u(\tau) = \frac{s(\bar{t}\tau)}{s_0}$. Mit der Kettenregel und (7.30) berechnen wir die Gleichungen, die von U und V gelöst werden:

$$\begin{aligned} \dot{U}(T) &= \frac{dU(T)}{dT} = \varepsilon \dot{u}(\varepsilon T) \\ &= \varepsilon (-u(\varepsilon T) (1 - v(\varepsilon T)) + k_l v(\varepsilon T)) \\ &= \varepsilon (-U(T) (1 - V(T)) + k_l V(T)) \\ \dot{V}(T) &= \frac{dV(T)}{dT} = \varepsilon \dot{v}(\varepsilon T) \\ &= u(\varepsilon T) (1 - v(\varepsilon T)) - \frac{k_m}{s_0} v(\varepsilon T) \\ &= U(T) (1 - V(T)) - \frac{k_m}{s_0} V(T) \end{aligned}$$

Oder in kompakterer Notation:

$$\begin{aligned} \dot{U} &= \varepsilon (-U (1 - V) + k_l V) \\ \dot{V} &= U (1 - V) - \frac{k_m}{s_0} V \end{aligned}$$

Der kleine Parameter steht jetzt nicht mehr vor einer Ableitung. Entwickeln wir dieses System asymptotisch nach ε und betrachten den asymptotischen Limes $\varepsilon = 0$, erhalten wir die Gleichungen.

$$\begin{aligned} \dot{U}_0 &= 0 \\ \dot{V}_0 &= U_0 (1 - V_0) - \frac{k_m}{s_0} V_0. \end{aligned}$$

Zusammen mit den Anfangswerten $U_0(0) = \frac{s(0)}{s_0} = 1$ und $V_0(0) = 0$ können wir dieses System explizit lösen, die zweite Gleichung ist ja dann linear in V_0 und wir erhalten die Lösung

$$U_0(T) = 1 \quad \text{und} \quad V_0(T) = \frac{1}{1 + \frac{k_m}{s_0}} \left(1 - e^{-\left(1 + \frac{k_m}{s_0}\right)T} \right).$$

Diese Lösungen stellen die Anfangsphase mit dem Aufbau des Komplexes dar (siehe Abb. 7.8).

Asymptotic Matching Die beiden Lösungen (U, V) und (u, v) beschreiben die Vorgänge der enzymatischen Reaktion auf unterschiedlichen Zeitskalen. Die Lösung (U, V) hat die korrekten Anfangswerte und beschreibt den Prozess für kurze Zeiten nach dem Prozentsstart, die Lösung (u, v) hat die falschen Anfangswerte, beschreibt den Prozess aber gut für große Zeiten t .

Lässt sich aus den beiden Lösungen auch eine Lösung zusammensetzen, die für alle Zeiten eine gute Annäherung liefert? Das geht mit der Methode des *Asymptotic Matching*. Die Idee ist folgende: Die Lösung auf der schnellen Zeitskala muss die richtigen Anfangswerte liefern, ist aber für große Zeiten nicht mehr korrekt. Die Lösung auf der langsamen großen Skala soll für das Langzeitverhalten zuständig sein, aber woher bekommen wir die Anfangsdaten? Die Antwort ist das **Asymptotic Matching**: Wir bestimmen die Anfangsdaten für die langsame Zeitskala aus dem Limes der Lösung für die schnelle Skala:

$$\lim_{T \rightarrow \infty} U(T) = 1 \stackrel{!}{=} u(0)$$

$$\lim_{T \rightarrow \infty} V(T) \frac{1}{1 + \frac{k_m}{s_0}} \stackrel{!}{=} v(0).$$

Damit erhalten wir $u(0) = 1$ und $v(0) = \frac{1}{1 + \frac{k_m}{s_0}}$. Man beachte dabei, dass unsere anfängliche Wahl von $u(0) = 1$ reines Glück war, mit **Asymptotic Matching** können wir die richtigen Anfangswerte berechnen. Jetzt können wir auch eine Lösung zusammenbauen, die für alle Zeiten gut approximiert, indem wir beide Lösungen addieren und den gemeinsamen Teil einmal subtrahieren:

Abb. 7.8 Substrat und Komplex im asymptotischen Limes mit der schnellen Zeitskala

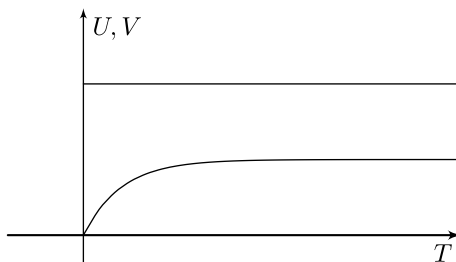
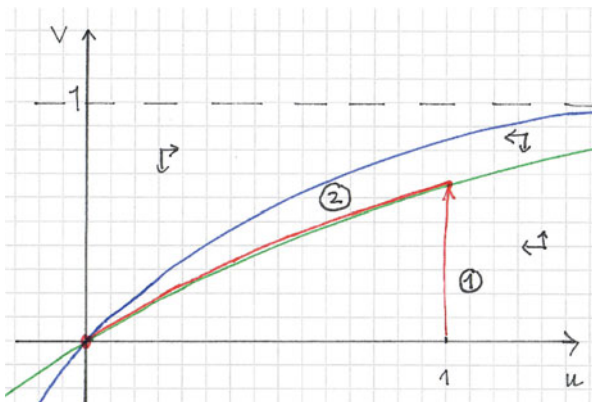


Abb. 7.9 Die gesamte Lösung im Phasenporträt: In Phase ① gilt $\dot{v} \approx 1/\varepsilon$, $\dot{u} \approx 1$, die Lösung wird durch (U_0, V_0) beschrieben. In Phase ② gilt $u(1-v) - \frac{k_m}{s_0}v \approx \varepsilon$, die Lösung läuft annähernd auf der Nullkline entlang und wird durch (u_0, v_0) beschrieben



$$u_{0,\text{full}}(\tau) = U_0(\tau/\varepsilon) + u_0(\tau) - 1$$

$$v_{0,\text{full}}(\tau) = V_0(\tau/\varepsilon) + v_0(\tau) - \frac{1}{1 + \frac{k_m}{s_0}}$$

Wir wollen die gefundene Lösung im Phasenporträt nachvollziehen (siehe Abb. 7.9). Für die Darstellung entscheiden wir uns für die langsame Zeitskala mit der Zeitvariablen τ , in Aufg. 7.10 soll die Diskussion mit der schnellen Zeitvariablen T geführt werden. Die volle Lösung wird also durch die Gl. (7.30) beschrieben. Zu Beginn der Entwicklung startet die Lösung in dem Punkt $(1, 0)$. Die rechte Seite der v -Gleichung ist zu dieser Zeit von der Größenordnung 1, die Ableitung $\dot{v}$ also von der Größenordnung $1/\varepsilon \approx 1000$. Die rechte Seite der u -Gleichung ist von der Größenordnung 1, also auch die Ableitung $\dot{u}$. Die Lösungstrajektorie geht also fast senkrecht nach oben und die Lösung wird durch die Funktionen $\tau \mapsto (U(\tau/\varepsilon), V(\tau/\varepsilon))$ approximiert. Das Verhältnis der beiden Ableitungen ändert sich erst, wenn die Lösung in die Nähe der v -Nullkline kommt: Damit $\dot{v}$ von der selben Größenordnung wie $\dot{u}$ ist, muss die rechte Seite der v -Differentialgleichung von der Größenordnung ε sein, die Lösung also fast auf der v -Nullkline entlanglaufen. Dies ist der zweite Teil der Entwicklung, der durch die Funktion $\tau \mapsto (u(\tau), v(\tau))$ beschrieben wird.

7.3 Der Tannenwickler und die Balsamtanne – zweiter Teil

Hier soll die in Abschn. 5.4 begonnene Diskussion um die Interaktion zwischen den Beständen der Balsamtanne und den Tannenwicklern wieder aufgegriffen werden. Die zeitliche Entwicklung der Tannenwickler-Population modellieren wir durch die Gl. (5.15), die wir hier noch einmal notieren:

$$\dot{P} = r_P P \left(1 - \frac{P}{K_P}\right) - S \frac{P^2}{V^2 + P^2} \quad (7.32)$$

Der Zustand des Balsamtannenbestands geht dabei insbesondere in die Parameter K_P , die Kapazität des Gebiets für Tannenwickler, und V , die Populationsdichte mit halber maximaler Bejagungsrate, ein.

Jetzt sollen die zeitliche Entwicklung des Tannenbestandes und die Interaktion nach [60] modelliert werden.

7.3.1 Die Modellbildung

Der Tannenbestand

Um ein umfassendes Bild über den Zustand eines Tannenbestands in einem bestimmten Gebiet zu erhalten, sollte man etwas über die Altersstruktur, den Zustand der Belaubung (hier ist natürlich von Nadeln die Rede) und über weitere Gegebenheiten, wie die Versorgung mit Wasser und Nährstoffen, wissen. Im Sinne der Einfachheit soll der Bestand aber nur durch zwei Größen gekennzeichnet werden: Mit $Z(t)$ bezeichnen wir die Menge an Zweigen pro Fläche, die der Tannenbestand zur Zeit t ausgebildet hat. Damit das eine vernünftige Größe ist, muss definiert werden, was unter einem Zweig verstanden wird: Unter einem Standardzweig verstehen wir einen Nadeln tragenden Teil des Geästs eines Baums, der, wenn von einem Polygonzug umrandet, eine Fläche von einem Quadratfuß, das sind 0.09 m^2 , einnimmt. Schon einmal zur groben Orientierung: Für einen gesunden, voll entwickelten Tannenbestand geht man von einer Zweiganzahl von ca. $25\,000 \text{ ha}^{-1}$ aus, siehe [60].

Die zweite Größe beschreibt den physiologischen Zustand der benadelten Zweige und soll Werte zwischen null und eins annehmen. Dabei bedeutet der Zustand eins, dass die Zweige voll benadelt sind, und der Zustand null, dass die Zweige entnadelt und nicht mehr lebensfähig sind. In [60] ist hier von einer „Energiereserve“ der Bäume die Rede. Da es sich aber offensichtlich nicht um eine Energie im physikalischen Sinn handelt, wollen wir lieber von dem Zustand der Benadelung sprechen, auch wenn es schwierig sein dürfte, dafür eine konkrete Messvorschrift festzulegen. Wir bezeichnen diese dimensionslose Größe mit E .

Nun benötigen wir Differentialgleichungen, welche die Entwicklung der beiden Größen beschreiben sollen. Sicherlich ist die Anzahl der Zweige pro Fläche eine begrenzte Größe und wir wählen das Modell des logistischen Wachstums

$$\dot{Z} = r_Z Z \left(1 - \frac{Z}{K_Z E} \right).$$

Der Faktor E im Nenner drückt dabei aus, dass die langfristige Anzahl der Zweige auch von dem Zustand der Benadelung abhängt, also bei schwacher Benadelung niedriger bleibt als bei optimaler Benadelung. Auch hier schon einmal eine Abschätzung für den Wert von r_Z : Ein Tannenbestand braucht ca. 80 Jahre von der Anpflanzung bis zum voll ausgewachsenen Zustand. Die Lösung der entdimensionalisierten logistischen Gleichung $\dot{x} = x(1 - x)$ zum Anfangswert $x_0 = \frac{1}{2}$ ist gegeben durch

$$x(\tau) = \frac{e^\tau}{1 + e^\tau},$$

siehe Aufg. 5.14. Es wird also die dimensionslose Zeit $\tau = \ln(9) = 2.20$ benötigt, damit der Bestand von 50 % auf 90 % der Kapazität anwächst. Dieser Zeit sollten mit Dimensionen 40 Jahre entsprechen. Die zur Entdimensionalisierung benutzte Zeitskala ist gerade $\bar{t} = \frac{1}{r_Z}$. Wir erhalten also aus der Gleichung $2.20\bar{t} = 40$ a die grobe Abschätzung

$$r_Z = 0.055 \text{ a}^{-1}.$$

Wir vergleichen diesen Wert mit dem Ratenparameter der Tannenwickler r_P . Innerhalb eines Jahres kann sich die Tannenwickler-Dichte ausgehend von einer sehr geringen Dichte und ohne Bejagung ungefähr verfünffachen. Für geringe Populationsdichten ist das logistische Modell im Wesentlichen exponentiell. Wir können also eine Abschätzung für den Ratenparameter r_P aus der Gleichung $e^{r_P 1 \text{ a}} = 5$ gewinnen, mit der Lösung $r_P = 1.61 \text{ a}^{-1}$. P ist in diesem Sinne eine schnelle Variable und Z eine langsame Variable und das Verhältnis der Ratenparameter r_Z/r_P ist von der Größenordnung Hundertstel.

Nun zum Zustand der Benadelung E . Ein schädlingsfreier Bestand sollte sich in den Zustand $E = 1$ entwickeln. Wird der Bestand einmal geschädigt, so erholt er sich davon in einer anschließenden schädlingsfreien Zeit in der Regel wieder. Wir setzen auch hier ein logistisches Wachstum an:

$$\dot{E} = r_E E(1 - E)$$

(In Aufg. 7.12 soll ein entsprechendes Modell diskutiert werden, in dem für die Benadelungsvariable E ein beschränktes Wachstum, $\dot{E} = r_E(1 - E)$, angesetzt wird.) Um eine Größenordnung für den Ratenparameter r_E abzuschätzen, setzen wir eine Erholungszeit von sieben Jahren an (eine Tanne trägt in der Regel die Nadeln von sieben Jahrgängen; wird einmal ein Jahrgang ganz abgefressen, sieht man sieben Jahre später nichts mehr davon). Wir setzen grob an, dass das Abfressen des jüngsten Jahrgangs den Zustand auf $E = 0.5$ reduziert. Es dauert sehr grob 5 a, bis der Zustand sich wieder auf $E = 0.9$ erholt hat. Dem entspricht im logistischen Wachstumsmodell ein Ratenparameter $r_E = \frac{2.20}{5 \text{ a}} = 0.44 \text{ a}^{-1}$. Das Verhältnis r_E/r_P ist also von der Größenordnung Zehntel. Im nächsten Schritt soll die Interaktion des Tannenwicklers mit dem Tannenbestand modelliert werden.

Die Wickler-Tannen-Interaktion

Die Wickler fressen im Wesentlichen die Nadeln des jüngsten Jahrgangs. Sind diese vollkommen abgefressen, so fressen sie auch die älteren Nadeljahrgänge. Limitiert sind sie im Wesentlichen durch den Platz auf dem Zweig und erst ganz am Schluss, kurz vor dem vollständigen Abfraß, durch das Nahrungsangebot an Nadeln. Es erscheint also sinnvoll, die Kapazität K_P als im Wesentlichen proportional zur Anzahl der Zweige anzusetzen $K_P \sim Z$. Die Abhängigkeit vom Benadelungszustand E spielt erst dann eine Rolle, wenn E sehr nah an null ist. Auch hier können wir wieder gut die *S-förmige Funktion vom Typ III* verwenden,

indem wir einen zusätzlichen Faktor der Form $E^2/(T^2 + E^2)$ hinzufügen. Hierbei wird die Schwelle T einen eher kleinen Wert erhalten, weil der Vorrat an Nadeln das Wachstum der Wicklerlarven erst dann bremst, wenn kaum noch Nadeln am Zweig sind. Insgesamt modellieren wir die Kapazität für die Wicklerlarven durch

$$K_P = KZ \cdot \frac{E^2}{T^2 + E^2}$$

mit einer Konstanten K , welche die Kapazität für Wicklerlarven pro voll benadelten Zweig angibt.

Nun zu V , der Populationsdichte, ab der die Bejagungsrate durch die Vögel den halben maximalen Wert erreicht hat. Ein sinnvoller Ansatz ist, hier die Besiedelungsdichte der Zweige durch Wicklerlarven als den entscheidenden Wert zu betrachten: Wenn im Schnitt mehr als so und so viele Larven einen Zweig besiedeln, fangen die Vögel an, die Larven zu fressen. Wir setzen also

$$V = WZ$$

mit einer Konstanten W , welche die Bejagungsschwelle pro Zweig angibt.

Nun müssen wir noch modellieren, mit welcher Rate die Wicklerlarven die Tannen schädigen. Da die Larven die Nadeln und nicht die Zweige fressen, wirkt sich die Interaktion auf die Variable E und nicht auf die Variable Z aus. Naheliegender ist es, die Fressrate als proportional zur Larvenanzahl pro Zweig anzusetzen. Die Larven können sehr lange bis zur Sättigungsgrenze fressen, und erst wenn die Benadelung unter den Schwellenanteil absinkt, nimmt auch die Fressrate pro Larve wieder ab. Als Fressrate setzen wir also

$$C \frac{P}{Z} \frac{E^2}{T^2 + E^2}$$

mit einer Konstante C , welche die Schädigungsrate am Benadelungszustand pro Larve am Zweig angibt. Insgesamt erhalten wir damit das Modell

$$\begin{aligned} \dot{P} &= r_P P \left(1 - \frac{P}{KZ \frac{E^2}{T^2 + E^2}} \right) - \frac{SP^2}{W^2 Z^2 + P^2} \\ \dot{Z} &= r_Z Z \left(1 - \frac{Z}{K_Z E} \right) \\ \dot{E} &= r_E E(1 - E) - C \frac{P}{Z} \frac{E^2}{T^2 + E^2}. \end{aligned} \tag{7.33}$$

Wir wollen jetzt dieses System analysieren und überprüfen, ob sich das Phänomen der periodischen Wicklerausbrüche mit einem damit zusammenhängenden Zusammenbruch des Tannenbestandes finden lässt und in welchem Bereich die Parameter dafür liegen müssen.

Da es sich hier um ein recht kompliziertes dreidimensionales System handelt, wir also das Werkzeug der Phasenporträtanalyse nicht in der Form wie bei zweidimensionalen Systemen

Tab.7.1 Parameterwerte für das Tannenwickler-Tanne-System (7.33) nach [60] und [64]. L = Larven, Z = Zweig(e)

Parameter	r_P	r_Z	r_E	K	W	K_Z	S	C
Maßzahl	1.52	0.095	0.92	355	1.11	25 440	45 300	0.00195
Einheit	a^{-1}	a^{-1}	a^{-1}	L/Z	L/Z	Z ha^{-1}	$\text{L ha}^{-1} \text{a}^{-1}$	Z/L a^{-1}

anwenden können, wollen wir zunächst das System numerisch nach unterschiedlichen Phänomen absuchen, die wir anschließend mit analytischen und asymptotischen Betrachtungen erläutern. Insgesamt gehen acht Parameter in das Modell ein. Um hier den Überblick zu behalten wollen wir davon sieben Parameter festhalten und nur einen einzigen, nämlich T , variieren (in Aufg. 7.11 sollen auch die anderen Parameter variiert werden). Wir übernehmen die Parameterwerte $r_P, r_Z, r_E, K, K_Z, W, S$ und C von Ludwig et al. aus [60], die sich bei der Festlegung dieser Parameterwerte auf jahrzehntelange Feldstudien berufen, die in [64] dokumentiert sind, siehe Tab. 7.1 (und vergleiche auch mit unseren sehr hemdsärmlichen Abschätzungen).

Für den Schwellenwert T schreiben Ludwig et al. lediglich, dass er „sehr klein“ sein solle. Es ist also interessant zu untersuchen, inwiefern unterschiedliche Werte von T das qualitative Verhalten des Gesamtsystems beeinflussen. Wie häufig ist es aber zunächst einmal sinnvoll, zu entdimensionalisieren, um „kleine Größen“ zu identifizieren.

Entdimensionalisierung

Wir bezeichnen die dimensionslosen Variablen für t, P, Z, E mit τ, p, z, e (man beachte, dass E eigentlich schon dimensionslos ist, also $E = e$), die zunächst unbestimmten Skalen mit $\bar{t}, \bar{P}$ und $\bar{Z}$ (weil E schon dimensionslos ist, wird keine eigene Skala benötigt) und berechnen die dimensionslosen Gleichungen:

$$\begin{aligned}\dot{p} &= \bar{t} r_P p \left(1 - \frac{\bar{P}}{K \bar{Z}} \frac{p}{z} \frac{T^2 + e^2}{e^2} \right) - \frac{\bar{t} S}{\bar{P}} \frac{p^2}{\frac{W^2 \bar{Z}^2}{\bar{P}^2} z^2 + p^2} \\ \dot{z} &= \bar{t} r_Z z \left(1 - \frac{\bar{Z}}{K_Z} \frac{z}{e} \right) \\ \dot{e} &= \bar{t} r_E e (1 - e) - \frac{\bar{t} C \bar{P}}{\bar{Z}} \frac{p}{z} \frac{T^2 + e^2}{e^2}\end{aligned}$$

Es ist kanonisch, für die Zweige und Larven die jeweiligen Kapazitäten als Skalen zu benutzen:

$$\bar{Z} = K_Z \quad \text{und} \quad \bar{P} = K K_Z$$

Für die Zeitskala müssen wir zwischen den Ratenkonstanten der drei Variablen wählen. Wir entscheiden uns für die Konstante der „langsamsten“ Variablen, also für

Tab. 7.2 Dimensionslose Parameter für das Tannenwickler-Tanne-System (7.33) mit numerischen Werten bezogen auf die dimensionsbehafteten Parameterwerte in Tab. 7.1

Parameter	ε_1	ε_2	ε_3	ε_4	α	T
Definition	$\frac{r_P}{r_Z}$	$\frac{S}{r_Z \bar{K} \bar{K}_Z}$	$\frac{r_E}{r_Z}$	$\frac{CK}{r_Z}$	$\frac{W}{\bar{K}}$	T
Wert	16	0.053	9.7	7.3	0.0031	???

$$\bar{t} = \frac{1}{r_Z}.$$

Dies hat den Vorteil, dass auf dieser Zeitskala die beiden anderen Variablen schnell sind und wir dadurch asymptotische Methoden anwenden können. Mit dieser Wahl erhalten wir das folgende dimensionslose System, wobei die Definition der dimensionslosen Parameter und ihre konkreten numerischen Werte auf Basis der Parameterwerte in Tab. 7.1 in 7.2 zusammengefasst sind. (Wir würden aus typografischen Gründen eigentlich gerne T durch t ersetzen, kämen dadurch aber in Konflikt mit der Belegung von t als dimensionsbehafteter Zeit. Deshalb verzichten wir darauf.)

$$\begin{aligned} \dot{p} &= \varepsilon_1 p \left(1 - \frac{p}{z} \frac{T^2 + e^2}{e^2} \right) - \varepsilon_2 \frac{p^2}{\alpha^2 z^2 + p^2} \\ \dot{z} &= z \left(1 - \frac{z}{e} \right) \\ \dot{e} &= \varepsilon_3 e(1 - e) - \varepsilon_4 \frac{p}{z} \frac{e^2}{T^2 + e^2} \end{aligned} \quad (7.34)$$

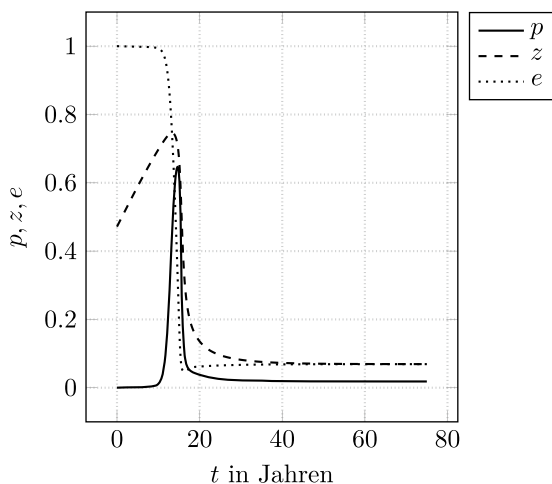
7.3.2 Numerische Simulationen

Um einen ersten Überblick zu gewinnen, berechnen wir numerische Lösungen zu (7.34) mit dem expliziten Euler-Verfahren zu unterschiedlichen Parameterwerten des Schwellenparameters T . (Um dieses Kapitel gut nachvollziehen zu können, sollte auch die geneigte Leserin das System (7.34) jetzt numerisch implementieren (was nicht ganz unhakelig ist)). Wir starten dabei immer mit einem halb entwickelten Tannenbestand bei voller Gesundheit mit wenigen Schädlingen, konkret $P_0 = 1000$ Larven pro Hektar, $Z_0 = 12\,000$ Zweige und $E_0 = 1$, entsprechend

$$p_0 = 0.00011, \quad z_0 = 0.472, \quad e_0 = 1. \quad (7.35)$$

Um die Zeit leichter lesen zu können, skalieren wir sie allerdings mit der Einheit ein Jahr anstatt mit der eigentlich gewählten Einheit $\bar{t} = \frac{1}{r_Z} = 10.53$ a. Ludwig et al. fordern T sehr

Abb. 7.10 Numerisch berechnete Lösung von (7.34), (7.35) mit $T = 0.1$, explizites Euler-Verfahren mit Schrittweite $\Delta t = 0.1$ a



klein, wir probieren es also mit $T = 0.1$ und $T = 0.01$. Die numerischen Ergebnisse sind in Abb. 7.10 und 7.11 dargestellt.

In beiden Szenarien gibt es einen Zusammenbruch des Tannenbestandes. Im ersten Szenario mit $T = 0.1$ scheint sich das System nach dem Zusammenbruch in einem Gleichgewichtszustand mit kleinem Tannenbestand bei schlechter Benadelung und einer kleinen Wicklerpopulation zu stabilisieren.

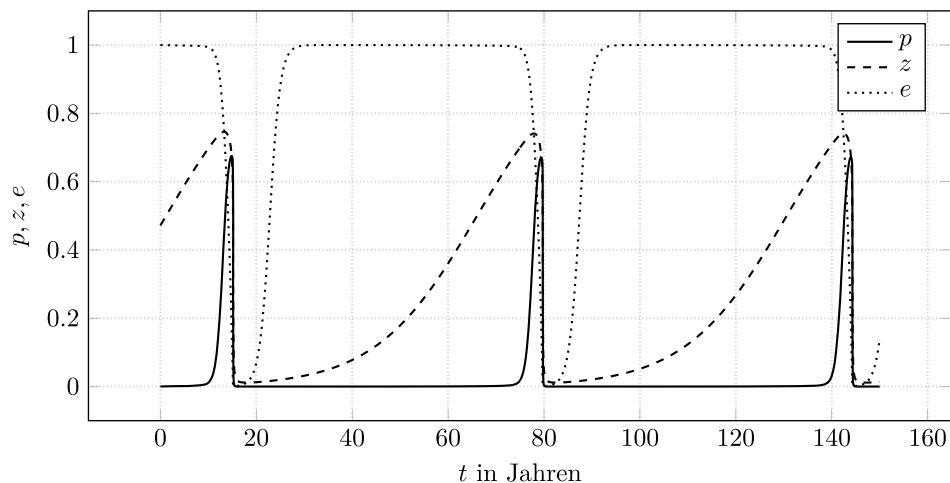


Abb. 7.11 Numerisch berechnete Lösung von (7.34), (7.35) mit $T = 0.01$, explizites Euler-Verfahren mit Schrittweite $\Delta t = 0.1$ a

Im zweiten Szenario, $T = 0.01$, brechen alle drei Größe fast auf null zusammen, die Benadelungsgröße e erholt sich und wächst innerhalb von ein paar Jahren wieder auf den Wert „volle Benadelung“. Daran anschließend wächst der Zweigbestand anscheinend logistisch, während der Wicklerbestand ca. 60 Jahre klein bleibt, dann in sehr kurzer Zeit wächst, worauf sich der nächste katastrophale Zusammenbruch anschließt. Das Modell enthält also für $T = 0.01$ die gewünschten Eigenschaften der langen gleichmäßigen Wachstumsphase, des heftigen Ausbruchs der Wicklerbestände, des Zusammenbruchs und der periodischen Wiederholung, wobei bei den benutzten Parameterwerten die Periode ca. 65 Jahre beträgt, also deutlich größer ist als die beobachteten 35 bis 40 Jahre.

Was haben wir nun davon außer einem mathematischen Modell, welches bei bestimmten Parameterwerten ein in bestimmten Beziehungen ähnliches Verhalten aufweist wie die Realsituation? Wir können anhand des Modell versuchen, die Vorgänge in der Realsituation zu *erklären*, d.h. Hypothesen aufzustellen, warum es zu einem bestimmten Zeitpunkt zu diesem heftigen Zusammenbruch kommt, warum bei der einen Parameterkonstellation ein periodisches Verhalten, in der anderen ein stationäres Verhalten eintritt. Das soll im nächsten Abschnitt geschehen.

7.3.3 Analyse des Systems mit asymptotischen Methoden

Da es sich um ein dreidimensionales System handelt, können wir nicht so einfach das Phasenporträt analysieren. Um die Komplexität zu reduzieren, wollen wir die unterschiedlichen Geschwindigkeiten der einzelnen Variablen ausnutzen. Dazu vergleichen wir die Änderungsraten der einzelnen Variablen und starten mit z . Aus der Differentialgleichung

$$\dot{z} = z \left(1 - \frac{z}{e} \right)$$

können wir ablesen, dass sich z immer in Richtung e entwickelt. Wenn z in Richtung e wächst, ist die Änderungsrate durch $e/4 \leq 1/4$ beschränkt. Wir werden gleich sehen, dass das im Vergleich zu den Änderungsraten der anderen Variablen langsam ist. Wenn z in Richtung e fällt, kann der Betrag der Änderungsrate dann groß werden, wenn e sehr klein und z nicht sehr klein ist. Nur unter diesen Umständen wird die Variable z „schnell“.

Nun zu p . Mit den konkreten Parameterwerten (T halten wir noch unbestimmt) haben wir

$$\dot{p} = 16p \left(1 - \frac{p}{z} \frac{T^2 + e^2}{e^2} \right) - 0.050 \frac{p^2}{9.7 \times 10^{-6} z^2 + p^2}.$$

Diese Änderungsrate ist also groß mit den Ausnahmen, wenn $p \approx 0$ oder $p \approx \frac{z^2 e^2}{T^2 + e^2}$ gilt.

Um die Änderungsrate von e besser abschätzen zu können, stellen wir etwas um:

$$\begin{aligned}\dot{e} &= \varepsilon_3 e \left(1 - e - \frac{\varepsilon_4}{\varepsilon_3} \frac{p}{z} \frac{e/T^2}{1 + (e/T)^2} \right) \\ &= 9.7e \left(1 - e - 0.752 \frac{p}{z} \frac{e/T^2}{1 + (e/T)^2} \right)\end{aligned}$$

Die Benadelung e ist also eine schnelle Variable bis auf die Fälle, in denen $e \ll 1$ oder der Term in der Klammer sehr klein ist.

Wir trennen die Situation asymptotisch: Auf einer schnellen Zeitskala können wir z als konstant betrachten (mit der Ausnahme der oben erwähnten Situation mit $e \approx 0$ und $z \approx 1$). In dem zweidimensionalen p - e -System, in das z als konstanter Parameter eingeht, stellen sich p und e in einer relativ kurzen Zeitspanne in einen stabilen stationären Punkt des zweidimensionalen Systems ein. z entwickelt sich auf einer langsamen Zeitskala in Richtung e . Dadurch ändern sich die stationären Punkte von p und e langsam. Schnelle Veränderungen werden immer dann auftreten, wenn sich durch die langsame Änderung von z die Konstellationen der stationären Punkte von p und e ändern.

Dummerweise ist auch das zweidimensionale p - e -System mit konstantem z noch recht kompliziert (man versuche einmal für festes z die Nullklinen zu berechnen). Beide Gleichungen haben jedoch eine Struktur, die wir bereits in Abschn. 5.4 studiert haben und die wir jetzt ausnutzen wollen. In Abschn. 5.4 haben wir eine Differentialgleichung der Form

$$\dot{x} = \varepsilon x \left(a \left(1 - \frac{x}{b} \right) - \frac{x}{1+x^2} \right) \quad (7.36)$$

studiert und gesehen, dass diese Gleichung je nach Parameterkombination (a, b) einen bis drei positive stationäre Punkte hat, siehe Abb. 7.12. Im Fall von drei positiven stationären Punkten $x_1^* < x_2^* < x_3^*$ sind die beiden äußeren stabil und der mittlere instabil. Die Parameterpaare (a, b) , bei denen es genau zwei stationäre Punkte gibt, also der Übergang von einem

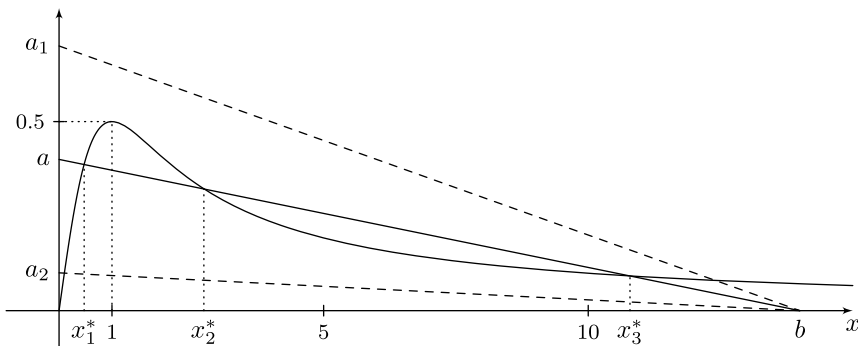


Abb. 7.12 Die DGL $\dot{x} = \varepsilon x \left(a \left(1 - \frac{x}{b} \right) - \frac{x}{1+x^2} \right)$ hat einen bis drei positive stationäre Punkte

zu drei stationären Punkten oder umgekehrt stattfindet, sind von besonderem Interesse und lassen sich durch die *kritische Kurve* darstellen, siehe Abb. 7.13. Für eine ausführlichere Diskussion siehe Abschn. 5.4.2. Aus der Parameterdarstellung (in der a im zweiten und b im ersten Slot dargestellt werden)

$$(1, \infty) \ni y \mapsto (2y^3/(1-y^2), 2y^3/(1+y^2)^2)^t$$

lässt sich ablesen, dass für b hinreichend groß die beiden Äste der kritischen Kurve, $\bar{a}(b)$ der obere Ast und $\underline{a}(b)$ der untere Ast, durch einfache Ausdrücke gut approximiert werden können:

$$\bar{a}(b) \approx \frac{1}{2} \quad \text{und} \quad \underline{a}(b) \approx \frac{4}{b} \quad (7.37)$$

Diese Approximationen sind in Abb. 7.13 durch die gestrichelten Linien dargestellt. Die Fixpunktgleichung $a(1-x/b) = x/(1+x^2)$ lässt sich zwar algebraisch lösen (es geht um eine Gleichung dritten Grades), die resultierenden Ausdrücke hängen jedoch recht kompliziert von a und b ab. In einigen Fällen können wir jedoch gute einfache Approximationen direkt aus Abb. 7.12 ablesen:

1. Sind a und b hinreichend groß, so gilt

$$x_3^*(a, b) \approx b. \quad (7.38)$$

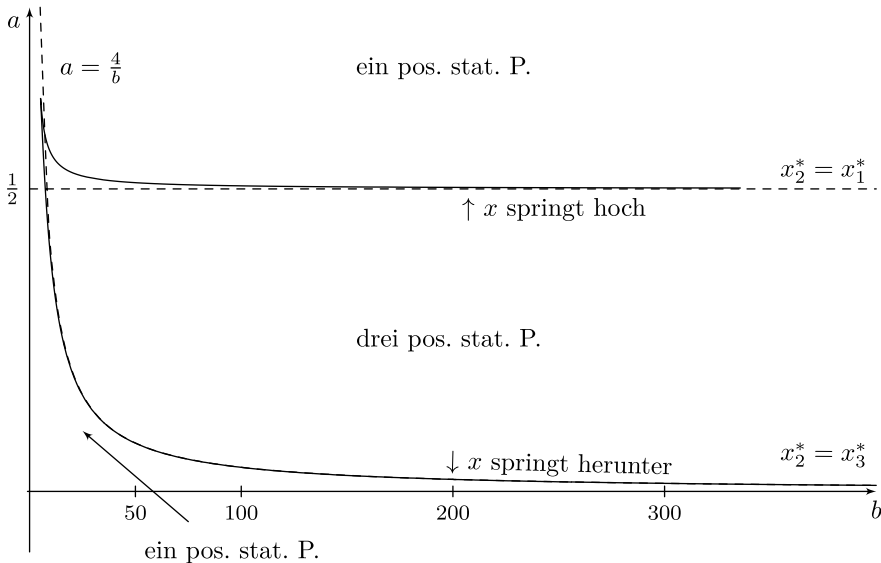


Abb. 7.13 Die kritische Kurve der DGL (7.36) mit den Approximationen für die beiden Äste (gestrichelt), der Anzahl der positiven stationären Punkte und dem Sprungverhalten beim Überschreiten der kritischen Kurve

2. Für b hinreichend groß und a hinreichend klein gilt

$$x_1^*(a, b) \approx a, \quad (7.39)$$

dabei beachte man, dass die Tangente an den Graphen der Funktion $x \mapsto x/(1+x^2)$ im Ursprung gerade der Graph der Identität ist.

Mit diesem Rüstzeug kehren wir nun zu dem Wickler-Tannen-System zurück und schreiben die p - und e -Gleichungen so um, dass wir die eben gesammelten Fakten anwenden können:

$$\begin{aligned} \dot{p} &= \frac{\varepsilon_2 p}{\alpha z} \left(a_p \left(1 - \frac{p/\alpha z}{b_p} \right) - \frac{p/\alpha z}{1 + (p/\alpha z)^2} \right) \\ \dot{e} &= \frac{\varepsilon_4 p e}{T z} \left(a_e \left(1 - \frac{e/T}{b_e} \right) - \frac{e/T}{1 + (e/T)^2} \right) \end{aligned}$$

mit

$$a_p = \frac{\varepsilon_1 \alpha}{\varepsilon_2} z = 0.97 z, \quad b_p = \frac{1}{\alpha} \cdot \frac{(e/T)^2}{1 + (e/T)^2} = 323 \frac{(e/T)^2}{1 + (e/T)^2} \quad (7.40)$$

$$a_e = \frac{\varepsilon_3 T z}{\varepsilon_4 p} = 1.33 T \frac{z}{p}, \quad b_e = \frac{1}{T}. \quad (7.41)$$

Zustände, in denen $p/\alpha z$ im oder nah am stationären Punkt x_1^* ist, nennen wir *endemisch*; Zustände, in denen $p/\alpha z$ im oder nah am stationären Punkt x_3^* ist, nennen wir *epidemisch*.

Wir wollen jetzt untersuchen, wann p die kritische Kurve kreuzt. Die Bedingung $b_p \gg 1$ ist nur verletzt, wenn $e \ll T$ gilt. In allen anderen Fällen

- verschwindet (bzw. erscheint) der endemische Zustand x_1^* , wenn $a_p \approx \frac{1}{2}$ gilt, also bei

$$z \approx \frac{\varepsilon_2}{2\varepsilon_1 \alpha} = 0.53.$$

Das ist dann der Fall, wenn der Tannenbestand ungefähr halb ausgewachsen ist.

- verschwindet (bzw. erscheint) der epidemische Zustand x_3^* bei $a_p = 4/b_p$, also bei

$$z \approx \frac{4\varepsilon_2}{\varepsilon_1} \cdot \frac{(1 + (e/T)^2)}{(e/T)^2} = 0.013(1 + (T/e)^2).$$

Der epidemische Zustand verschwindet also erst, wenn der Tannenbestand nur noch wenige Prozente der Kapazität enthält (oder wenn $T/e \geq 1$, dann ist aber die verwendete Approximation nicht mehr gut).

Nun zu den Übergängen für e . Zustände, in denen e/T im oder nah am stationären Punkt x_1^* ist, nennen wir *krank*; Zustände, in denen e/T im oder nah am stationären Punkt x_3^* ist, nennen wir *gesund*. Die Bedingung $b_e = \frac{1}{T}$ hinreichend groß ist stets erfüllt.

- Der gesunde Zustand x_3^* verschwindet (bzw. erscheint), wenn

$$\frac{p}{z} \approx \frac{\varepsilon_3}{4\varepsilon_4} = 0.33.$$

Das ist dann der Fall, wenn ca. 100 Larven an einem Zweig sitzen ($K = 355$ Larven pro Zweig).

- Der kranke Zustand e_1^* verschwindet (bzw. erscheint), wenn

$$\frac{p}{z} \approx \frac{2\varepsilon_3 T}{\varepsilon_4} = 1.33T.$$

Für $T = 0.1$ verschwindet der kranke Zustand bei ca. 50 Larven pro Zweig, für $T = 0.01$ erst bei ca. 5 Larven pro Zweig.

Nach diesen Vorbereitungen können wir die Entwicklung der Wickler-Tannen-Populationen in diesem asymptotischen Bild beschreiben.

Der Ausbruch der Epidemie Wir starten mit den Anfangswerten $p_0 = 0.00011$, $z_0 = 0.472$ und $e_0 = 1$ aus (7.35). p befindet sich im endemischen Zustand und e in einem gesunden Zustand. Die Variable z entwickelt sich langsam in Richtung $e \approx 1$. Mit der wachsenden Variable z passiert a_p den kritischen Wert 0.5 (ungefähr dann, wenn $z = 0.53$) und der endemische Zustand für p verschwindet. p entwickelt sich dann auf der schnelleren Zeitskala in Richtung seines epidemischen Zustands $p/\alpha z \approx x_3^*$, also $p \approx z \frac{(e/T)^2}{1+(e/T)^2}$. Diese Entwicklung ist anfänglich noch langsam, weil p von einem sehr kleinen Wert startet.

Der Zusammenbruch der Gesundheit und des Bestandes Mit wachsendem p sinkt der gesunde Zustand $e/T \approx x_3^*$, also $e \approx 1$, zunächst langsam ab. Bei $\frac{p}{z} \approx 0.33$ verschwindet der gesunde Zustand und e/T entwickelt sich in Richtung des kranken Zustands x_1^* . Beim Übergang zum kranken Zustand gilt $a_e \approx 4/b_e = 4T \ll 1$. Damit ist die Approximation $x_1^* \approx a_e$ zulässig und e entwickelt sich in Richtung

$$e \approx \frac{\varepsilon_3 T^2}{\varepsilon_4} \frac{z}{p}.$$

Dieses Abfallen geschieht auf einer schnellen Zeitskala.

Gelegentlich passiert e bei diesem Sinken den Wert von z . Sobald das passiert, beginnt z in Richtung e mit $\dot{z} \approx -\frac{z}{e}$ zu fallen. Wenn dieser Quotient betragsmäßig groß ist, wird z schnell. Mit dem fallenden e und z sinkt auch die epidemische Wicklerpopulation $p \approx z \frac{e/T}{1+(e/T)^2}$. Das ist der Zusammenbruch der Gesundheit, des Tannenbestandes und auch der Wicklerpopulation. Kurz gesagt passiert also Folgendes: p wechselt vom endemischen in den epidemischen Zustands, worauf e vom gesunden in den kranken Zustand wechselt und der Bestand zusammenbricht. Bis zu diesem Zeitpunkt der Entwicklung spielte der konkrete

Wert des Parameters T noch keine große Rolle, im Weiteren entscheidet er aber über den folgenden Verlauf.

Stationärer kranker epidemischer Zustand oder vollständiger Zusammenbruch mit anschließendem Neuaufbau Wie kommt nun der große Unterschied der folgenden Entwicklung je nach Wert von T zustande?

Sobald $\dot{z} = z(1 - \frac{z}{e}) \approx 1$ gilt, können wir wieder das asymptotische Bild verwenden, mit $e/T \approx x_1^*$, $p/\alpha z \approx x_3^*$ und dem e langsam folgenden z .

Wir fragen uns nun, ob p wieder in den endemischen Zustand $p/\alpha z \approx x_1^*$ wechselt oder im epidemischen Zustand bleibt. Dabei approximieren wir weiterhin nach (7.37) bzw. (7.38) und (7.39)

$$\underline{a}(b) = 4/b \quad \text{sowie} \quad e_1^*/T \approx a_e \quad \text{und} \quad p_3^*/\alpha z \approx x_3^* \approx b_p. \quad (7.42)$$

Dabei muss uns aber bewusst sein, dass diese Approximationen nur noch Größenordnungen liefern und keine vernünftigen numerischen Werte mehr, weil die Voraussetzungen für diese Approximationen nicht mehr gut erfüllt sind.

Solange p im epidemischen Zustand ist, gilt mit (7.42) und (7.40)

$$p \approx z \cdot \frac{(e/T)^2}{1 + (e/T)^2}, \quad \text{also} \quad \frac{z}{p} \approx \frac{1 + (e/T)^2}{(e/T)^2}. \quad (7.43)$$

Für e folgt mit (7.42), (7.41) und (7.43)

$$e \approx \frac{\varepsilon_3}{\varepsilon_4} T^2 \frac{z}{p} \approx \frac{\varepsilon_3}{\varepsilon_4} T^2 \frac{1 + (e/T)^2}{(e/T)^2}. \quad (7.44)$$

Damit p in den endemischen Zustand wechseln kann, muss die kritische Kurve gekreuzt werden. In der (hier schlechten) Approximation $\underline{a}_p(b) = 4/b_p$ muss z nach (7.40) also den kritischen Wert

$$z_{\text{krit}} \approx \frac{4\varepsilon_2}{\varepsilon_1} \frac{1 + (e/T)^2}{(e/T)^2}$$

erreichen können. Dafür muss $z_{\text{krit}} > e$ gelten (denn z folgt langsam e), was mit (7.44) auf

$$T < 2 \sqrt{\frac{\varepsilon_2 \varepsilon_4}{\varepsilon_1 \varepsilon_3}} = 0.01$$

führt. Während der numerische Wert aufgrund der diversen schlechten Approximation nicht besonders aussagekräftig ist, zeigt diese Überlegung, dass T klein genug sein muss, damit p in den endemischen Zustand zurückwechseln kann.

Diese Überlegung wird durch die Numerik bestätigt. In Abb. 7.14 folgen wir in den beiden Fällen $T = 0.1$ und $T = 0.01$ der p -Trajektorie in der b - a -Ebene und erkennen, dass nur im letzteren Fall p wieder in den endemischen Zustand wechselt. Dafür muss ja der untere Ast der kritischen Kurve von oben nach unten gekreuzt werden. Das passiert im Fall $T = 0.01$ zwischen den Zeiten 15.25 und 15.3 Jahren, im Fall $T = 0.1$ gar nicht.

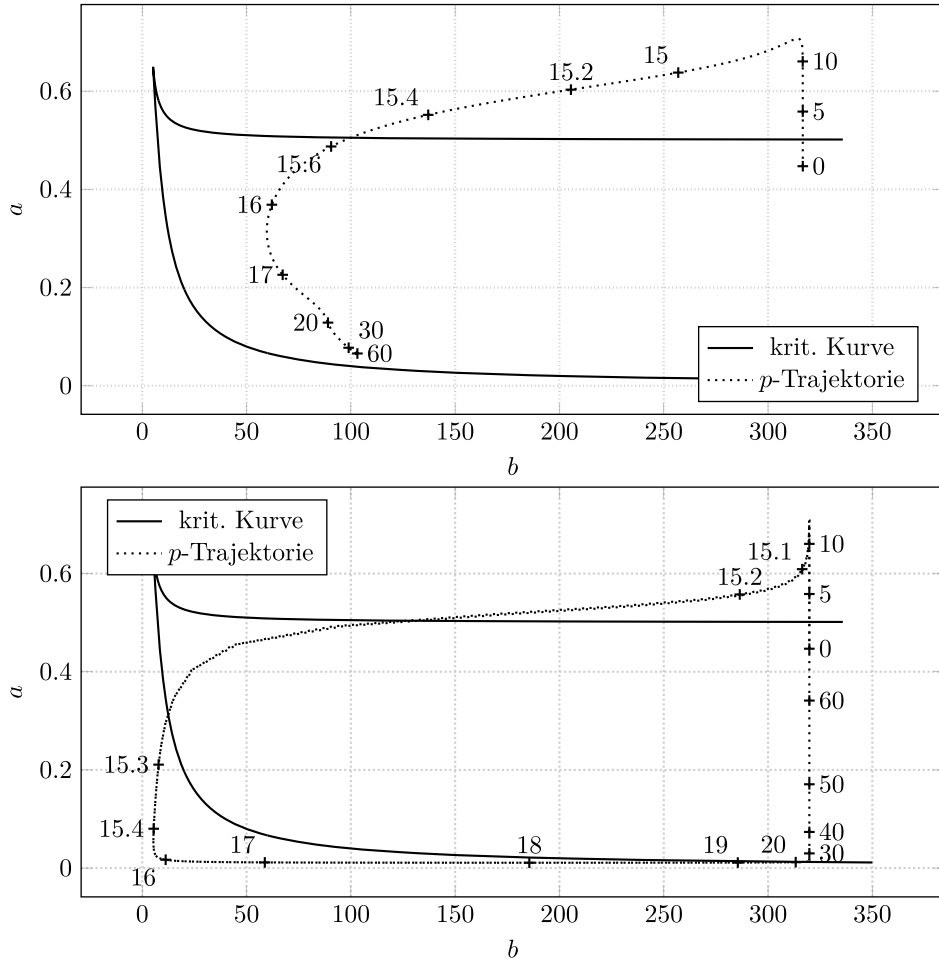


Abb. 7.14 Die p -Trajektorie im b - a -Diagramm für die Anfangswerte (7.35). Oben $T = 0.1$, unten $T = 0.01$. Die Jahreszahlen sind an der Trajektorie notiert

Wir stellen uns jetzt die Frage, unter welchen Umständen e in den gesunden Zustand zurückkehren kann. Dafür muss $a_e = \frac{\varepsilon_3 T z}{\varepsilon_4 p}$ im Fall $T = 0.1$ den Wert 0.56 und im Fall $T = 0.01$ den Wert 0.505 überschreiten (diese Werte kann man zum Beispiel aus einem GeoGebra-Plot der kritischen Kurve an der Stelle $b = 1/T$ auslesen). Der t - a_e -Graph (dieses t bezeichnet die Zeit) muss dafür die Gerade $a = 0.56$ bzw. $a = 0.505$ von unten schneiden. Aus der numerischen Evidenz schließen wir, dass das offensichtlich nur im endemischen Fall, aber nicht im epidemischen Fall passiert, siehe Abb. 7.15.

Wir fassen unsere Ergebnisse zusammen: Laut dem gewählten Modell kommt es zur Tannenwickler-Epidemie, weil ab einer gewissen Größe des Tannenbestandes die Wickler-

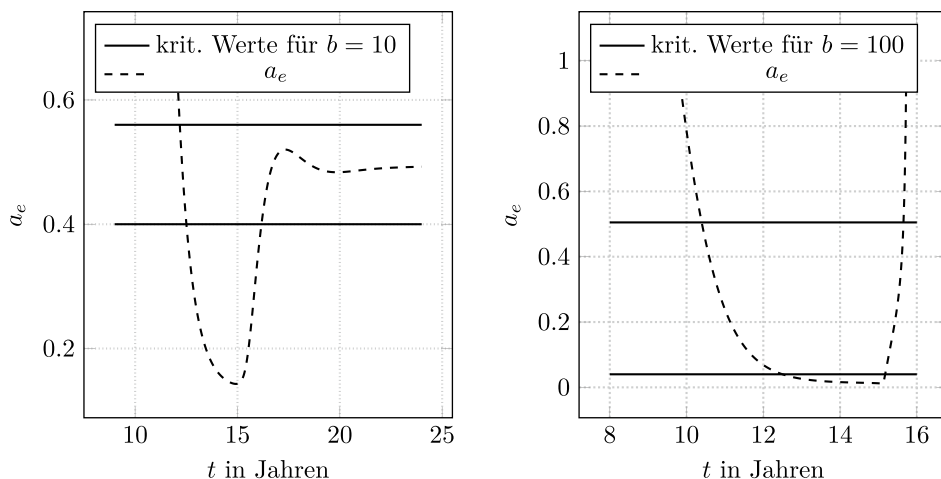


Abb. 7.15 Ein Ausschnitt aus dem t - a_e -Diagramm für die Anfangswerte (7.35). Links $T = 0.1$, rechts $T = 0.01$. Im ersten Fall verbleibt e im kranken Zustand, im zweiten Fall wechselt e in den gesunden Zustand

population von einem endemischen Zustand, in dem ihre Größe im Wesentlichen durch den Bejagungsdruck der Fressfeinde kontrolliert wird, in einen epidemischen Zustand wechselt, in dem die Größe durch das Nadelangebot und den Gesundheitszustand der Bäume limitiert ist. Der Übergang vom endemischen in den epidemischen Zustand vollzieht sich schnell. Im epidemischen Zustand wechselt der Gesundheitszustand der Bäume ab einem gewissen Besiedelungsgrad der Zweige durch Larven von gesund auf krank, wodurch der Zweigbestand und mit ihm der Wicklerbestand zusammenbrechen. Auch diese Veränderung vollzieht sich auf einer schnellen Zeitskala.

Abhängig vom gewählten Schwellenwert T , ab dem der Benadelungsgrad die Wicklerpopulation einschränkt, wechselt der Wicklerbestand in den endemischen Zustand zurück (wenn T hinreichend klein ist) oder verbleibt im epidemischen Zustand (wenn T hinreichend groß ist). Im ersten Fall kehrt der Benadelungszustand E in den gesunden Zustand zurück, der Tannenbestand baut sich neu auf und es kommt zu einem periodischen Durchlaufen des Zyklus. Im zweiten Fall verbleibt das System in einem stationären Kümmerzustand.

7.3.4 Bewertung des Modells

Die Tannenwickler verursachen enorme wirtschaftliche Schäden in den Wäldern Nordamerikas. Mathematische Modelle sind also sehr gefragt, mit denen sich Entwicklungen vorhersehen oder Auswirkungen von Interventionen simulieren lassen. Dabei werden teilweise Modelle mit ein paar Zehntausend Differenzengleichungen genutzt, welche die örtliche Ver-

teilung und die Altersstruktur der Bäume mit berücksichtigen, siehe z. B. [50]. Im Rahmen solcher Modelle lassen sich empirische Messdaten gut reproduzieren und darauf aufbauend umfangreiche Simulationen durchführen. Warum bestimmte Phänomene aber auftauchen und welche Einflüsse dafür maßgeblich sind, lässt sich aufgrund der Komplexität der betrachteten Modelle kaum erklären.

Im Gegensatz dazu kommt das hier benutzte Modell mit drei Variablen aus, hat also den Vorteil der Einfachheit. Qualitativ zeigt es das Phänomen des zyklischen Zusammenbrechens, allerdings mit einer zu großen Periode von über 60 Jahren, im Gegensatz zu beobachteten Zyklen von 35 bis 40 Jahren. Ein Teil dieser Diskrepanz lässt sich dadurch erklären, dass in der Natur die alten Bäume absterben und die neu gekeimten Bäume direkt mit optimaler Gesundheit $E = 1$ ins Leben treten. Im Modell benötigt der Bestand ca. zehn Jahre bis die Variable E auf den Wert 1 angewachsen ist.

Das vorliegende Modell liefert also keine numerisch zuverlässige Voraussage für empirische Daten, dafür aber einen Erklärungsrahmen, warum es zu dem zyklischen Verhalten kommt oder auch nicht. Es kann ein Verständnis dafür erzeugen, welche Variablen tatsächlich ausschlaggebend für die Hervorbringung des beobachteten Phänomens sind. (Die Ansichten über solche qualitativ erklärenden, aber nicht quantitativ vorhersagenden Modelle gehen weit auseinander.) Die Auswirkungen möglicher Interventionen gegen die Wickler, wie das Versprühen von Insektiziden oder das Aussetzen weiterer Fressfeinde, sollen im Rahmen dieses Modells in Aufg. 7.12 durchgespielt werden. Für eine weitergehende Diskussion dieses Modells siehe auch [32].

7.3.5 Exkurs: der Übergang vom stationären Punkt zum oszillierenden Verhalten

Wir haben in den vorherigen Abschnitten gesehen, dass der Charakter des Systems (7.34) sich mit kleiner werdendem T von einem stationären Endzustand hin zu einem oszillierenden Verhalten mit einer Periode von ca. 60 Jahren ändert. Es stellt sich die Frage, was am Übergang zwischen diesen beiden Verhaltensweisen passiert.

Um dieser Frage nachzugehen, tasten wir uns numerisch an einen T -Wert heran, an dem das Verhalten gerade umschlägt. Das scheint ungefähr bei $T = 0.0538$ der Fall zu sein. Für $T = 0.0538$ verfolgen wir die Geschichte des Anfangsdatums (7.35) numerisch zunächst über 150 Jahre, siehe Abb. 7.16. Während es zunächst so aussieht, also ob sich das System einem stationären Zustand annähert, setzen tatsächlich Oszillationen mit einer Periode von ca. acht Jahren und einer kleinen, aber anwachsenden Amplitude ein (man beachte den Wechsel in der Skalierung der vertikalen Achse auf Hundertstel nach 75 Jahren). Schaukeln sich dies Oszillationen weiter auf, bis p in den endemischen und e in den gesunden Zustand wechselt und das Spiel von vorne losgeht? Wenn wir die Lösung weitere 150 Jahre verfolgen, siehe Abb. 7.17, bemerken wir, dass das nicht der Fall zu sein scheint. Stattdessen scheint die Amplitude der Oszillationen selber zu oszillieren, und zwar mit einer Periode von ca. 40 Jahren. Was passiert in unserem asymptotischen Bild, in dem p und e nah an

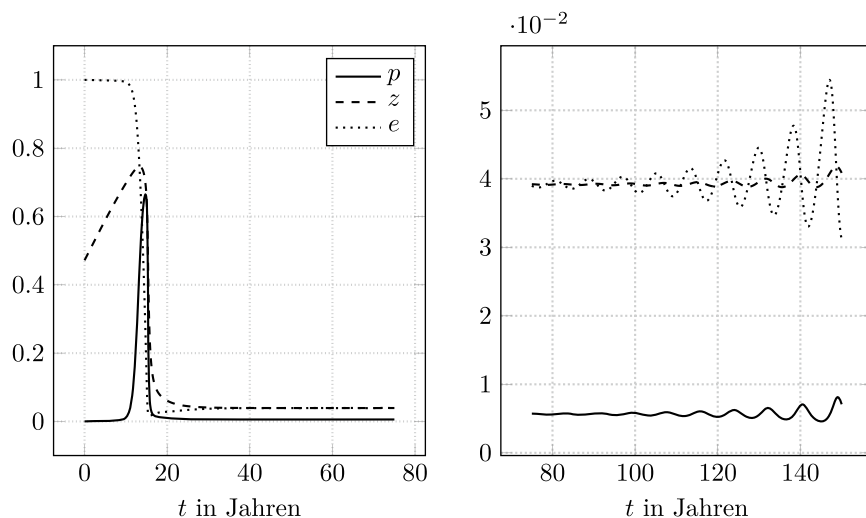


Abb. 7.16 Das Tannenwickler-Balsamtannen-System für $T = 0.0538$. Man beachte den Wechsel in der Skalierung der vertikalen Achse

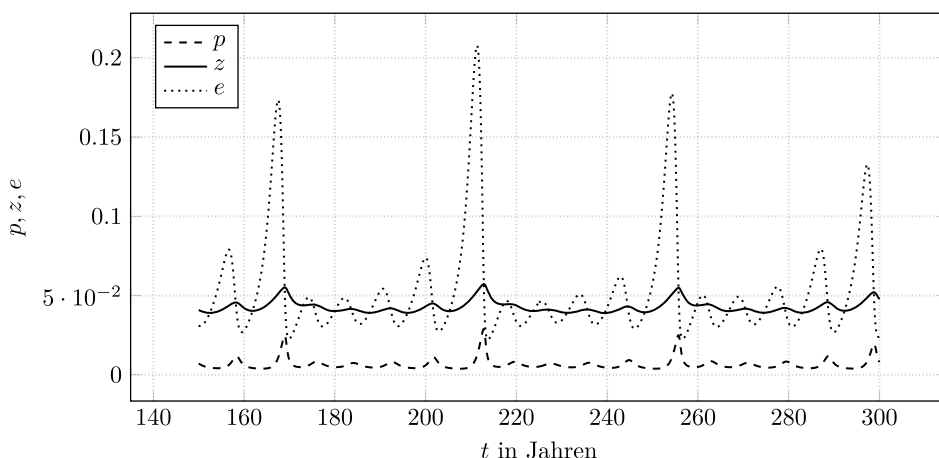


Abb. 7.17 Das Tannenwickler-Balsamtannen-System für $T = 0.0538$, die Jahre 150 bis 300

stationären Punkten sind und z langsam e folgt? Um diese Frage zu klären, schauen wir uns zunächst das b - a -Diagramm der ersten 60 Jahre an, siehe Abb. 7.18. Die p -Trajektorie scheint schließlich genau auf der kritischen Kurve zu verlaufen. Sehen wir uns die später einsetzenden Oszillationen genauer an, wir haben die Jahre 148–160 gewählt, sehen wir, dass die Trajektorie tatsächlich um die kritische Kurve herum oszilliert, siehe Abb. 7.19. Die p -Trajektorie schneidet den unteren Ast der kritischen Linie von oben nach unten, kommt also für ca. drei Jahre in den Einzugsbereich des endemischen Zustands, müsste eigentlich

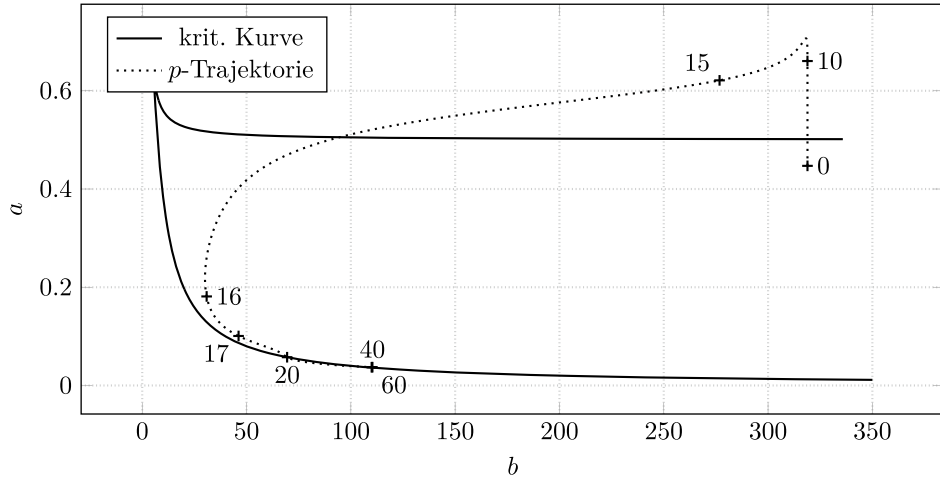


Abb. 7.18 Der Weg der p -Trajektorie in den ersten 60 Jahren. Aus einem epidemischen Zustand kommend scheint die Trajektorie genau auf der kritischen Kurve zu landen, $T = 0.0538$

in diesen kippen und in diesem beim neuerlichen Überschreiten der kritischen Kurve in der entgegengesetzten Richtung verbleiben. Offensichtlich ist p aber während der gesamten Entwicklung nicht endemisch. Wie lässt sich das erklären?

Betrachten wir das t - a_e -Diagramm desselben Zeitraums, sehen wir ein analoges Verhalten, der Graph oszilliert um den kritischen Wert für $b = 1/0.0538$, siehe Abb. 7.20. Auch

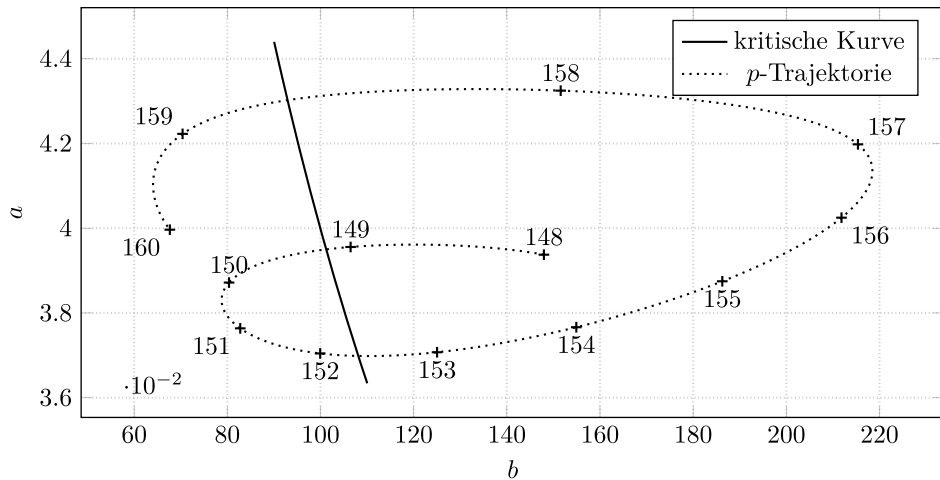


Abb. 7.19 Oszillation im b - a -Diagramm um den unteren Ast der kritischen Kurve mit einer Periode von knapp zehn Jahren, $T = 0.0538$. Man beachte die Skalierung der a -Achse in Hundertstel

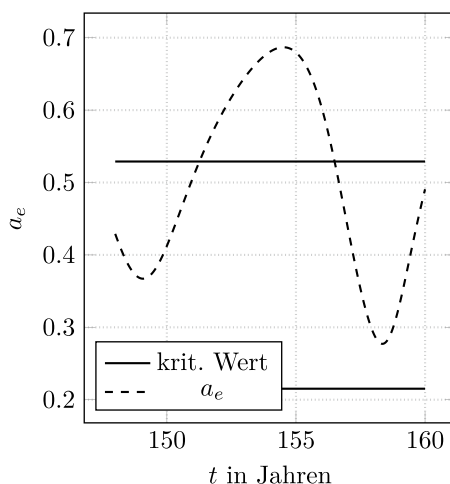
hier ist e ca. zwischen den Jahren 152 und 156 im Einzugsbereich des gesunden Zustands, müsste eigentlich in diesen kippen und in diesem beim neuerlichen Überschreiten des kritischen Wertes verbleiben. Offensichtlich ist e aber während der gesamten Entwicklung nicht im gesunden Zustand. Wie kann das sein?

Wir diskutieren die Situation für die Variable p : Zur Zeit $t = 148$ Jahre befindet sich das System im Zustand $p = 0.0069$, $z = 0.042$, $b = 148.0$, $a = 0.039$ und damit $p/z = 0.167$. Die Werte entnehmen wir unserer Tabellenkalkulation. Es hängen also ca. 50 Larven an jedem Zweig, genauer $p/\alpha z = 53.3$. Die stationären Zustände für die Variable $p/\alpha z$ für die genannten Parameter a und b können wir aus einem GeoGebra-Plot der Funktionen in Abb. 7.12 mit entsprechend eingestellten Parametern a und b ablesen, es sind die x -Koordinaten der Schnittpunkte (natürlich könnte man die Werte für b und a auch in die algebraische Lösungsformel einsetzen): $x_1^* = 0.040$, $x_2^* = 32.50$ und $x_3^* = 115.36$. Weil $x_2^* = 32.50 < p/\alpha z$, befindet sich $p/\alpha z$ zu diesem Zeitpunkt im Anziehungsbereich von x_3^* und hat eine positive Änderungsrate. Bei der Reise durch das b - a -Diagramm verschwindet der stationäre Punkt x_3^* und $p/\alpha z$ kommt in den Einzugsbereich des (einzig verbliebenen) Punktes x_1^* . Warum verbleibt $p/\alpha z$, obwohl der stationäre Punkt x_1^* zu keiner Zeit verschwindet, nicht in dessen Einzugsbereich?

Wir schauen uns die Zeit $t = 153$ Jahre an. Zu dieser Zeit liefern unsere numerischen Berechnungen die Werte $p = 0.00434$, $z = 0.0391$, $b = 125.02$, $a = 0.0371$ und damit $p/\alpha z = 35.34$. Für die stationären Punkte zu diesen Parameterwerten b und a erhalten wir $x_1^* = 0.0371$, $x_2^* = 39.274$ und $x_3^* = 85.71$. Es gilt $p/\alpha z < x_2^*$ und damit ist $p/\alpha z$ zu diesem Zeitpunkt im Anziehungsbereich von x_1^* und hat eine negative Änderungsrate.

Wie sieht es nun zwei Jahre später aus, $t = 156$ Jahre? Hier haben wir $p = 0.00430$, $z = 0.0409$, $b = 186.25$, $a = 0.0388$ und $p/\alpha z = 33.74$. Für die stationären Punkte erhalten wir $x_1^* = 0.0388$, $x_2^* = 30.96$ und $x_3^* = 155.25$, also $x_2^* < p/\alpha z$. Obwohl $p/\alpha z$

Abb. 7.20 Oszillation von a_e um den oberen kritischen Wert mit einer Periode von knapp zehn Jahren, $T = 0.0538$



keine kritische Kurve gekreuzt hat, ist es aus dem Anziehungsbereich des endemischen Zustands in den Anziehungsbereich des epidemischen Zustands gewechselt.

Damit ist das gesamte asymptotische Konzept der Unterteilung in schnelle Variablen, die zwischen den stationären Zuständen endisch/epidemisch bzw. gesund/krank wechseln, und einer langsamen z -Variablen, die der e -Variablen folgt, in diesem kritischen Grenzfall $T = 0.0538$ nicht mehr anwendbar. Der Weg durch das b - a -Diagramm hat keine besondere Aussagekraft mehr, weil sich der Zustand $p/\alpha z$ gar nicht in der Nähe der stabilen stationären Punkte x_1^* oder x_3^* aufhält, sondern in der Nähe des instabilen Punkts x_2^* , wenn es ihn denn gibt.

Aufgaben

7.1 Das Phasenporträt für das Galilei'sche Modell des senkrechten Wurfs

Wir betrachten das Galilei'sche Modell des senkrechten Wurfs:

$$\begin{aligned}\dot{h} &= v \\ \dot{v} &= -g,\end{aligned}\tag{7.45}$$

dabei ist g die Erdbeschleunigung.

- Bestimmen Sie ein erstes Integral $H = H(x, v)$, das konstant auf allen Lösungen von (7.45) ist.
- Zeichnen Sie das Phasenporträt von (7.45).
- Ein Körper wird aus der Höhe $h_0 = 0$ mit der Geschwindigkeit $v_0 > 0$ abgeworfen. Berechnen Sie, welche maximale Höhe der Körper erreicht, einmal durch Verwendung der Bewegungsgleichung (7.1) und einmal mit Hilfe des in Aufgabenteil a) bestimmten ersten Integrals.

7.2 Die Energie einer senkrecht geworfenen Kugel im Newton'schen Gravitationsfeld

Zeigen Sie, dass die Funktion

$$E(h, v) = \frac{1}{2}v^2 - \frac{Gm_E}{R_E + h}$$

ein erstes Integral für das Newton'sche Modell

$$\begin{aligned}\dot{h} &= v \\ \dot{v} &= -\frac{Gm_E}{(R_E + h)^2};\end{aligned}$$

des senkrechten Wurfs liefert.

7.3 Alter Trick

In dieser Aufgabe soll ein bekannter Trick aus der Analysis 1 in Erinnerung gerufen werden, um die Formel (7.13) herzuleiten.

- Drücken Sie die Funktion $z \mapsto \frac{1}{1+z}$ für $|z| < 1$ mit Hilfe der geometrischen Reihe aus.
- Folgern Sie (7.13) durch gliedweises Differenzieren.

7.4 Eine asymptotische Entwicklung zur Übung

Wir betrachten die Schar von AWP

$$\ddot{x}_\varepsilon = -\varepsilon \dot{x}_\varepsilon - 1 \quad x_\varepsilon(0) = 0 \quad \dot{x}_\varepsilon(0) = 1.$$

für $\varepsilon > 0$.

- Berechnen Sie für diese Schar die asymptotische Entwicklung bis einschließlich der Ordnung zwei.
- Untersuchen Sie das Verhältnis der Lösungen und der asymptotischen Lösungen numerisch mit GeoGebra.
- Beweisen Sie analytisch, dass die Approximation erster Ordnung die tatsächlichen Lösungen bis auf einen Fehler der Ordnung ε^2 approximiert.

7.5 Die asymptotische Abschätzung erster Ordnung

Wir definieren die Approximation erster Ordnung zum Newton-Modell durch

$$x_\varepsilon^I = x_0 + \varepsilon x_1, \quad y_\varepsilon^I = y_0 + \varepsilon y_1,$$

wobei (x_0, y_0) und (x_1, y_1) gemäß (7.14) bzw. (7.16) mit den zugehörigen Lösungen definiert sind. Beweisen Sie die folgende asymptotische Abschätzung: Für alle $\varepsilon \geq 0$ und alle $0 \leq \tau \leq 2$ gilt

$$|x_\varepsilon(\tau) - x_\varepsilon^I(\tau)| \leq 9\varepsilon^2; \quad (7.46)$$

dabei ist x_ε die Lösung des Newton-Systems (7.11).

7.6 Fehlerschranken in konkreten Situationen

Berechnen Sie Fehlerschranken für die nullte und erste Approximation zum Newton-Modell des senkrechten Wurfs mit dem zugehörigen Zeitintervall, auf dem die Abschätzung gültig ist für

- $v_0 = 10 \text{ m s}^{-1}$. Das entspricht einem kräftigen Wurf eines Balls.
- $v_0 = 100 \text{ m s}^{-1}$.
- $v_0 = 10 \times 10^4 \text{ m s}^{-1}$. Das entspricht ungefähr der Geschwindigkeit, mit der ein Körper von der Erde entfliehen kann, also nicht mehr auf die Erde zurückfällt.

Nutzen Sie dazu die Abschätzungen (7.20) bzw. (7.46).

7.7 Beweis der asymptotischen Abschätzung zweiter Ordnung für den senkrechten Wurf

In dieser Aufgabe soll die asymptotische Abschätzung zweiter Ordnung bewiesen werden. Dabei ist x_ε eine Lösungsschar für (7.11) und x_ε^Π ist in (7.17) definiert.

- a) Bestimmen Sie eine Konstante $c_2 > 0$, so dass für alle $0 \leq \tau \leq 2$ und alle $0 \leq \varepsilon$ die folgende Ungleichung gilt:

$$|x_\varepsilon(\tau) - x_\varepsilon^\Pi \tau| \leq C_2 \varepsilon^3$$

- b) Schätzen Sie den Fehler ab, der entsteht, wenn die Lösung des Newton-Modells durch die Lösung der asymptotischen Entwicklung der Ordnung zwei ersetzt wird, wenn von folgenden Parametern ausgegangen wird: $R = 10 \times 10^7 \text{ m}$, $g = 10 \text{ m s}^{-2}$, $v_0 = 1000 \text{ m s}^{-1}$.

7.8 Die Fahnenstange bei niederfrequenten Erdbeben

In Abschn. 1.3 wurde die Modellierung der Bewegung einer Fahnenstange bei Erdbeben vorgestellt, die wir hier kurz wiederholen: Das Beben wird durch eine sinusförmige horizontale Bewegung des Bodens mit der Frequenz ω modelliert, und für die Auslenkung $u(t, x)$ der Fahnenstange zur Zeit t in der Höhe x werden die folgenden Gleichungen und Randwerte angenommen:

- $\rho \frac{\partial^2 u}{\partial t^2} + Ek^2 \frac{\partial^4 u}{\partial x^4} = 0$ für $0 < x < l$.
- u ist periodisch in der Variablen t mit Periode $T = 2\pi/\omega$.
- u erfüllt die Randbedingungen

$$u(t, 0) = a \sin(\omega t), \quad \frac{\partial u}{\partial x}(t, 0) = 0, \quad (7.47)$$

$$\frac{\partial^2 u}{\partial x^2}(t, l) = 0, \quad \frac{\partial^3 u}{\partial x^3}(t, l) = 0. \quad (7.48)$$

Dabei ist a die Amplitude und ω die Frequenz des Bebens, ρ die Massendichte und E der Young'sche Modul des Fahnenstangenmaterials, l die Länge und k der Trägheitsradius der Fahnenstange. Eine Entdimensionalisierung des Modells lautet mit den dimensionslosen Variablen $\tau = \omega t$, $\xi = x/l$ und $v = u/a$:

- $\frac{\partial^2 v}{\partial \tau^2} + J \frac{\partial^4 v}{\partial \xi^4} = 0$ für $0 < \xi < 1$.
- v ist periodisch in der Variablen τ mit Periode 2π .
- v erfüllt die Randbedingungen

$$v(\tau, 0) = \sin(\tau), \quad \frac{\partial v}{\partial \xi}(\tau, 0) = 0, \quad (7.49)$$

$$\frac{\partial^2 v}{\partial \xi^2}(\tau, 1) = 0, \quad \frac{\partial^3 v}{\partial \xi^3}(\tau, 1) = 0, \quad (7.50)$$

mit dem dimensionslosen Parameter $J = \frac{Ek^2}{\rho\omega^2 l^4}$.

- Berechnen Sie die asymptotische Entwicklung für $J \gg 1$, also $1/J \ll 1$, und beschreiben Sie die dahinterliegende physikalische Situation.
- Probieren Sie eine asymptotische Entwicklung für $J \ll 1$.

7.9 Der senkrechte Wurf mit Luftwiderstand

Ein Körper der Masse m wird von der Erdoberfläche mit der Geschwindigkeit v_0 senkrecht nach oben geworfen. Es soll das Galilei'sche Modell benutzt werden, diesmal aber zusätzlich der Luftwiderstand berücksichtigt werden: Der Luftwiderstand soll mit dem *Stoke'schen Gesetz* $F_R = -cv$ für den Strömungswiderstand in viskosen Fluiden berücksichtigt werden, das für kleine Geschwindigkeiten sinnvoll ist. Dabei ist c ein von der Gestalt des Körpers abhängiger Koeffizient. Die Bewegung hänge von der Masse m , der Anfangsgeschwindigkeit v_0 sowie dem Reibungskoeffizienten c mit der Dimension $[c] = M/T$ ab. M sei dabei die Abkürzung für die Dimension Masse.

- In der Physik versucht man häufig schon allein aus den in einer physikalischen Situation auftretenden Parametern geeignete Skalen zu bestimmen (ohne Kenntnis bzw. Nutzung einer Differentialgleichung). Bestimmen Sie jeweils α , β , γ und δ so, dass
 - $\varepsilon = m^\alpha c^\beta g^\gamma v_0^\delta$ dimensionslos ist.
 - $\bar{l} = m^\alpha c^\beta g^\gamma v_0^\delta$ eine Länge ist.
 - $\bar{t} = m^\alpha c^\beta g^\gamma v_0^\delta$ eine Zeit ist.
- Wir betrachten das Anfangswertproblem für die Höhe des Körpers

$$\dot{h} = v, \quad \dot{v} = -cv - mg, \quad h(0) = 0, \quad v(0) = v_0. \quad (7.51)$$

Entdimensionalisieren Sie die Differentialgleichung auf alle verschiedenen vernünftigen Arten. Es gibt wieder mehrere Möglichkeiten.

- Es gelte $\beta := \frac{cv_0}{mg} \ll 1$. Untersuchen Sie, mit welchen Skalen Sie zu einem sinnvollen asymptotischen Modell für $\beta = 0$ kommen. Entwickeln Sie das Modell bis zur Ordnung β^2 und berechnen Sie die zugehörigen Lösungen.
- Beweisen Sie, dass die Approximation erster Ordnung die tatsächliche Lösung mit einem Fehler der Größenordnung β^2 approximiert.

7.10 Asymptotische Verhältnisse im Phasenporträt

Führen Sie die Diskussion, die zu Abb. 7.9 führt, mit Hilfe der schnellen Zeitvariablen T .

7.11 Numerische Untersuchung des Wickler-Tannen-Systems

Implementieren Sie das explizite Euler-Verfahren für das dimensionslose Wickler-Tannen-System (7.34) in ein Tabellenkalkulationsprogramm. Gehen Sie dabei so vor, dass Sie die dimensionsbehafteten Parameter vorgeben und leicht variieren können und das Programm die dimensionslosen Parameter berechnet. Gestalten Sie das Blatt so, dass Sie auch die Anfangswerte und die benutzte Schrittweite bequem variieren können.

- a) Reproduzieren Sie die im Text dargestellten numerischen Ergebnisse. Variieren Sie außerdem
 - die Anfangsdaten,
 - die Schrittweite.

Erklären Sie, warum der Algorithmus in die Knie geht, wenn Sie die Schrittweiten zu groß wählen. Variieren Sie die Schrittweite und versuchen Sie die beobachteten Effekte zu erklären.

- b) Überlegen Sie, wie sich eine Änderung einzelner dimensionsbehafteter Parameter auf das System auswirkt. Testen Sie Ihre Vermutung numerisch. Versuchen Sie das Verhalten nach Möglichkeit zu erklären. In einigen Situationen kann die asymptotische Betrachtung aus dem Text anwendbar sein, in anderen eher nicht.
- c) Tasten Sie sich insbesondere „von unten“ an den kritischen Wert von T heran. Erklären Sie das Verhalten der Lösung im Rahmen der asymptotischen Näherung.
- d) Erläutern Sie, welche Parameter Sie wie ändern müssen, damit der Fehler, den Sie durch die Anwendung der asymptotischen Näherung machen, kleiner wird.

7.12 Variationen des Wickler-Tannen-Systems und Modellierungen von Bekämpfungsstrategien

Variieren Sie das Wickler-Tannen-System und untersuchen Sie das neue System so weit wie möglich:

1. Ersetzen Sie das logistische Wachstumsmodell für den Grad der Benadelung/die Gesundheit E durch ein Modell mit beschränktem Wachstum: $\dot{E} = r_E(1 - E)$.
2. Modellieren Sie das Sprühen von Insektiziden, indem Sie eine weitere Sterberate $-r_i P$ zur Wicklergleichung hinzufügen.
3. Zusätzliche Fressfeinde werden ausgesetzt.
4. Zusätzliche Fressfeinde werden ausgesetzt, die ohne Sättigungseffekt jagen. (Man könnte auch sagen, dass Wickler manuell abgesammelt werden.)
5. Es werden sterile Wickler ausgesetzt.

6. Es werden Insektizide versprüht, aber nur wenn der Wicklerbestand eine gewisse Größe überschreitet.
7. Die Bäume werden so besprüht, dass die Larven die Nadeln schlechter fressen können.
8. Denken Sie sich weitere Bekämpfungsstrategien aus und modellieren Sie diese im Rahmen des dargestellten Modells.



Das Modell der Reizverarbeitung in Nervenzellen nach Hodgkin und Huxley – Experimente, Modelle und Spielzeugmodelle

8

Inhaltsverzeichnis

8.1 Die Nervenzelle und das Aktionspotential	272
8.2 Erklärungsansätze vor Hodgkin und Huxley	274
8.3 Die Versuchsreihe von Hodgkin und Huxley	274
8.4 Die mathematische Modellierung und die quantitative Berechnung der Aktionspotentials	298
8.5 Toy-Models – das vereinfachte Modell nach FitzHugh-Nagumo	309
8.6 Vom Spielzeugmodell zum großen Modell: neuronales Feuern im Modell von Hodgkin und Huxley	317
Aufgaben	318

In diesem Kapitel stellen wir ein berühmtes Beispiel für eine gelungene mathematische Modellierung aus der Mitte des 20. Jahrhunderts vor: die Modellierung der Bildung von Aktionspotentialen in Nervenzellen durch Alan Lloyd Hodgkin und Andrew Fielding Huxley. Wir kürzen die beiden Namen ab jetzt durch H & H ab. In einer Serie von fünf Veröffentlichungen [40–43, 45] entwickeln sie H & H eine physiologische Deutung des *Aktionspotentials*, eines typischen Spannungsverlaufs an der Membran von Nervenzellen, stellen auf dieser Basis ein mathematisches Modell auf und zeigen, dass die numerisch aus dem Modell gewonnenen Ergebnisse die experimentellen Befunde gut widerspiegeln. Diese Arbeiten legten einen wichtigen Grundstein für die moderne Elektrophysiologie und wurden 1963 mit dem Nobelpreis für Medizin und Physiologie geehrt. In Abschn. 8.1 bis 8.4 wollen wir diese Entwicklung von den zugrunde liegenden Fragestellungen über die Experimente bis zum mathematischen Modell, einem Differentialgleichungssystem mit vier Variablen, nachvollziehen, das Differentialgleichungssystem numerisch implementieren und die numerischen Ergebnisse von H & H reproduzieren.

Das aufgestellte mathematische Modell bezieht sich konkret auf spezielle Nervenzellen von Tintenfischen. In anderen Nervenzellen, auch von anderen Spezies, liegen die physiologischen Verhältnisse durchaus anders, trotzdem kommt es zur Ausbildung typischer

Aktionspotentiale. Vom Standpunkt des mathematischen Modellierens stellt sich daher die Frage, welche Eigenschaften ein mathematisches Modell haben muss, um Phänomene vom Typ „Aktionspotential“ generieren zu können. Insbesondere stellt sich die Frage, wie ein möglichst einfaches Modell aussehen könnte, das dann nicht nur numerischen, sondern auch analytischen Untersuchungen zugänglich wäre. Man spricht in diesem Zusammenhang auch von *Toy-Models*. In Abschn. 8.4 stellen wir das vereinfachte Modell nach FitzHugh vor und studieren es mit den bereits vertrauten Methoden der Phasenporträt-Analyse und asymptotischen Grenzwerten. Es stellt sich heraus, dass dieses Spielzeugmodell nicht nur Aktionspotentiale produzieren kann, sondern auch weitere Phänomene enthält, die sich als *biochemischer Schalter* und *neuronales Feuern* interpretieren lassen. Angeregt durch diese Befunde machen wir uns in Abschn. 8.6 auf die Suche nach vergleichbaren Phänomenen im „großen“ Modell von H & H.

8.1 Die Nervenzelle und das Aktionspotential

Worum geht es? Nervenzellen sind in einem großen Netzwerk, dem neuronalen Netz, zusammengeschlossen, siehe [13]. Dabei müssen Signale von einer Nervenzelle zur nächsten übertragen werden. Nervenzellen bestehen in der Regel aus einem Zellkörper, der den Zellkern enthält, und zwei verschiedenen Arten von Fortsätzen, über die mit den benachbarten Zellen kommuniziert wird. Das sind zum einen die sogenannten Dendriten, welche die Signale von benachbarten Zellen aufnehmen, und zum anderen die Axone, welche die Signale an die weiteren Nervenzellen weitergeben (siehe Abb. 8.1). Grob gesprochen nimmt die Zelle über die Dendriten Signale von den Nachbarzellen auf. Diese Signale laufen in Form von zeitlich veränderlichen elektrischen Spannungen zwischen der Innen- und der Außenseite der Zellmembran, der Grenzfläche zwischen dem Inneren und dem Äußeren der Zelle, entlang. Die Spannungen von verschiedenen Dendriten summieren sich am Übergang zum Axon, dem

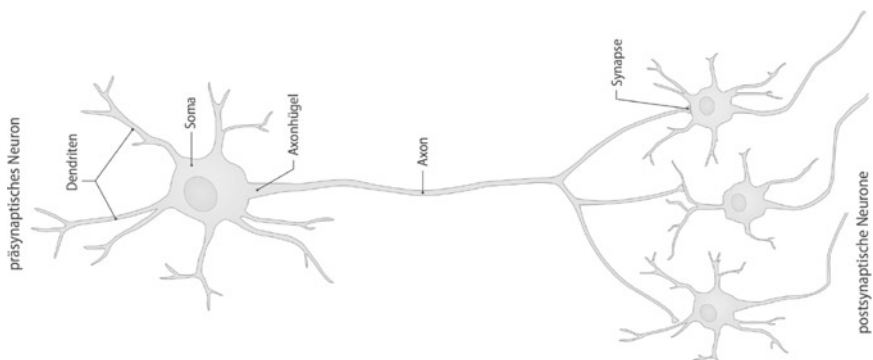


Abb. 8.1 Der Aufbau einer menschlichen Nervenzelle (Aus [22, Abb. 5.2], ©Springer-Verlag GmbH Deutschland 2019)

Axonhügel, und lösen dort – aber nur, wenn sie einen bestimmten Wert überschreiten – ein sogenanntes Aktionspotential aus. Ein Aktionspotential ist dabei ein typischer Spannungsverlauf an der Axonmembran, der wie eine Welle das Axon entlangläuft und nun wiederum die Dendriten von benachbarten Zellen reizen kann. Wenn der Spannungsimpuls, der am Axonhügel ankommt, einen kritischen Wert nicht überschreitet, wird kein Aktionspotential in Gang gesetzt, der Spannungspike klingt wieder ab. Überschreitet der Spannungsimpuls jedoch den kritischen Wert, durchläuft das Aktionspotential vier Phasen (siehe Abb. 8.2). Die im Folgenden angegebenen konkreten Spannungswerte beziehen sich auf menschliche Nervenzellen.

- (I) In der Initiationsphase treibt ein Reiz (die von den Dendriten zusammenkommenden Spannungen) die im Ruhezustand negative Spannung von ca. -70 mV zwischen der Innenseite des betrachteten Membranstücks und der Außenseite über einen bestimmten Schwellenwert von circa -40 mV. Diese Anregung kann schnell oder langsam erfolgen. (Bleibt die Spannung unter dem Schwellenwert, klingt sie einfach wieder auf den Ruhewert ab.)
- (II) In der zweiten Phase erhöht sich die Spannung auf circa $+30$ bis $+50$ mV: Dieser Anstieg vollzieht sich schnell in circa 1 ms, man spricht vom *Aufstrich* oder *Upstroke*.
- (III) In der dritten Phase fällt die Spannung wieder ab, man spricht von der *Repolarisationsphase*. Diese Repolarisation vollzieht sich in der Regel langsamer als der Upstroke.
- (IV) In der vierten Phase unterschreitet die Spannung die Ruhespannung und kehrt dann sehr viel langsamer in die Ruhespannung zurück, man spricht von der *Hyperpolarisation*. In dieser Zeit ist die Zelle nicht erregbar.

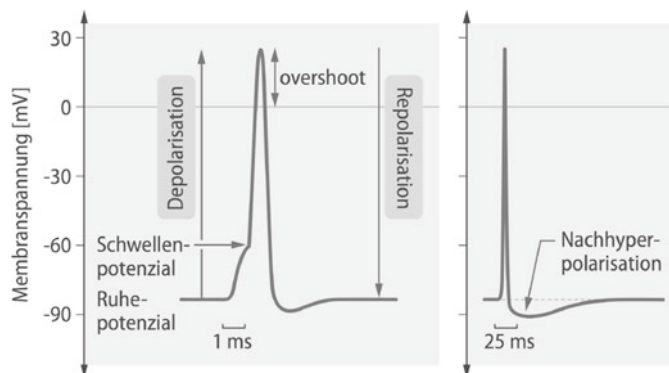


Abb. 8.2 Die vier Phasen eines Aktionspotentials nach [25, Abb.6.4], (Abdruck mit freundlicher Genehmigung der ©Springer-Verlag GmbH Deutschland, 2019)

8.2 Erklärungsansätze vor Hodgkin und Huxley

Zur Zeit der Versuche von H & H waren der prinzipielle Aufbau der Nervenzelle und die Phänomenologie des Aktionspotentials, wie im vorherigen Abschnitt beschrieben, bereits bekannt. Ferner wusste man, dass das Axon im Ruhezustand negativ polarisiert ist sowie im Inneren eine höhere Konzentration an Kaliumionen als im Äußeren und im Inneren eine niedrigere Konzentration an Natriumionen als im Äußeren aufweist. Sowohl Kalium- als auch Natriumionen sind positiv geladene Ionen. Die einzelnen Konzentrationen und auch die Spannungsverläufe an der Zellmembran konnten bereits experimentell gemessen werden. Unklar war die Natur und das Zustandekommen der Ströme durch die Zellmembran, die zur Ausbildung von Aktionspotentialen führen. Eine erste Theorie von Bernstein, siehe [8], legt nahe, dass die Ausbildung des Ruhepotentials durch eine selektive Leitfähigkeit der Zellmembran für Kalium zustande komme und diese Selektivität bei der Ausbildung eines Aktionspotentials zusammenbräche. Es bleibt dabei allerdings unklar, warum die Zelle sich dann sogar positiv auflädt. Zur Zeit der Versuche von H & H hatten sich Hinweise darauf verdichtet, dass der Strom, der zum Upstroke führt, im Wesentlichen aus Natriumionen besteht und der Strom, der die Repolarisation veranlasst, aus Kaliumionen. Die Mechanismen, die den Stromfluss ermöglichen oder behindern, waren aber unklar. Warum wird die Membran durch einen Reiz auf einmal durchlässig für Natrium? Welche Natur hat der Transport? Fließen die Ionen einfach „passiv“ aufgrund des elektrochemischen Antriebs oder werden sie „aktiv“ von bestimmten Trägerteilchen durch die Membran getragen, wie es zunächst von H & H in [44] vermutet wurde?

8.3 Die Versuchsreihe von Hodgkin und Huxley

8.3.1 Beschreibung des Versuchsaufbaus

Zur Untersuchung der Ströme durch die Zellmembran eines Axons einer Nervenzelle verwenden H & H eine von Marmont [62] und Cole [18] kurz vorher entwickelte Technik, die wir in diesem Abschnitt skizzieren werden, die sogenannte *Voltage-Clamp*-Technik. Genaue experimentelle Details finden sich in [45].

Die Axone von menschlichen Nervenzellen haben lediglich einen Durchmesser von $0.08\text{ }\mu\text{m}$ bis $20\text{ }\mu\text{m}$ und waren in den 40er Jahren des letzten Jahrhunderts zu klein für die geplanten Experimente. Aus diesem Grund wichen die Forscher auf das Riesenaxon des Tintenfischs *Loligo ferosi* aus, das einen Durchmesser von ungefähr $600\text{ }\mu\text{m} = 0.6\text{ mm}$ aufweist. In diesen Größenordnungen konnte auch damals schon gut experimentiert werden.

Die Versuchsanordnung ist in Abb. 8.3 schematisiert dargestellt. In ein präpariertes Axon werden zwei auf einer Glaskapillare aufgebrachte, voneinander isolierte Silberdrähte eingebracht. Der eine Draht, in Abb. 8.3 ist das Draht *a*, dient dazu, von außen einen Elektronenstrom in das Axon führen zu können, der zweite Draht *b*, dient dazu, die elektrische

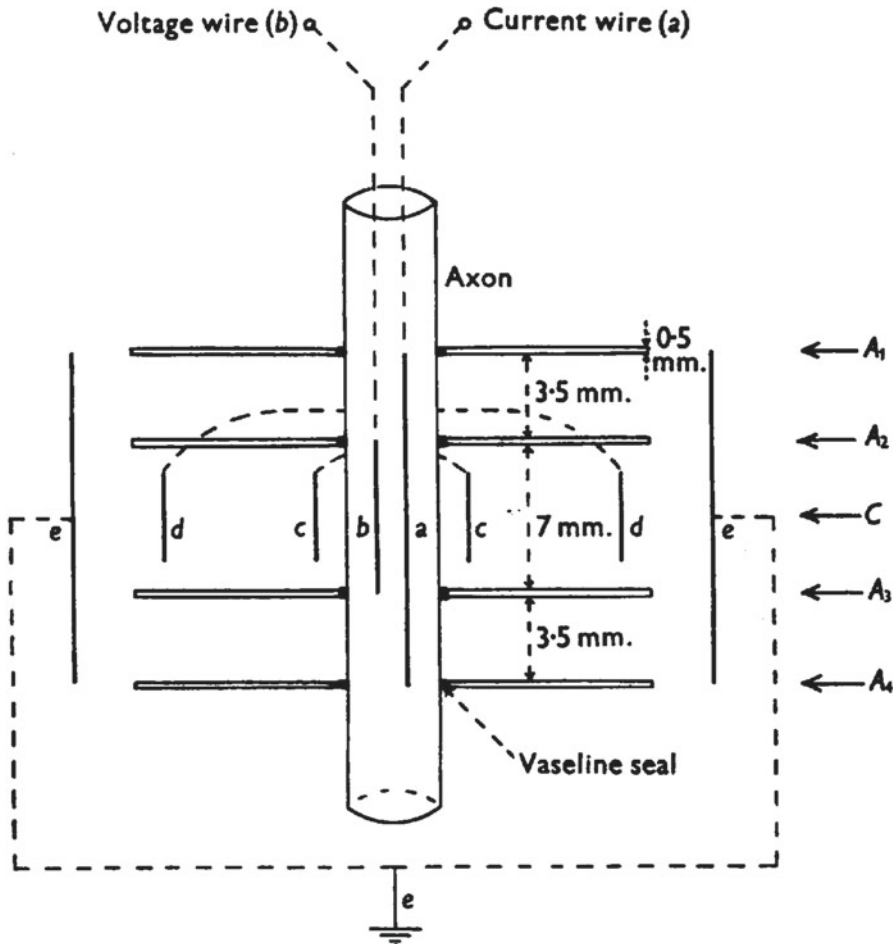


Abb. 8.3 Anordnung der inneren und äußeren Elektroden: a , b , c und d sind Elektroden, A_1 , A_2 , A_3 und A_4 sind Plexiglasunterteilungen. Mit Schellack isolierte Elektrodenteile sind gestrichelt dargestellt. Der Abstand zwischen den Elektroden c und b beträgt 2 mm, zwischen c und d 10 mm sowie zwischen d und e 8 mm (Abbildung aus [45, Fig. 1], ©(1952) John Wiley & Sons)

Spannung, die an der Axonmembran zwischen dem Inneren und dem Äußeren des Axons anliegt, zu messen. Das so präparierte Axon wird in einen Plexiglasbehälter eingebracht, in dem drei weitere Elektroden (also Silberdrähte) c , d und e angebracht sind und der mit verschiedenen Flüssigkeiten befüllt und auf einer konstanten Temperatur gehalten werden kann.

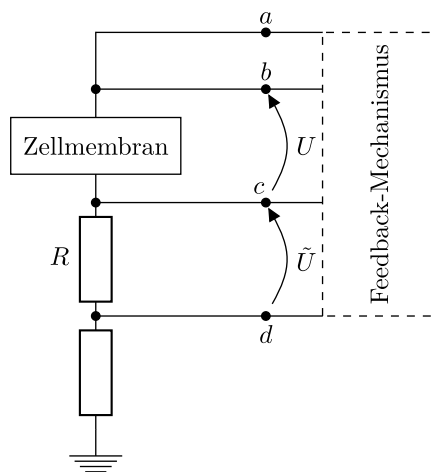
Über die Elektroden b und c wird die an der Membran anliegende Spannung gemessen. Dabei wählen H & H die Orientierung der Spannungsmessung so, dass die Ruhespannung eine positive Größe ist: Die Spannung U ist definiert als Außenpotential minus Innenpo-

tential. Eine vereinfachte Schaltskizze des Versuchsaufbaus ist in Abb. 8.4 dargestellt. Der Stromfluss wird aus dem Spannungsabfall $\tilde{U}$ zwischen den Elektroden c und d nach der Formel $I = \frac{\tilde{U}}{R}$ bestimmt, wobei R der elektrische Widerstand der 10 mm langen Strecke durch die Umgebungsflüssigkeit von der Elektrode d zur Elektrode c ist. Beispielsweise rechnen H & H hier mit einem Widerstand von $R = 74 \Omega$ bei Meerwasser mit einer Temperatur von 20°C . Mit dieser Orientierung ergibt ein Strom von positiv geladenen Ionen von außen nach innen eine positive Stromstärke.

Wesentlich für die Experimente ist der hier lediglich angedeutete Feedback-Mechanismus, eine recht komplizierte elektrische Schaltung, die es erlaubt, die Spannung U an der Zellmembran zu fixieren, das sogenannte *Voltage-Clamp*. Baut sich beispielsweise ein einwärts gerichteter Strom positiv geladener Ionen auf, so sorgt der Feedback-Mechanismus dafür, dass über die Elektrode a ein solcher Elektronenstrom in die Zelle geführt wird, sodass die anliegende Spannung konstant bleibt. Eine detaillierte Beschreibung des Feedback-Mechanismus befindet sich in [45, Abb. 6].

In einem ersten Versuch testen H & H, ob das so präparierte Axon noch die typischen elektrochemischen Eigenschaften hat und in der Lage ist, ein Aktionspotential auszubilden. Dazu wird ohne den Feedback-Mechanismus die Membran durch kurze Stromstöße über den Draht a depolarisiert. Die experimentellen Ergebnisse sind positiv, einige der aufgenommenen Messkurven sind in Abb. 8.5 dargestellt. Die präparierten Axone haben Ruhepotentiale von durchschnittlich $U_r = 60 \text{ mV}$ und bilden Aktionspotentiale mit einer Amplitude von ca. 110 mV aus, der Schwellenwert für die Ausbildung eines Aktionspotentials liegt bei einer Auslenkung von ca. -12 mV aus der Ruhespannung; sogenannte Hyperpolarisationen führen zu keinen Aktionspotentialen.

Abb. 8.4 Vereinfachte Schaltskizze des Versuchsaufbaus



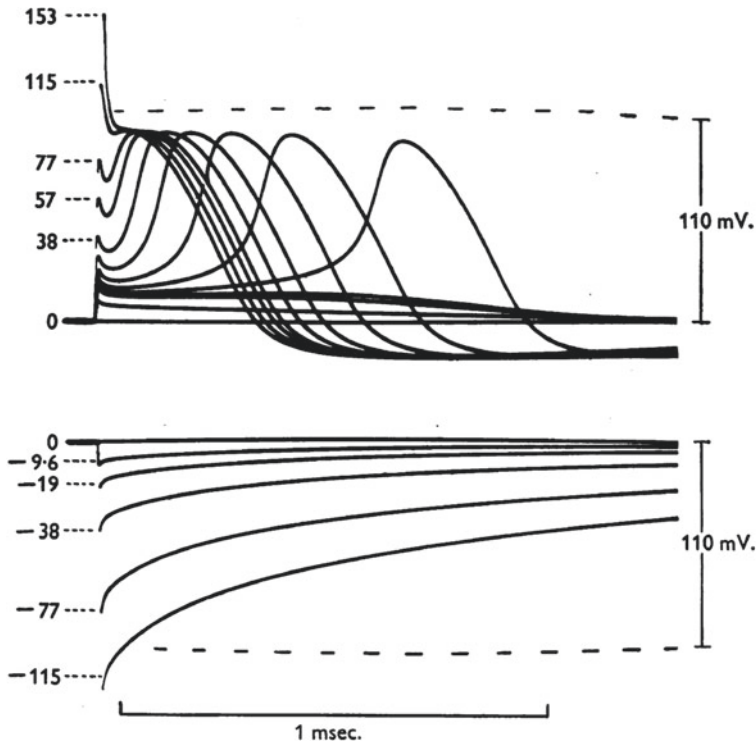


Abb.8.5 Gemessene Aktionspotentiale an präparierten Axonen. Die Zahlen an den Kurven geben die anfängliche Störung der Spannung aus der Ruhespannung an. Die Spannungen sind in mV angegeben. Unglücklicherweise ist in dieser Abbildung die Spannung andersherum orientiert als in den restlichen Ausführungen. Mit der im weiteren Verlauf benutzten Konvention ist also $-U$ anstelle von U gegen die Zeit aufgetragen. (Abbildungen aus [45, Abb. 8], ©1952, John Wiley & Sons)

Spannungskonventionen und Sprechweisen

Wir werden uns ab jetzt der Konventionen von H & H anschließen und die Spannungen in Relation zur Ruhespannung U_r angeben, $V = U - U_r$. Negative Spannungen V bedeuten also, dass das in der Ruhespannung negativ geladene Axon *depolarisiert* wird. Positive Spannungen V korrespondieren mit einem noch stärker negativ geladenen Axon als im Ruhezustand, wir nennen das Axon dann *hyperpolarisiert*.

Die jetzt folgenden Versuche werden alle unter *Voltage-Clamp*, also mit dem Feedback-Mechanismus durchgeführt. Bevor wir in die Details der Versuchsreihen eintauchen, soll ein in diesem Zusammenhang wichtiger Begriff erläutert werden, nämlich der des *elektrochemischen Antriebs* oder *elektrochemischen Potentials*.

Das elektrochemische Potential

Zum einen wirken auf ein Ion, ein elektrisch geladenes Teilchen, elektrische Kräfte: Die positiv geladenen Natrium- und Kaliumionen werden von Orten mit hohem elektrischem Potential zu Orten mit niedrigem elektrischem Potential getrieben (so wie umgekehrt das negative Elektron in einem Stromkreislauf vom Ort des niedrigen elektrischen Potentials, z. B. dem Minuspol der Batterie, zum Ort mit hohem elektrischem Potential, z. B. dem Pluspol der Batterie, getrieben wird). Das ist der elektrische Teil des Potentialgefälles.

Zum anderen diffundieren die Teilchen aber auch von Orten, an denen sie in hoher Konzentration vorliegen, zu Orten, an denen sie in niedriger Konzentration vorliegen. Das ist der chemische Teil des Potentialgefälles. Im elektrochemischen Gleichgewicht heben sich diese Effekte gerade auf.

Wir wollen das an einem Beispiel erläutern (ohne ernsthaft in die statistische Physik einzusteigen). Dazu stellen wir uns zwei Kammern vor, die durch einen Durchgang verbunden sind. In Kammer 1 soll das elektrische Potential φ_1 vorliegen und in Kammer 2 das elektrische Potential φ_2 mit $\varphi_1 > \varphi_2$. Diese beiden Potentiale sollen durch irgendeinen Mechanismus von außen aufrechterhalten werden. In diese beiden verbundenen Kammern entlassen wir nun einen Schwarm von ungefähr 10^{21} gleichen, positiv geladenen Teilchen, von denen sich jedes einzelne in zufälliger Richtung durch die beiden Kammern bewegt (Stöße mit den Wänden und anderen Teilchen sollen elastisch, also ohne Energieverlust, sein). Ohne den elektrischen Potentialunterschied wäre die Konzentration der Teilchen in beiden Kammern nach kurzer Zeit gleich groß. Wären die Konzentrationen nämlich unterschiedlich, würden durch die zufälligen Bewegungen mehr Teilchen von der Kammer mit der höheren Konzentration in die Kammer mit der niedrigeren Konzentration fliegen als umgekehrt. Dadurch gleichen sich die Konzentrationen aus.

Durch den elektrischen Potentialunterschied wird nun aber die Bewegung in Richtung Kammer 2 gegenüber der Bewegung in Richtung Kammer 1 bevorzugt. Dadurch vergrößert sich die Konzentration in der Kammer mit dem niedrigeren elektrischen Potential. Das passiert so lange, bis sich die beiden Effekte wieder aufheben. Die statistische Physik liefert eine Formel für den Zusammenhang zwischen elektrischem Potentialunterschied und Konzentrationsunterschied im Gleichgewichtszustand, die Nernst-Gleichung:

$$\varphi_2 - \varphi_1 = \frac{RT}{F} \ln \left(\frac{v_1}{v_2} \right), \quad (8.1)$$

dabei sind v_1 und v_2 die Konzentrationen in den Kammern 1 und 2, T ist die (absolute) Temperatur und R und F sind gewisse Konstanten (die universelle Gaskonstante und die Faraday-Konstante). Ist die linke Seite größer als die rechte Seite, resultiert daraus ein Teilchenstrom in die Kammer 1; ist umgekehrt die rechte Seite größer als die linke, gibt es einen Teilchenstrom in die Kammer 2.

8.3.2 Die Trennung in kapazitive Ströme und Ionenströme

In allen Versuchen unter *Voltage-Clamp* legen H & H stufenförmige Spannungsprofile an der Membran an: Das heißt, zu einem bestimmten Zeitpunkt wird die Ruhespannung $V = 0$ zu einer anderen Spannung V_1 verändert und diese dann eine gewisse Zeitspanne (meistens im Millisekundenbereich) konstant gehalten. Dabei zeigen sich immer zwei Arten von Strömen: Die erste Art besteht aus einem sehr kurzfristigem Strom (Dauer ca. 30 μs), der stets ein ähnliches Profil aufweist, nämlich bei negativen angelegten Spannungen nach außen und bei positiven angelegten Spannungen nach innen weist und dessen Stärke proportional zur angelegten Spannung ist. Zwei typische Messkurven dieses Stroms sind in Abb. 8.6 dargestellt. Diesen Teil des Stroms interpretieren H & H folgendermaßen: Ein Teil der Membran verhält sich wie ein *elektrischer Kondensator*. Ein elektrischer Kondensator ist ein physikalisches Modell für einen Gegenstand, der in der Lage ist, elektrische Ladungen aufzunehmen, und für den ein proportionaler Zusammenhang $Q \sim U$ zwischen der an dem Kondensator anliegenden Spannung U und der aufgenommenen elektrischen Ladung Q besteht. Die Proportionalitätskonstante C (also $Q = CU$) gibt die sogenannte *Kapazität* des Kondensators an. Ein Kondensator hat folglich dann eine große Kapazität, wenn er auch schon bei kleinen angelegten Spannungen eine große Ladung aufnimmt. Aus diesen Kurven schließen H & H, dass der kapazitive Teil der Membran sich wie ein Kondensator mit einer Kapazität von ca. 35 mF cm^{-2} verhält. (Die übliche Einheit der elektrischen Ladung ist das Coulomb. Ein Coulomb entspricht der elektrischen Ladung von 10^{15} Elektronen, die gerade eine negative Elementarladung e tragen. Ein Kondensator mit einer Kapazität von einem Farad (1 F) trägt also bei einer Spannung von einem Volt eine Ladung von einem Coulomb.)

Bei dieser Art von Strömen treten im engeren Sinne gar keine Ladungsträger durch die Membran hindurch, sondern die Ladungsträger verteilen sich lediglich im Inneren und Äußeren der Zelle räumlich anders und verursachen so den gemessenen Strom zwischen den Elektroden c und d . Hier eine Erläuterung der Situation bei einer Depolarisation des Axons im Detail: Bei einer Depolarisation werden negativ geladene Elektronen aus der Elektrode a im Innern des Axons abgezogen. Dadurch werden die positiv geladenen Ionen weniger stark von dem Draht angezogen bzw. sogar abgestoßen, falls die Elektrode a positiv geladen ist, und konzentrieren sich dadurch stärker an der Membran, was zu einer positiven

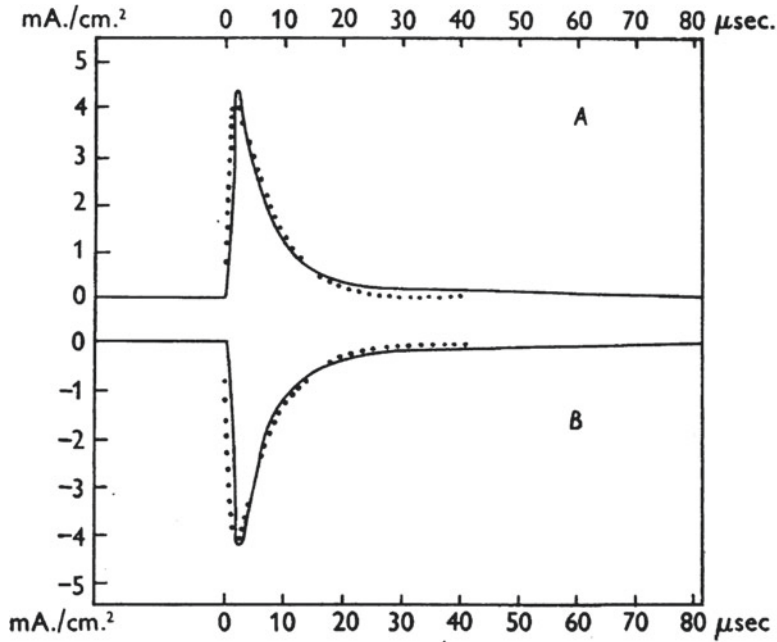


Abb. 8.6 Der Strom durch den kapazitiven Teil der Membran. Zur Zeit $t = 0$ wurde die Spannung von 40 mV für die Kurve A und -40 mV für die Kurve B angelegt. Die durchgezogenen Linien sind experimentell aufgenommene Stromverläufe, die gepunkteten Kurven wurden mit der Funktion $I^*(t) = \pm 6.8 (\exp(-0.159t) - \exp(-t))$ erstellt. (Abbildung aus [45, Fig. 16], ©1952, John Wiley & Sons)

Ladung auf der Innenseite der Membran führt. Diese positive Ladung auf der Innenseite der Membran verursacht eine Abstoßung der ebenfalls positiv geladenen Ionen auf der Außenseite der Membran, die sich dadurch stärker im von der Membran weit entfernten Bereich konzentrieren. Diese Verschiebung der Ionen verursacht insgesamt einen kurzen, nach außen gerichteten Strom, der somit als negativ gemessen wird und andauert, bis sich die Konzentrationen passend zur angelegten Spannung eingestellt haben.

Die zweite Art von Strömen hat eine langsamere Zeitskala und zeigt ein stark nichtproportionales Verhalten: Eine Auswahl von Stromkurven bei Hyper- und Depolarisationen des Axons von 65 mV bis -130 mV ist in Abb. 8.7 dargestellt. Die schnellen kapazitiven Ströme können in dieser Zeitaufösung nicht dargestellt werden und verbergen sich in den anfänglichen Unstetigkeiten der aufgenommenen Stromkurven. Für Hyperpolarisationen stellt sich ein einwärts gerichteter Strom ein, der zeitlich konstant erscheint und dessen Stärke proportional zur Auslenkung aus der Ruhespannung wirkt. Depolarisationen zwischen -5 und -10 mV bewirken einen auswärts gerichteten Strom, dessen Stärke mit der Zeit zuzunehmen scheint (das ist auf Aufnahmen mit längerer Laufzeit in Abb. 8.8 deutlich klarer zu

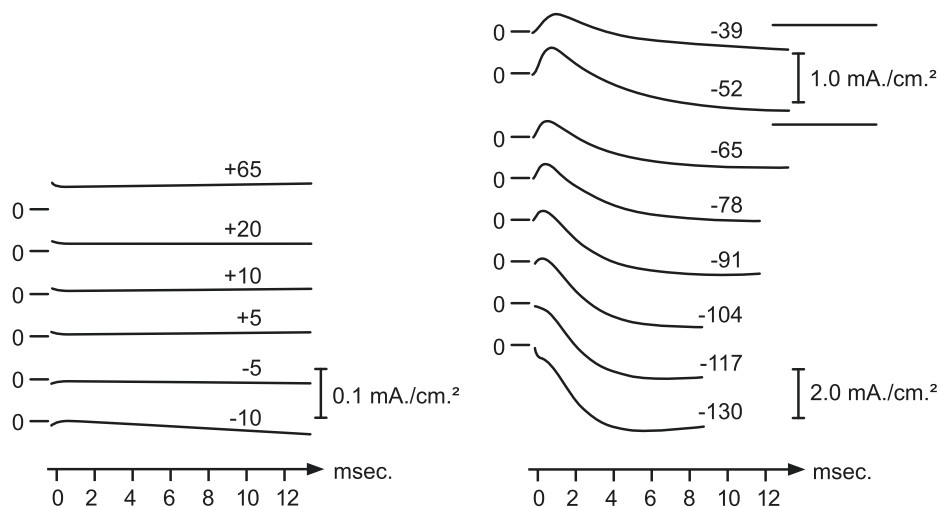
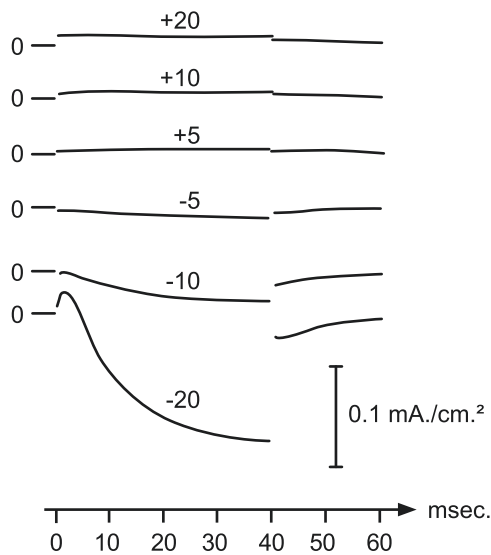


Abb. 8.7 Aufnahmen von Strömen unter *Voltage-Clamp*. Die ab der Zeit $t = 0$ angelegte Polarisation des Axons ist an der Kurve vermerkt. Man beachte die unterschiedlichen Skalen für die Stromstärke. (Abbildung nach [45, Abb. 12])

Abb. 8.8 Aufnahmen von Strömen unter *Voltage-Clamp*. Die ab der Zeit $t = 0$ angelegte Polarisation des Axons ist an der Kurve vermerkt. Zur Zeit $t = 40$ ms wird wieder auf die Ruhespannung $V = 0$ geschaltet. (Abbildung nach [45, Abb. 12])



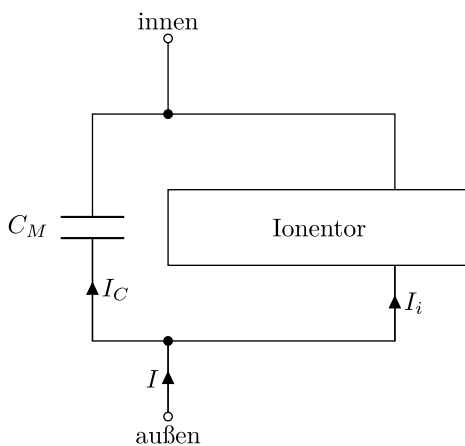
sehen). Depolarisationen zwischen -20 mV und -104 mV erzeugen zunächst einen einwärts gerichteten Strom, der sein Maximum in den ersten zwei Millisekunden erreicht, wobei das Maximum bei größeren Depolarisationen früher erreicht wird. An den kurzfristigen, einwärts gerichteten Strom schließt sich ein lang anhaltender, auswärts gerichteter Strom an, dessen Stärke mit wachsender Depolarisation zunimmt und einen Endzustand anzustreben scheint, siehe Abb. 8.8. Ab Depolarisationen von -117 mV beginnt die Kurve direkt mit einem auswärts gerichteten Strom, der bei -117 mV nahezu geradlinig startet und bei noch höheren Depolarisationen um -130 mV einen anfänglichen Buckel aufweist.

H & H stellen die These auf, dass diese Ströme durch Ionen unterschiedlicher Art hervorgerufen werden. Bei diesen Strömen passieren die Ladungsträger die Membran tatsächlich. Sie modellieren die Membran physikalisch als eine Parallelschaltung von einem Kondensator und einem weiter zu untersuchenden Element, welches die Ionenströme passieren lässt, einem Ionentor. Ein schematisches Schaubild ist in Abb. 8.9 dargestellt. Der gemessene Gesamtstrom I setzt sich aus dem kapazitiven Strom I_C und dem Ionenstrom I_i zusammen. Aus der Beziehung $Q = CU$ erhalten wir durch Differentiation nach der Zeit zudem die Beziehung $I_C = \dot{Q} = C\dot{U} = C\dot{V}$ und somit

$$I = C\dot{V} + I_i.$$

Diese Gleichung unterstreicht den Nutzen der *Voltage-Clamp*-Technik: Sie erlaubt es, die Ionenströme durch die Membran direkt zu messen, weil die kapazitiven Ströme während eines *Voltage-Clamp* verschwinden.

Abb. 8.9 Die Membran als Parallelschaltung eines Kondensators und eines noch zu untersuchenden Bauteils, durch das die Ionenströme fließen



8.3.3 Die Trennung in Natrium- und Kaliumionenströme

Die aufgezeichneten Stromkurven in Abb. 8.7 und 8.8 zeigen insbesondere für Depolarisationen zwischen -20 mV und -104 mV die Auffälligkeit, dass die Membran auf eine Depolarisation mit einem kurzfristigen Strom reagiert, der anscheinend in die „falsche“ Richtung zeigt: Obwohl das Innere des Axons durch die Depolarisation „positiver“ wird, strömen zunächst positive Ladungsträger in das Innere. H & H stellen die folgende Hypothese auf: Da im Inneren des Axons die Natriumionenkonzentration nur ein Fünftel der Natriumionenkonzentration im umgebenden Seewasser beträgt, liegt im Ruhezustand des Axons und auch bei Depolarisationen bis -104 mV ein elektrochemischer Antrieb für einen Transport von Natriumionen in das Axon vor. Für eine Erläuterung des elektrochemischen Potentials oder Antriebs siehe den Kasten im Abschn. 8.3.1. Wenn die Zellmembran durch irgendeinen zu erforschenden Mechanismus durch die angelegte Spannung kurzfristig durchlässig für Natriumionen würde, würde das zu einem Natriumstrom in die Zelle führen. Bei großen Depolarisationen von -120 mV und mehr würde diese Durchlässigkeit für Natriumionen dann für einen kurzfristigen, auswärts gerichteten Strom sorgen, der zu dem anfänglichen Buckel in der Stromkurve bei sehr hohen Depolarisationen in Abb. 8.7 führen würde. Die These lautet also: Der kurzfristige Strom in den ersten Millisekunden ist auf den Strom von Natriumionen zurückzuführen.

H & H prüfen diese Hypothese wie folgt: Sie führen Messungen durch, in denen in der Axonumgebung das Meerwasser durch Flüssigkeiten ersetzt wird, die entweder gar keine Natriumionen oder aber Natriumionen in einer sehr viel geringeren Konzentration enthalten. Sie wählen Konzentration von 0 %, 10 % und 30 % von der in Meerwasser vorliegenden Natriumkonzentration. Die drei in Abb. 8.10 aufgenommenen Stromkurven entstanden alle bei einer Depolarisation von -65 mV. Die erste Messung fand in der üblichen Meerwasserumgebung statt, für die zweite Messung wurde das Meerwasser durch Cholin ersetzt, einer Elektrolytlösung, die kein Natrium enthält, das ist die 0 %-Lösung. Für die dritte Messung wurde wieder Meerwasser benutzt.

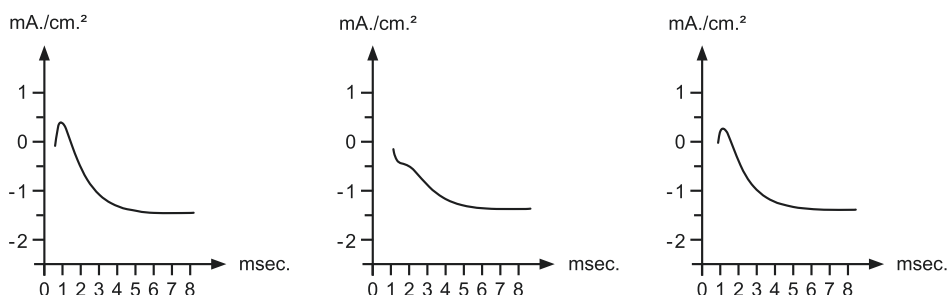


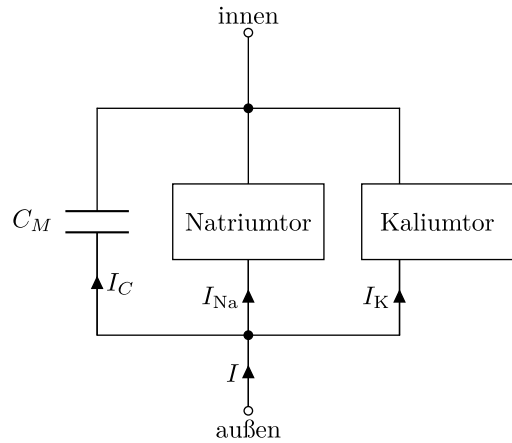
Abb. 8.10 Aufzeichnungen von Stromkurven unter *Voltage-Clamp* bei einer Depolarisation von -65 mV in einer Umgebung von Meerwasser, Cholinlösung und Meerwasser (Abbildungen nach [41, Fig. 1])

Die erste Kurve zeigt den üblichen Stromverlauf mit dem kurzfristigen, einwärts gerichteten Strom. In der zweiten Aufzeichnung fällt der einwärts gerichtete Strom weg, stattdessen setzt der auswärts gerichtete Strom mit einer Stufe ein, der sich zunächst ein Plateau anschließt, welches dann in einen Auswärtsstrom übergeht, wie er auch in der ersten Aufzeichnung zu erkennen ist. Die dritte Aufnahme gleicht dann im Wesentlichen wieder der ersten Aufnahme. Diese experimentellen Ergebnisse stützen die Annahme, dass die kurzfristigen Ströme durch Natriumionen getragen werden. Der einwärts gerichtete Strom fällt in der zweiten Aufnahme weg, weil in der Umgebung gar keine Natriumionen vorliegen, die einwärts strömen könnten.

Der kurze, fast senkrechte Abfall in der zweiten Aufnahme wird dadurch erklärt, dass kurzfristig Natriumionen aufgrund des Konzentrationsgefälles (außen gibt es so gut wie keine Natriumionen) ausströmen. Dafür spricht, dass dieser senkrechte Teil der Stromkurve nicht auftritt, wenn für die mittlere Messung nicht Cholin, sondern Meerwasser mit einem reduzierten Natriumanteil von 10 % des normalen Natriumanteils verwendet wird. Zugehörige Messkurven finden sich in [41, Abb. 3. und Abb. 4], sind hier aber nicht abgebildet.

Der langfristige Stromverlauf ändert sich im Wesentlichen nicht, egal ob die Umgebung aus Meerwasser oder aus Cholin besteht. H & H folgern daraus, dass Natriumionen zu diesem Strom nicht beitragen. Ferner ist auf der mittleren Aufnahme in Abb. 8.10 und auf vielen anderen, in dieser Art entstandenen Aufnahmen zu erkennen, dass der Strom ohne Natriumionen „verspätet“ einsetzt, nämlich erst nach ca. 2 ms. H & H stellen die Hypothese auf, dass dieser Strom hauptsächlich durch ausströmende Kaliumionen und zu einem geringen Anteil durch weiterer Ionenströme, darunter vor allem einströmende Chloridionen, zustande kommt. Der Ionenstrom besteht also aus zwei Teilen, dem Natrium- und dem Kaliumstrom (alle anderen Ströme werden zunächst einmal dem Kaliumstrom zugerechnet). Die Mechanismen innerhalb der Membran, die für die Durchlässigkeit für einzelne Ionensorten sorgen, scheinen unabhängig voneinander zu arbeiten: Die Durchlässigkeit für Natriumionen ist kurzfristig und erfolgt schnell auf die veränderte Spannung, die Durchlässigkeit für die übrigen Ionen reagiert langsamer auf die angelegte Spannung, scheint aber über den Messzeitraum erhalten zu bleiben. Das schematische Schaubild wird also verfeinert: Der ionendurchlässige Bereich wird in eine Parallelschaltung von zwei unabhängigen Ionentoren, einem für Natrium und einem für die übrigen Ionen, analytisch ausdifferenziert, siehe Abb. 8.11. H & H stellen nun drei weitere Hypothesen auf, mit deren Hilfe sie aus dem gemessenen Gesamtionenstrom I_i , der in der Umgebung Meerwasser fließt, durch Vergleich mit dem Gesamtionenstrom I'_i , der in der Umgebung mit reduziertem Natriumanteil fließt, die Natriumströme I_{Na} und I'_{Na} sowie die Kaliumströme I_K und I'_K in der Meerwasserumgebung bzw. in der Umgebung mit verändertem Natriumgehalt bestimmen können. (Gestrichene Stromgrößen beziehen sich hier immer auf die modifizierte Umgebung. Die Kaliumströme beinhalten immer auch noch einen Anteil an anderen Ionen.) Die drei Hypothesen lauten wie folgt:

Abb. 8.11 Die Membran als Parallelschaltung eines Kondensators mit zwei noch zu untersuchenden Bauteilen, durch die der Natrium- bzw. der Kalium- zuzüglich Reststrom fließt



- 1) Der zeitliche Verlauf des Kaliumstroms ist in beiden Umgebungen identisch: $I_K = I'_K$
- 2) Die Form des Natriumstroms ist in beiden Umgebungen gleich, d. h., es gibt eine (positive oder auch negative) Konstante k , so dass $I'_{Na} = k I_{Na}$ gilt.
- 3) Im ersten Drittel des Zeitraums, den der Natriumstrom benötigt, um seine maximale Stromstärke zu erreichen, gilt $\dot{I}_K = 0$.

Diese Annahmen sind mit den aufgenommenen Stromkurven gut vereinbar. Mit ihnen lassen sich die Natrium- und Kaliumströme aus den beobachteten Ionenströmen I_i und I'_i bestimmen. Als Erstes wird die Konstante k bestimmt, die in Annahme 2) postuliert wird. Dazu wird für sehr kleine t I'_i gegen I_i geplottet. Werden nur Wertepaare $(I'_i(t), I_i(t))$ für hinreichend kleine t berücksichtigt, entsteht eine Gerade mit der Steigung k , denn mit den genial einfachen Bezeichnungen der Physiker gilt mit den Annahmen 1) bis 3)

$$\frac{dI'_i}{dI_i} = \frac{dI'_i}{dt} \cdot \frac{dt}{dI_i} = \frac{\dot{I}'_i}{\dot{I}_i} = \frac{\dot{I}'_{Na} + \dot{I}'_K}{\dot{I}_{Na} + \dot{I}_K} = \frac{k \dot{I}_{Na}}{\dot{I}_{Na}} = k.$$

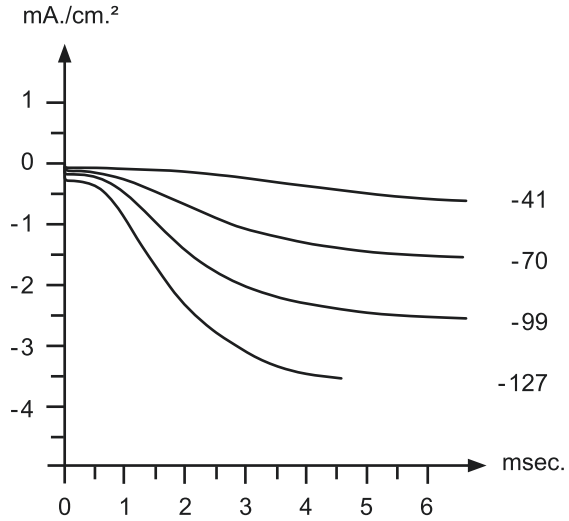
Dass bei solchen I'_i - I_i -Plots tatsächlich eine Gerade entsteht, wenn nur Wertepaare für sehr kleine Zeiten berücksichtigt werden, verleiht den Annahmen 1) bis 3) weitere Plausibilität. Mit dem so gefundenen k können wir I_{Na} berechnen: $I_i - I'_i = I_{Na} - I'_{Na} = I_{Na}(1 - k)$, also erhalten wir

$$I_{Na} = \frac{I_i - I'_i}{1 - k} \quad (8.2)$$

$$I_K = I_i - I_{Na} = \frac{I'_i - k I_i}{1 - k}.$$

Kurven des so berechneten Kaliumstroms sind in Abb. 8.12 für Depolarisationen von -41 mV bis -127 mV dargestellt.

Abb. 8.12 Der Kaliumstrom für verschiedene Werte von Depolarisationen, aufgenommen und berechnet aus dem Vergleich der Ionenströme in Meerwasser und in 30 %-Natrium-Meerwasser (Abbildung nach [41, Fig. 6])



8.3.4 Der „passive“ Ionentransport und die selektive Leitfähigkeit

Die nächste zu klärende Frage ist die nach der Natur des Natrium- und Kaliumtransports durch die Membran. Handelt es sich um einen rein „passiven“ Transport, bei dem die Ionen ausschließlich aufgrund des herrschenden elektrochemischen Antriebs durch die Membran hindurchtreten oder benötigen sie irgendeine Art von Trägerteilchen, die sie „aktiv“ und eventuell sogar gegen das elektrochemische Potential durch die Membran tragen? In einer früheren Veröffentlichung favorisierten H & H die zweite Möglichkeit, siehe [44]. Durch die beschriebenen *Voltage-Clamp*-Versuche sammeln sich aber nun viele Hinweise, die für die erste Variante sprechen. Wenn die These vom passiven Transport richtig ist, werden die Ionen lediglich durch den elektrochemischen Antrieb $V - V_{\text{Na}}$ bzw. $V - V_{\text{K}}$ angetrieben. Dabei sind U_{Na} und U_{K} gerade die elektrochemischen Gleichgewichtsspannungen für Natrium- und Kaliumionen bei dem vorherrschenden Verhältnis der Ionenkonzentration an der Innen- und Außenseite der Membran und wir setzen $V_{\text{Na}} = U_{\text{Na}} - U_r$ und $V_{\text{K}} = U_{\text{K}} - U_r$, siehe auch die Erläuterungen in dem Kasten zur Spannungsconvention in Abschn. 8.3.1. Die Membran weist dann zu jedem Zeitpunkt für jede Ionensorte eine bestimmte, aber sich eventuell zeitlich ändernde Durchlässigkeit auf. Weil es sich bei den Teilchen, die durchgelassen werden, um geladene Teilchen handelt, spricht man hier in der Regel anstatt von Durchlässigkeit von Leitfähigkeit, und so wollen wir es ab jetzt auch halten. Die Leitfähigkeit für Natrium und Kalium ließe sich dann für einen Zeitpunkt t durch einen Koeffizienten $g_{\text{Na}}(t)$ bzw. $g_{\text{K}}(t)$ beschreiben, die Koeffizienten sind dabei durch die Gleichungen $I_{\text{Na}}(t) = g_{\text{Na}}(t) \cdot (V(t) - V_{\text{Na}})$ und $I_{\text{K}}(t) = g_{\text{K}}(t) \cdot (V - V_{\text{K}})$ bestimmt und haben die Dimension Strom pro Fläche pro Spannung. Mit den üblichen Einheiten Ampere für die Stromstärke und Volt für die Spannung hat die Leitfähigkeit (pro Fläche) die Einheit $\text{A V}^{-1} \text{cm}^{-2}$.

Nun gut, g_{Na} und g_{K} lassen sich für jeden $I(t)$ - $V(t)$ -Zusammenhang auf diese Art und Weise definieren, welche Eigenschaften müssen diese beiden Größen nun haben, wenn die Annahme vom „passiven Ionenstrom“ zutrifft?

- 1) g_{Na} und g_{K} müssen stets nichtnegativ sein.
- 2) Wenn die Membran in einem definierten Leitfähigkeitszustand präpariert ist, sagen wir $\hat{g}_{\text{Na}}$ und $\hat{g}_{\text{K}}$, dann muss für alle möglichen Spannungen V und die sich *unmittelbar nach dem Umschalten* einstellenden Ionenströme die Strom-Antrieb-Relation proportional sein, also $I_{\text{Na}} \sim (V - V_{\text{Na}})$ mit der Proportionalitätskonstante $\hat{g}_{\text{Na}}$ und $I_{\text{K}} \sim V - V_{\text{K}}$ mit der Proportionalitätskonstante $\hat{g}_{\text{K}}$.

Die erste Bedingung lässt sich sofort überprüfen. Für die zweite Bedingung muss man etwas tricksen, denn die Leitfähigkeiten der Membran scheinen sich ja laufend mit der Zeit zu ändern. Andernfalls wären ja I_{Na} und I_{K} in einer *Voltage-Clamp*-Situation konstant. Wie schafft man es also, die Membran immer wieder in ein und denselben Zustand zu manövrieren und dann unterschiedliche Spannungen anzulegen? Wir diskutieren hier die Versuche, die H & H bezüglich des Kaliumstroms durchführen. Dazu experimentieren sie in der Cholinumgebung. In einer Reihe von Versuchen präparieren sie die Membran durch eine immer gleiche Depolarisation von $V_0 = -84$ mV über einen Zeitraum von 0.63 ms. Nach 0.63 ms legen sie dann in jedem Versuch der Versuchsreihe eine unterschiedliche Spannung V an, sie wählen $V = 28, 21, 14, 7, 0, -28$ und -56 mV. Der zugehörige Ionenstrom wird aufgezeichnet. Im Moment des Umschaltens sollte sich die Membran über die Versuchsreihe hinweg immer im selben Zustand befinden. Ist die Passivitätshypothese richtig, sollte die Kaliumstromstärke unmittelbar nach dem Umschalten der Spannung linear von der neu gewählten Spannung abhängen. Die zugehörigen, um die Kapazitätsströme korrigierten Ionenströme sind in Abb. 8.13 zu sehen. Da es sich um Versuche in der Cholinumgebung handelt, sollten die Beiträge des Natriumstroms im Vergleich zum Kaliumstrom sehr klein sein und nicht ins Gewicht fallen. Durch ein sorgfältiges Auslesen der hier dargestellten und weiterer Stromkurven in [40, Fig. 11] kommen wir auf die Messwerte in Tab. 8.1. Wir tragen diese Messwerte in ein V - I -Diagramm ein und sehen, dass sie hinreichend nah auf einer Geraden liegen, siehe Abb. 8.14. Die eingezeichnete Ausgleichsgerade hat eine Steigung von $0.018 \text{ A V}^{-1} \text{ cm}^{-2}$ und schneidet die V -Achse bei 13.9 mV. Damit hat sich die Passivitätshypothese bewährt und wir können auch die Kalium-Gleichgewichtsspannung ablesen, nämlich $V_{\text{K}} = 14$ mV.

Vergleichbare Experimente führen H & H auch bezüglich des Natriumstroms durch. Diese Experimente bestätigen die Passivitätshypothese für den Natriumstrom und liefern einen Wert von $V_{\text{Na}} = -112$ mV für das Natrium-Gleichgewicht. Dieser Wert steht in guter Übereinstimmung mit den experimentellen Kurven in Abb. 8.7: Dort kehrt sich die Richtung des kurzzeitigen Stroms, den wir als Natriumstrom identifiziert haben, zwischen den Depolarisationen -104 mV und -117 mV von einem einwärts gerichteten Strom in einen auswärts gerichteten Strom um. Auch eine Bestimmung über die unterschiedlichen

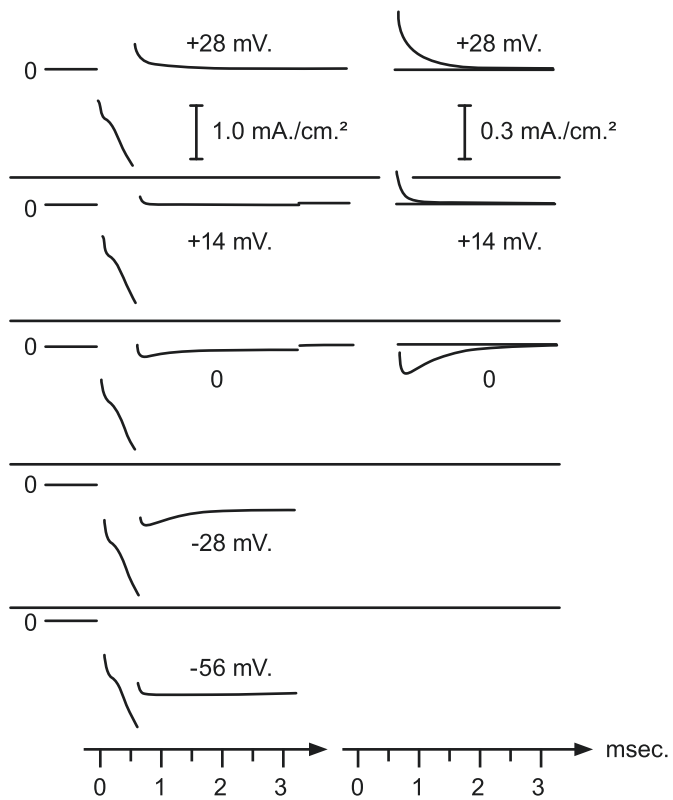


Abb. 8.13 Experimente zur Überprüfung der Passivitätshypothese für Kaliumströme. Strom durch die Membran nach einer Depolarisation um -84 mV , gefolgt von einer Auslenkung, die an den jeweiligen Kurven angegeben ist. Die Kurven links zeigen die Ströme während beider Schritte, in den Kurven rechts sind die Ströme nach der zweiten Auslenkung vergrößert dargestellt. Die Experimente fanden in einer Cholinumgebung statt. (Abbildung nach [40, Abb. 11])

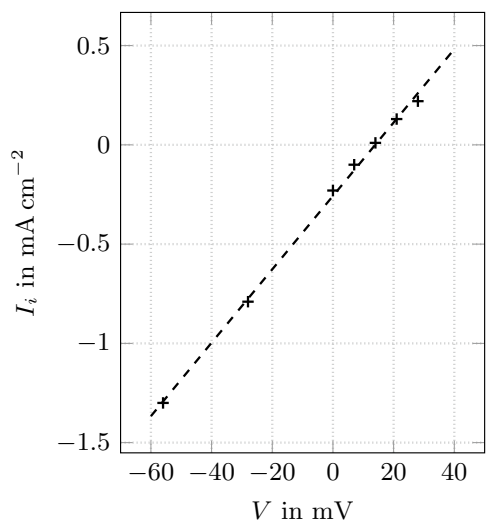
Tab. 8.1 Aus Abb. 8.13 bzw. [40, Abb. 11] ausgelesene Messwerte für den Ionenstrom unmittelbar nach Umschalten des angelegten *Voltage-Clamp* auf die angegebene Spannung V

$\frac{V}{\text{mV}}$	28	21	14	7	0	-28	-56
$\frac{I_i}{\text{mA}}$	0,22	0,13	0,01	-0,10	-0,23	-0,79	-1,35

Natriumkonzentrationen im Innern des Axons und in der Umgebung mit Hilfe der Nernst-Gleichung (8.1) führt auf einen vergleichbaren Wert, siehe Aufg. 8.1.

Die wesentlichen Größen, die von jetzt an studiert werden, sind die Natriumleitfähigkeit g_{Na} und die Kaliumleitfähigkeit g_{K} . Experimentelle Kurven für diese Größen erhält man nach dem eben Gesagten aus den Stromkurven des Natriumstroms und des Kaliumstroms mittels einfacher Division durch $V - V_{\text{Na}}$ bzw. $V - V_{\text{K}}$. Kurven für die zeitliche Entwicklung

Abb. 8.14 Strom-Spannung-Relation zwischen dem unmittelbaren Ionenstrom und der angelegten Spannung bei einer vorherigen Depolarisation auf -84 mV. Gemessene Werte aus Tab. 8.1, Ausgleichsgerade: $I = 0.01848V - 0.2573$ mit der Nullstelle $V_0 = 13.9$, V in mV, I in mA cm^{-2}



der Leitfähigkeiten sind in Abb. 8.15 wiedergegeben. Dabei wurde das Axon unter *Voltage-Clamp* aus der Ruhelage auf die am Graphen angegebene relative Spannung V gebracht. In den nächsten beiden Abschnitten werden wir die Rückschlüsse darstellen, die H & H aus diesen und weiteren Zeit-Leitfähigkeit-Kurven ziehen, getrennt für die Kalium- und die Natriumleitfähigkeit.

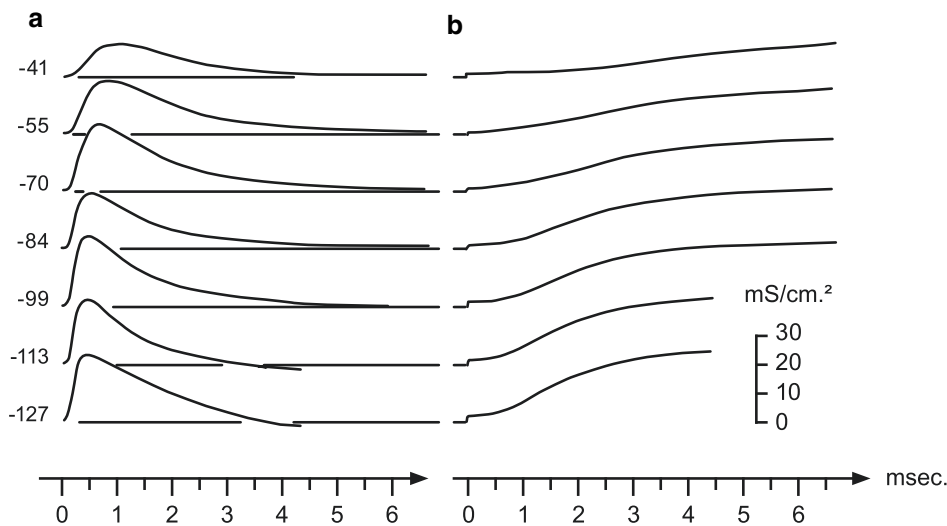


Abb. 8.15 Kurven der Natriumleitfähigkeit (**a**) und der Kaliumleitfähigkeit (**b**) der Membran. Die dabei angelegte Auslenkung V aus der Ruhespannung ist an jeder Kurve in mV angegeben. (Abbildung nach [41, Abb. 8])

8.3.5 Die Kaliumleitfähigkeit

Aus den Kurven in Abb. 8.15b entnehmen wir, dass die Kaliumleitfähigkeit bei angelegten Depolarisationen S-förmig ansteigt und sich einem Grenzwert anzunähern scheint. Diese langfristig angenommenen Leitfähigkeitsgrenzwerte aus den in Abb. 8.15b dargestellten Experimenten und einer Reihe entsprechender Experimente sind in Abb. 8.16 notiert. Die Kurven aus Abb. 8.15b liefern die mit einem Kreis markierten Messwerte in Abb. 8.16 (Axon 20).

Für kleine Auslenkungen V aus der Ruhelage scheint die Leitfähigkeit nahezu exponentiell anzusteigen (die logarithmisch aufgetragenen Werte liegen annähernd auf einer Geraden). Im Ruhezustand der Zelle beträgt die Kaliumleitfähigkeit ca. 2 % der bei -100 mV erreichten Leitfähigkeit. Bei einer Depolarisation von -10 mV hat sie sich bereits verfünffacht und erreicht 50 % bei einer Auslenkung von ca. -30 mV.

Wichtig ist nun, wie schnell die Membran bei einer bestimmten Auslenkung ihre maximale Kaliumleitfähigkeit erreicht. Aus Abb. 8.15b können wir bereits qualitativ ablesen, dass bei hohen Depolarisationen die Hälfte der maximalen Leitfähigkeit schneller erreicht ist als bei niedrigen Depolarisationen. Quantitativ ist dieser Zusammenhang in Abb. 8.17 dargestellt. In dieser Abbildung ist die maximale Steigung der Leitfähigkeitsgraphen, also

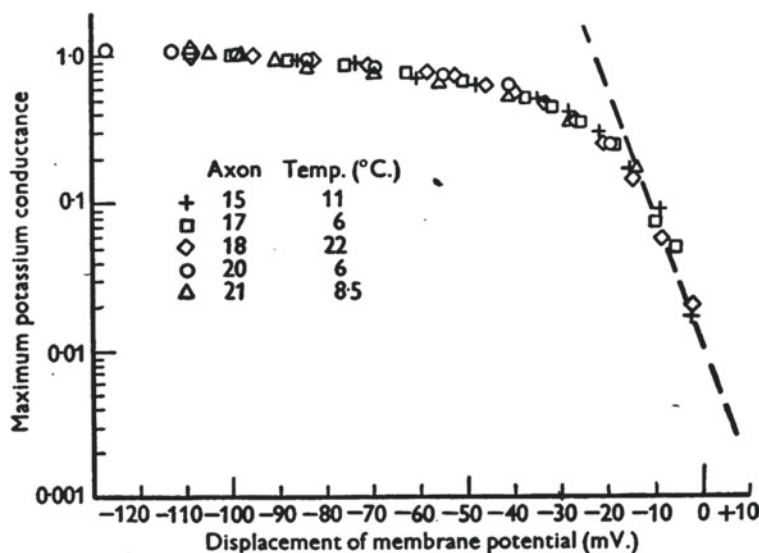


Abb. 8.16 Langfristige Kaliumleitfähigkeit, die während eines *Voltage-Clamp* erreicht wurde. Vertikale Achse: Verhältnis der langfristigen Leitfähigkeit bei der angelegten Spannung zur langfristigen Leitfähigkeit, die bei einer Depolarisation von -100 mV erreicht wird, logarithmische Auftragung. Horizontale Achse: Auslenkung aus der Ruhespannung während des *Voltage-Clamp* in mV (Abbildung aus [41, Abb. 10] ©(1952) John Wiley & Sons)

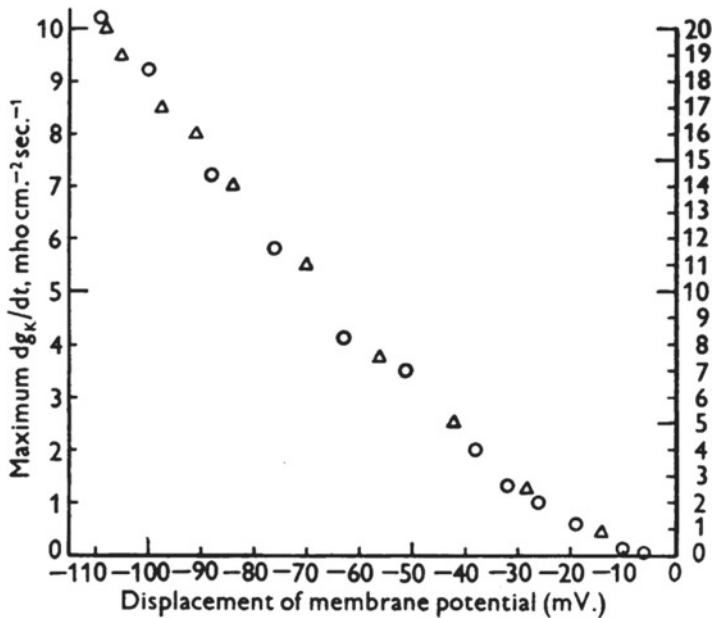


Abb. 8.17 Links: Maximale Steigung der Kaliumleitfähigkeit in Abhängigkeit von der angelegten Spannung. Linke vertikale Skala für die Versuche, die bei einer Temperatur von 6 °C, rechte vertikale Skala für die Versuche, die bei einer Temperatur von 8.5 °C durchgeführt wurden. (Abbildung aus [41, Abb. 12], ©(1952) John Willey & Sons)

die Steigung in den Wendepunkten, gegen die angelegte Auslenkung V geplottet. Die beiden unterschiedlichen Skalen für die maximale Rate sind nötig, weil die Messwerte von zwei Experimenten stammen, die bei unterschiedlichen Temperaturen stattgefunden haben (bei höheren Temperaturen laufen viele Prozesse prinzipiell schneller ab). Es wird jedoch deutlich, dass beiden Messreihen derselbe funktionale Zusammenhang zugrunde zu liegen scheint. Die maximale Leitfähigkeit wird also bei kleinen Auslenkungen langsam und bei großen Auslenkungen schnell erreicht. Den Messwerten zufolge erscheint es plausibel, dass für noch größere Depolarisationen die Geschwindigkeit der Anpassung der Leitfähigkeit weiter zunimmt.

Bis jetzt sind lediglich Auslenkungen aus der Ruhespannung betrachtet worden. Wie sieht es aber aus, wenn die angelegte Spannung von einer Depolarisation wieder in die Ruhespannung zurückkehrt? Kehrt dann auch die Leitfähigkeit in ihren alten Zustand zurück? H & H gehen dieser Frage in einer Reihe von Experimenten nach, in denen die Membran zunächst für eine gewisse Zeit depolarisiert und anschließend wieder auf das Ruhepotential zurückgeführt wird. Eine zugehörige Leitfähigkeitskurve ist in Abb. 8.18 dargestellt. In diesem Experiment kehrt die Leitfähigkeit wieder auf ihren ursprünglichen Ruhewert zurück, allerdings ist die Form der Rückkehr nicht mehr S-förmig, sondern scheint einen exponentiellen Verlauf aufzuweisen.

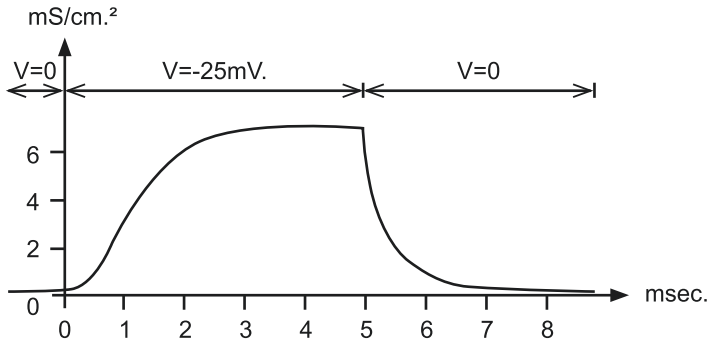


Abb. 8.18 Kaliumleitfähigkeit $g_K = I_K \cdot (V - V_K)$ mit $V_K = 12 \text{ mV}$. Für die Berechnung des Kaliumstroms wurde zudem ein geringfügiger Leckstrom abgezogen: $I_K = I_i - I_L$ mit $I_L = 0.4 \text{ mS cm}^{-2} \cdot (V + 4 \text{ mV})$. (Abbildungen nach [40, Fig. 13])

Wir fassen die Ergebnisse und die damit zusammenhängenden Vermutungen zusammen:

- Im Ruhezustand der Zelle hat die Membran eine geringe Kaliumleitfähigkeit in der Größenordnung von 1 % der Leitfähigkeit, die sie bei $V = -100 \text{ mV}$ erreicht.
- In Abhängigkeit von der Depolarisation steigt die (langfristige) Leitfähigkeit zunächst annähernd exponentiell an und erreicht bei ca. -30 mV die Hälfte der Leitfähigkeit, die bei -100 mV vorliegt.
- Die Leitfähigkeit ändert sich nicht instantan mit der Spannungsänderung, sondern stellt sich S-förmig auf den neuen Zustand ein. Die Geschwindigkeit dieses Einstellens ist abhängig von der Depolarisation und nimmt für größere Depolarisationen zu.
- Bei einer Repolarisation entwickelt sich die Leitfähigkeit wieder zurück, allerdings hier nicht S-förmig, sondern anscheinend exponentiell.

8.3.6 Die Natriumleitfähigkeit

Wie man aus den Kurven in Abb. 8.15a erkennen kann, ist die Situation bei den Natriumströmen komplexer als bei den Kaliumströmen: Bei einer Depolarisation steigt die Leitfähigkeit schnell (auch hier S-förmig) an, erreicht dann aber ein Maximum und geht anschließend deutlich langsamer wieder auf null zurück.

H & H nehmen nun an, dass es zwei voneinander unabhängige Mechanismen gibt, die gegeneinander spielen und gemeinsam die Leitfähigkeit für Natrium festlegen: Beide müssen auf hinreichend leitfähig geschaltet sein, damit die Membran für Natriumionen durchlässig wird. Sie vermuten, dass der eine, wir nennen ihn den *m*-Mechanismus, im Ruhezustand auf undurchlässig geschaltet ist und sich durch die Depolarisation schnell auf durchlässig umstellt. Der zweite Mechanismus, wir nennen ihn den *h*-Mechanismus, ist im Ruhezustand

auf durchlässig geschaltet und ändert sich bei Depolarisation langsam auf undurchlässig. Der schnelle Anstieg wäre somit durch die Öffnung des *m*-Mechanismus hervorgerufen, das langsame Abklingen hingegen durch die Schließung des *h*-Mechanismus.

H & H versuchen nun diese Mechanismen quantitativ zu beschreiben. Sie starten mit dem *m*-Mechanismus: Dazu werden die Werte der erreichten Leitfähigkeitsmaxima gegen die angelegte Spannung geplottet. Die Daten von 39 Messungen bei Depolarisationen zwischen 0 mV und -130 mV sind in Abb. 8.19 dargestellt. Die mit einem Kreis markierten Messwerte stammen von den Kurven in Abb. 8.15. Die erreichten Maxima werden als Maß für die langfristige erreichte Leitfähigkeit des *m*-Mechanismus bei der jeweils angelegten Auslenkung aus der Ruhespannung interpretiert. Auch hier scheint die langfristige *m*-Leitfähigkeit zunächst exponentiell zu wachsen (die Messwerte liegen in logarithmischer Auftragung nahezu auf einer Geraden). Die größtmögliche Leitfähigkeit scheint ab einer Auslenkung von um die -100 mV erreicht zu werden.

Ferner plotten sie die maximale Steigung, mit der die langfristige Leitfähigkeit erreicht wird, gegen die angelegte Spannung. Die Messwerte von zehn Messungen sind in Abb. 8.20 dargestellt. Sie nutzen diese Werte, um die Schnelligkeit der Einstellung des *m*-Mechanismus auf die neue Spannung einzuschätzen. Es ist aus den Werten auch quantitativ erkennbar, dass der *m*-Mechanismus sehr viel schneller reagiert als die Kaliumleitfähigkeit (vergleiche mit Abb. 8.17).

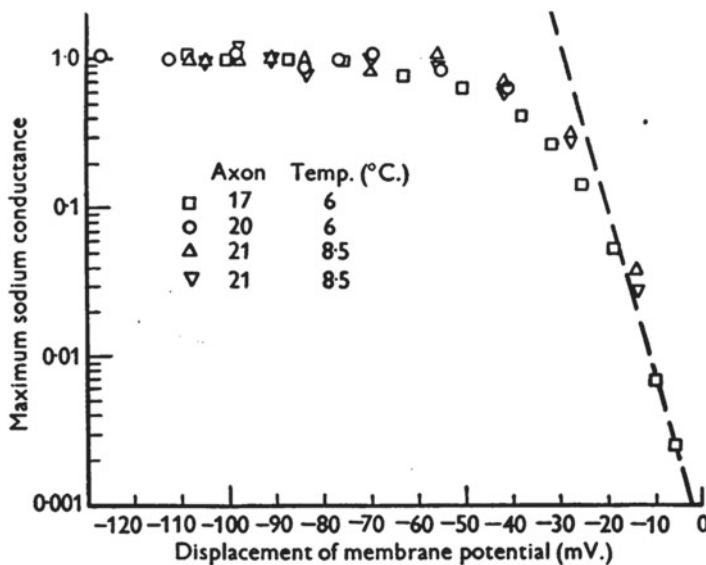


Abb. 8.19 Maximal erreichte Natriumleitfähigkeit unter *Voltage-Clamp*. Horizontale Achse: angelegte Auslenkung aus der Ruhespannung. Vertikale Achse: maximale Leitfähigkeit relativ zur maximalen Leitfähigkeit bei einer Auslenkung von -100 mV, logarithmische Auftragung (Abbildung aus [41, Abb. 9], ©(1952) John Wiley & Sons)

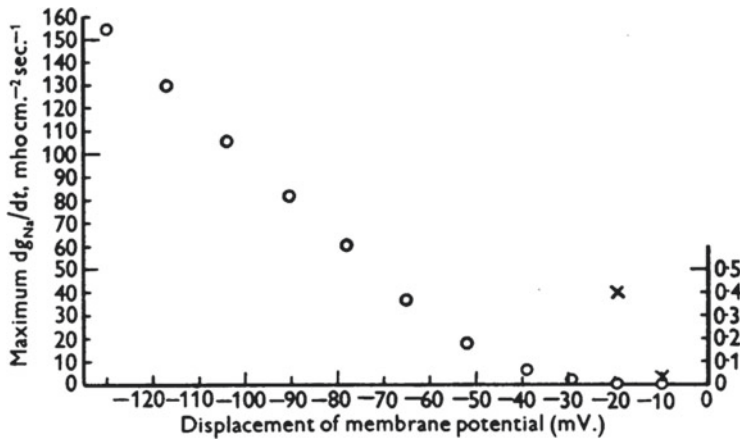


Abb. 8.20 Maximale Steigung der Natriumleitfähigkeit unter einem *Voltage-clamp*. Horizontale Achse: angelegte Auslenkung aus der Ruhespannung. Vertikale Achse: Maximale Änderungsrate der Natriumleitfähigkeit. Die durch Kreise dargestellten Messwerte beziehen sich auf die linke vertikale Skala. Die beiden Messwerte zu den Auslenkungen -10 mV und -20 mV sind zusätzlich durch Kreuze dargestellt, die sich auf die um den Faktor 100 vergrößerte Skala auf der rechten Seite beziehen. (Abbildung aus [41, Abb. 11], ©(1952) John Wiley & Sons)

Um den vermuteten langsamen *h*-Mechanismus zu untersuchen, lassen sich H & H etwas einfallen: Sie vermuten, dass auch schon kleine Depolarisationen langfristig den *h*-Mechanismus auf undurchlässiger stellen. Wenn dem so wäre, müsste bei einer anschließenden größeren Depolarisation nur ein geringeres Leitfähigkeitsmaximum erreicht werden, weil zwar der *m*-Mechanismus öffnet, der *h*-Mechanismus aber schlechter durchlässt. Sie untersuchen das mittels einer Reihe von Experimenten, in der einer „großen“ Depolarisation von -44 mV eine sogenannte Konditionierungsspannung von -8 mV unterschiedlicher Dauer vorgeschaltet ist. Drei Kurven dieser Versuchsreihe sind in Abb. 8.21a dargestellt. Es wird deutlich, dass durch die Vorkonditionierung das erreichte Maximum kleiner wird. In Abb. 8.21b sind Messkurven mit einer Vorkonditionierung von $+31$ mV dargestellt. Durch diese Vorkonditionierung wird das erreichte Maximum größer.

Analoge Versuchsreihen unternehmen sie auch mit anderen Konditionierungsspannungen und -dauern. Die Ergebnisse von vier Versuchsreihen werden in Abb. 8.22 zusammengefasst. Drei Messpunkte auf dem -8 mV-Ast stammen von den Kurven aus Abb. 8.21a und drei auf dem 31 mV-Ast von den Kurven aus Abb. 8.21b. Diese Daten lassen sich folgendermaßen interpretieren: Bei einer Depolarisation verringert der *h*-Mechanismus seine Durchlässigkeit relativ zur Durchlässigkeit bei $V = 0$ mV, bei einer Hyperpolarisation vergrößert er seine Durchlässigkeit relativ zur Durchlässigkeit bei $V = 0$ mV. Die Anpassung an die neue Spannung ist in der Umgebung von $V = 0$ mV langsam mit Zeitkonstanten von 8 ms bis 9 ms. Bei größeren De- und Hyperpolarisationen reagiert der *h*-Mechanismus schneller mit

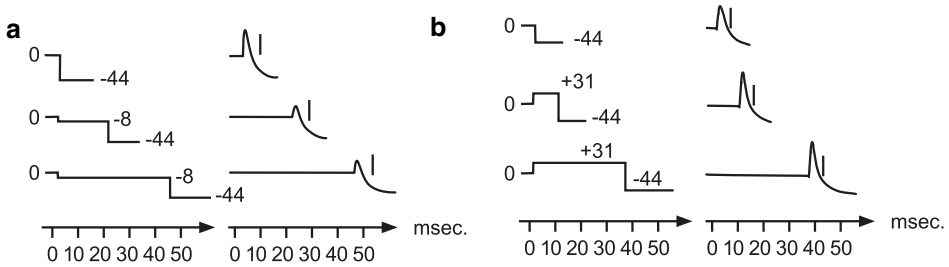


Abb. 8.21 Veränderung des kurzfristigen Stroms durch Vorschaltung einer Konditionierungsspannung von -8 mV in Abbildung (a) bzw. von 31 mV in Abbildung (b) unterschiedlicher Länge. Die linke Spalte zeigt jeweils den Spannungsverlauf, die rechte Spalte jeweils die Stromstärke, der vertikale Strich gibt die erwartete maximale Stromstärke ohne den Konditionierungsschritt an. (Abbildungen nach [42, Fig.1, Fig.2])

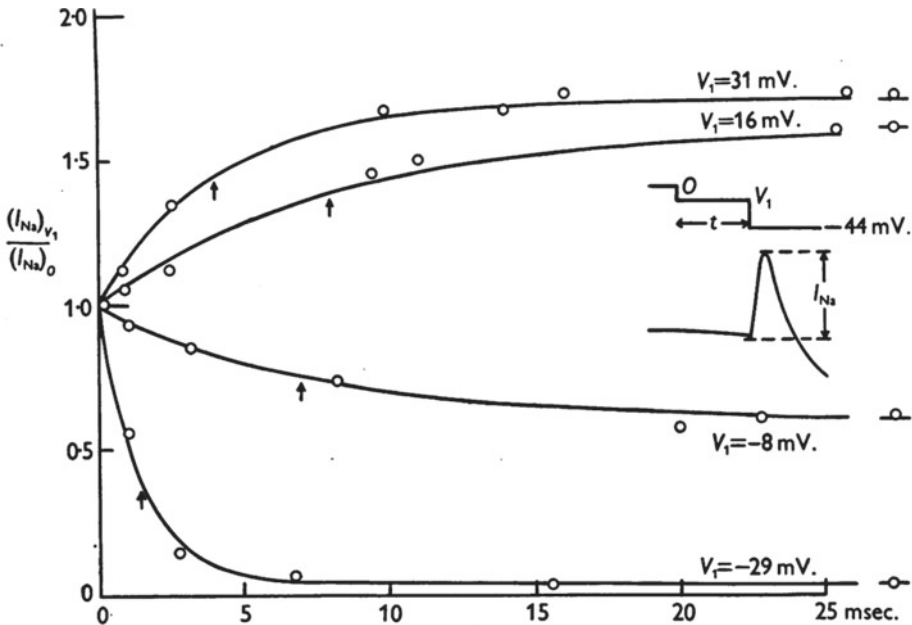


Abb. 8.22 Zeitlicher Verlauf der Änderung der Natriumdurchlässigkeit durch die Dauer der Vorkonditionierung für vier Konditionierungsspannungen, die an den Interpolationskurven vermerkt sind. Auf der horizontalen Achse ist die Dauer des Konditionierungsschritts angegeben, auf der vertikalen Achse ist die erreichte Stromstärke relativ zur unkonditionierten Stromstärke angegeben. Die Kreise geben die Messwerte wieder. Die Interpolationen entstehen aus $y(t) = y_{\infty} - (y_{\infty} - 1) \exp(-t/\tau_h)$ durch Anpassen der Parameter y_{∞} und τ_h an die gemessenen Werte. y_{∞} gibt dabei das Verhältnis von der Stromstärke zur unkonditionierten Stromstärke für $t = \infty$ an und τ_h hat die Rolle einer Zeitkonstante, die angibt, zu welcher Zeit die Hälfte der Änderung des Verhältnisses erreicht wird. Die Zeitkonstante ist durch einen Pfeil an den Graphen angedeutet. (Abbildung aus [42, Fig. 3], ©(1952) John Wiley & Sons)

Zeitkonstanten im Bereich von 1 ms bis 2 ms. In jedem Fall erreicht der h -Mechanismus seine langfristige Durchlässigkeit in einem Zeitraum von 40 ms.

H & H messen nun die *langfristige Durchlässigkeit* des h -Mechanismus durch eine Vorkonditionierung über einen Zeitraum von 40 ms bei unterschiedlichen Konditionierungsspannungen. Vier Messkurven sind in Abb. 8.23 wiedergegeben.

Nun tragen H & H das Verhältnis des konditionierten maximalen Stroms zum unkonditionierten maximalen Strom gegen die Konditionierungsspannung auf und erhalten die durch Kreise markierten Messwerte in Abb. 8.24 (vier dieser Messwerte stammen von den Kurven aus Abb. 8.23).

Die gemessenen Werte zeigen ein klares Grenzwertverhalten für große De- und Hyperpolarisationen und sind äußerst gut durch eine logistische Funktion interpolierbar. H & H führen die neue Variable h ein, welche die Öffnung des h -Mechanismus beschreibt und bestimmt, welche Durchlässigkeiten durch den m -Mechanismus überhaupt erreicht werden können. Die h -Variable nimmt Werte zwischen 0 und 1 an. Im Zustand $h = 0$ lässt sich

Abb. 8.23 Langfristiges Verhalten des h -Mechanismus bei unterschiedlichen Konditionierungsspannungen (Abbildung nach [42, Fig. 4])

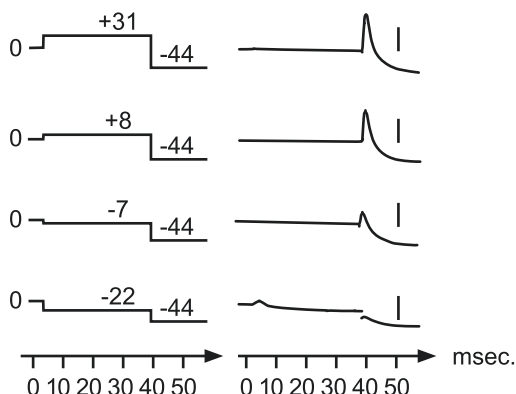
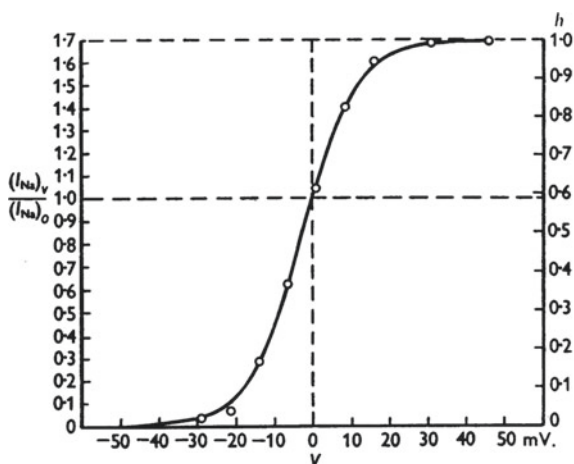


Abb. 8.24 Langfristige Leitfähigkeit des h -Mechanismus. Horizontale Achse: Auslenkung aus der Ruhespannung der Vorkonditionierung in mV. Linke vertikale Achse: Verhältnis des maximalen Stromflusses nach Vorkonditionierung zum maximalen Stromfluss ohne Konditionierung. Rechte vertikale Achse: Wert der h -Variablen, siehe (8.3) (Abbildung aus [42, Fig. 5], ©(1952) John Wiley & Sons)



die Membran nicht für einen Natriumstrom aktivieren. Dieser Zustand wird praktisch bei Depolarisationen größer als -25 mV langfristig angenommen. Im Zustand $h = 1$ lässt sich die Membran maximal für einen Natriumstrom aktivieren. Das ist praktisch langfristig für Hyperpolarisationen größer als 25 mV der Fall. H & H sprechen auch von $1 - h$ als einem Maß für die *Inaktivierung* der Membran und von h als einem Maß für die *Nicht-Inaktivierung* der Membran. Die gemessenen langfristigen Nicht-Inaktivierungen h lassen sich durch die Funktion

$$h_{\infty}(V) = \frac{1}{1 + \frac{\exp(V_h - V)}{7}} \quad (8.3)$$

interpolieren (die durchgezogene Kurve in Abb. 8.24). V_h ist dabei die Spannung, bei der h langfristig den Wert $1/2$ annimmt. In Experimenten an unterschiedlichen Axonen variiert V_h zwischen $-1,5\text{ mV}$ und $-3,5\text{ mV}$, die prinzipielle Form der h_{∞} -Kurve ist aber in jedem Fall gleich.

Wir fassen die Interpretation der Natriumleitfähigkeit zusammen:

- Sie wird bestimmt durch das Zusammenspiel von zwei Mechanismen, dem m -Mechanismus und dem h -Mechanismus.
- Der m -Mechanismus ist bei $V = 0\text{ mV}$ langfristig undurchlässig und wird mit steigender Depolarisation zunehmend durchlässig.
- Er passt sich relativ schnell und bei hohen Depolarisationen sehr schnell an die neue Spannung an.
- Der h -Mechanismus ist bei Hyperpolarisationen langfristig durchlässig und bei Depolarisationen langfristig undurchlässig.
- Er passt sich insbesondere um die Ruhespannung $V = 0\text{ mV}$ herum nur langsam an Spannungsänderungen an.

Mit diesen Ergebnissen und Interpretationen, die aus der *Voltage-Clamp*-Situation gewonnen wurden, lässt sich nun die Entstehung eines Aktionspotentials erklären.

8.3.7 Die qualitative Erklärung des Aktionspotentials

Im Ruhezustand der Membran, $V = 0\text{ mV}$, liegt zwar ein großer Natriumantrieb vor ($0\text{ mV} - (-112\text{ mV})$), die Membran ist aber durch den m -Mechanismus undurchlässig für Natrium. Der Kaliumantrieb ist gering ($0\text{ mV} - 12\text{ mV}$). Die Membran ist für Kalium nur wenig durchlässig, was zu einem geringen, auswärts gerichteten Kaliumstrom führt, der durch weitere Ionenströme, die als Leckströme zusammengefasst werden, kompensiert wird. Ein kurzer, von außen kommender Depolarisationsimpuls führt zu einer zunächst geringen Öffnung des m -Mechanismus, was zu einem einwärts gerichteten Natriumstrom führt. Wenn der anfängliche Impuls hinreichend groß ist, kann der anfängliche Natriumstrom die Depolarisation der Zelle aufrechterhalten, verstärkt sich durch die zunehmende

Natriumleitfähigkeit selber und führt zu einem großen Natriumstrom in die Zelle hinein, bis annähernd das Natriumgleichgewicht V_{Na} bei -112 mV erreicht ist. Dies ist die Phase des Upstrokes.

Nun erhöht sich zum einen aufgrund der hohen Depolarisation die Kaliumleitfähigkeit und es liegt ein hoher Kaliumantrieb vor. Das führt zu einem nach außen gerichteten Kaliumstrom. Gleichzeitig schaltet der h -Mechanismus, der bei hohen Depolarisationen deutlich schneller reagiert als in der Nähe der Ruhelage, die Membran für Natrium auf undurchlässig, was einen weiteren Natriumstrom unterbindet. Das ergibt die Phase der Repolarisation.

Weil die Kaliumleitfähigkeit sich nur mittelschnell anpasst, hyperpolarisiert die Membran, bis sie in die Nähe des Kaliumgleichgewichts bei $+12 \text{ mV}$ gelangt. Das ist die Phase der Hyperpolarisation.

Mit Hilfe der Leckströme wird dann in einem langsamen Prozess wieder die Ruhespannung angenommen. Dabei ist die Membran erst einmal nicht mehr durch einen kurzfristigen Reiz erregbar, weil der h -Mechanismus noch vom Ende der Upstroke-Phase her auf undurchlässig gestellt ist und in der Nähe der Ruhespannung lange benötigt, um seinen langfristigen Zustand, nämlich ungefähr halb durchlässig, zu erreichen.

8.4 Die mathematische Modellierung und die quantitative Berechnung der Aktionspotentials

Während diese qualitative Beschreibung vielleicht schon ganz plausibel klingt, gewinnt die Deutung von H & H ihre große Überzeugungskraft durch die anschließende mathematische Modellierung, die sie in dem Artikel [43] durchführen. Die Ergebnisse dieser Modellierung geben die experimentellen Befunde tatsächlich nicht nur qualitativ, sondern auch quantitativ gut wieder: Man kann den Verlauf eines Aktionspotentials berechnen und so Aktionspotentiale simulieren.

Aufgrund der Experimente beschreiben H & H die Axonmembran als Parallelschaltung von einem elektrischen Kondensator und drei „Bauteilen“, die für den Natriumstrom, den Kaliumstrom und die restlichen Ionenströme leitfähig sind, siehe Abb. 8.25. Der Gesamtstrom setzt sich dann aus dem kapazitiven Strom, dem Natriumstrom I_{Na} , dem Kaliumstrom I_{K} und dem Leckstrom I_{L} zusammen:

$$I = C \dot{V} + I_{\text{Na}} + I_{\text{K}} + I_{\text{L}}.$$

Die eingezeichneten Spannungsquellen sorgen für die unterschiedlichen elektrochemischen Antriebe. Gemäß den Experimenten scheinen die Ionenströme passiv zu sein, sie werden also angesetzt als ein Produkt aus einer Leitfähigkeit und dem elektrochemischen Antrieb:

$$I_{\text{Na}} = g_{\text{Na}} \cdot (V - V_{\text{Na}})$$

$$I_{\text{K}} = g_{\text{K}} \cdot (V - V_{\text{K}})$$

$$I_{\text{L}} = g_{\text{L}} \cdot (V - V_{\text{L}})$$

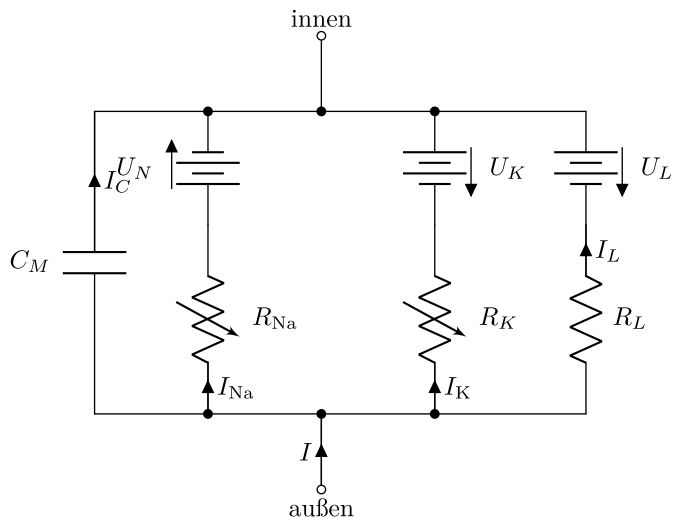


Abb. 8.25 Schaltzeichnung nach [43, Fig. 1], $R_{\text{Na}} = 1/g_{\text{Na}}$, $R_{\text{K}} = 1/g_{\text{K}}$, $R_{\text{L}} = 1/\hat{g}_{\text{L}}$, $U_{\text{Na}} = V_{\text{Na}} + U_r$, $U_{\text{K}} = V_{\text{K}} + U_r$, $U_{\text{L}} = V_{\text{L}} + U_r$

Die Natrium- und Kaliumleitfähigkeiten g_{Na} und g_{K} müssen gemäß den experimentellen Befunde modelliert werden, die Natrium- und Kalium-Gleichgewichtsspannungen werden experimentell gemessen oder über die Nernst-Gleichung bestimmt (was eine experimentelle Messung der Konzentrationen erfordert). Die Leitfähigkeit für die Leckströme wird von H & H als konstant angenommen und die zugehörige Gleichgewichtsspannung der Leckströme V_{L} am Ende so bestimmt, dass $V = 0 \text{ mV}$ tatsächlich ein Gleichgewichtszustand (also ein stationärer Punkt) des Modells wird. Die eigentliche Aufgabe ist also die Modellierung der Natrium- und Kaliumleitfähigkeiten g_{Na} und g_{K} . Wir beginnen wieder mit der Modellierung der Kaliumleitfähigkeit.

8.4.1 Die Modellierung der Kaliumleitfähigkeit

Es wird ein Modell benötigt, das Kurven wie in Abb. 8.15b und 8.18 erzeugt und bestenfalls auch eine physiologische Deutung zulässt. Dazu schreiben H & H der Membran eine maximale erreichbare Leitfähigkeit $\hat{g}_{\text{K}}$ für Kaliumströme zu. Durch einen Mechanismus, der durch die Variable n dargestellt wird, kann sich die Leitfähigkeit verringern. Die Variable n nimmt dabei Werte zwischen 0 (keine Leitfähigkeit) und 1 (maximale Leitfähigkeit) an. Aus den *Voltage-Clamp*-Versuchen schließen H & H, dass n sowohl von der Spannung V an der Membran als auch von der Zeit abhängt. Als mathematisches Modell kommt also eine Differentialgleichung für n in Betracht, bei der die Koeffizienten der Differentialgleichung von V abhängen.

Die erste Idee ist es die Gesamtleitfähigkeit für Kalium als ein einfaches Produkt anzusetzen, $g_K = \hat{g}_K \cdot n$. Wir betrachten die Kurven in Abb. 8.18 genauer: Der *S*-förmige Verlauf bei einer Depolarisation kann nicht durch eine lineare Differentialgleichung für n erzeugt werden, eine solche würde ja einen exponentiellen Verlauf erzwingen. Es gäbe nun die Möglichkeit, für n eine nichtlineare Differentialgleichung anzusetzen, z. B. eine, die quadratisch in n ist (wir wissen ja bereits, dass eine logistische Differentialgleichung Lösungen mit *S*-förmigem Verlauf erzeugen kann). H & H wählen aber einen anderen Weg, sie setzen die Leitfähigkeit mit einer *Potenz* von n an,

$$g_K = \hat{g}_K \cdot n^l, \quad (8.4)$$

und fordern dafür aber eine lineare Differentialgleichung für n , die sie in einer bemerkenswerten Form schreiben:

$$\dot{n} = \alpha_n(V)(1 - n) - \beta_n(V)n. \quad (8.5)$$

Dabei sollen die Koeffizienten α_n und β_n experimentell zu bestimmende Funktionen von V sein. Natürlich könnte man die Gl. (8.5) auch einfach in der Form $\dot{n} = \bar{\alpha} - \bar{\beta}n$ mit $\bar{\alpha} = \alpha$ und $\bar{\beta} = \alpha_n + \beta_n$ oder ähnlich notieren. H & H wählen aber ihre Darstellung, weil sie eine physiologische Deutung damit verbinden: Sie vermuten, dass die Änderung der Leitfähigkeit durch den Wechsel der Position spezieller Teilchen innerhalb der Membran zustande kommt. Damit diese speziellen Teilchen so empfindlich auf die angelegte Spannung V reagieren können, müssen sie elektrisch geladen sein oder zumindest einen elektrischen Dipol aufweisen, der sich im elektrischen Feld ausrichtet. n gibt dann den Anteil dieser Teilchen an, der auf leitfähig für Kaliumionen geschaltet ist, und $(1 - n)$ gibt den Anteil an, der auf nichtleitfähig geschaltet ist.

Wir werden gleich sehen, dass im Ansatz (8.4) die Potenz $l = 4$ als erste eine gute Übereinstimmung mit den gemessenen Kurven liefert. H & H deuten nun diese vierte Potenz so, dass in einem sehr kleinen räumlichen Bereich vier von diesen speziellen, für die Leitfähigkeit sorgenden Teilchen sich in der Position „leitfähig“ befinden müssen, damit die Membran tatsächlich an dieser Stelle durchlässig wird. (Wir werden im abschließenden Teil dieses Kapitels darauf eingehen, inwiefern diese physiologische Deutung von H & H durch spätere Experimente bestätigt oder widerlegt wurde.) Die Ratenparameter α_n und β_n geben an, mit welcher Rate sich die Teilchen von „nichtleitfähig“ auf „leitfähig“ bzw. von „leitfähig“ auf „nichtleitfähig“ umschalten. Nehmen wir einmal an, dass diese speziellen Teilchen negativ geladen sind und sich auf der Innenseite der Membran in der Position für „leitfähig“ befinden, an der Außenseite der Membran in der Position „nichtleitfähig“, dann sollte sich α_n mit zunehmender Depolarisation erhöhen (das negative Teilchen wird vom weniger negativen Membraninneren weniger abgestoßen bzw. bei starken Depolarisationen vom positiven Inneren stärker angezogen) und β_n sollte mit zunehmender Depolarisation abnehmen.

Warum also die Potenz $l = 4$ und keine höhere oder niedrigere? Die lineare Differentialgleichung (8.5) hat bekanntermaßen die Lösungen

$$n(t) = (n_0 - n_\infty) e^{\frac{t_0 - t}{\tau_n}} + n_\infty. \quad (8.6)$$

Dabei ist n_0 der Wert zur Zeit $t = t_0$, $n_\infty = \frac{\alpha}{\alpha + \beta}$ der stabile stationäre Wert und $\tau_n = \frac{1}{\alpha + \beta}$ die typische Zeit. n_∞ und τ sind dabei Funktionen der jeweils herrschenden Spannung V , also $n_\infty = n_\infty(V)$ und $\tau_n = \tau_n(V)$. Der korrespondierende Leitfähigkeitsverlauf ist dann durch die folgende Funktion gegeben (die globale Leitfähigkeitsskala $\hat{g}_K$ werden wir erst später festlegen):

$$g_K(t) = \left[(g_{K,0}^{1/l} - g_{K,\infty}^{1/l}) e^{(t_0 - t)/\tau_n} + g_{K,\infty}^{1/l} \right]^l, \quad (8.7)$$

wobei $g_{K,0}$ die Leitfähigkeit zur Zeit $t = t_0$ und $g_{K,\infty} = g_{K,\infty}(V)$ die auf lange Sicht bei der Spannung V angenommene Leitfähigkeit ist. Je nach gewähltem Modell ist $l = 1, 2, 3$ oder 4.

Wir laden die Messkurve aus Abb. 8.17 als Hintergrundbild in GeoGebra und versuchen die Parameter $g_{K,0}$, $g_{K,\infty}$ und τ (nach Augenmaß) mit Schieberegler für die beiden Kurventeile einzeln bestmöglich anzupassen. Für den ersten Kurventeil wählen wir $t_0 = 0$ ms, für den zweiten Kurventeil $t_0 = 5$ ms. Die Ergebnisse der Anpassung für $l = 1$, $l = 2$, $l = 3$ und $l = 4$ sind in Abb. 8.26 zu sehen. Die vierte Potenz ist die erste Potenz, die eine gute Übereinstimmung zwischen der gemessenen und der berechneten Kurve liefert. Höhere Potenzen verbessern die Anpassung nicht wesentlich. Damit ist das Modell für ein festes V bestimmt.

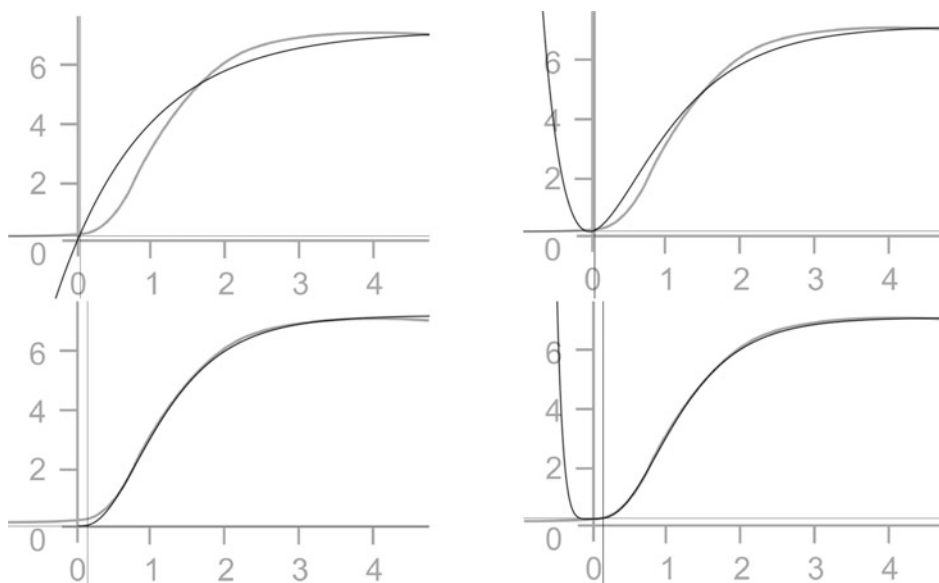


Abb. 8.26 Anpassung der Parameter $g_{K,0}$, $g_{K,\infty}$ und τ_n in (8.7) an die gemessene Kurve aus Abb. 8.18 für $l = 1, 2, 3$ und 4. Benutzte Parameterwerte für $l = 4$: $g_{K,0} = 0.02 \text{ mS cm}^{-2}$, $g_{K,\infty} = 6.58 \text{ mS cm}^{-2}$, $\tau_n = 0.61 \text{ ms}$

Wie hängen nun aber die Parameter τ und $g_{K,\infty}$ von V ab? Um diese Frage zu beantworten, bestimmen H & H die Parameter $\tau(V)$ und $g_{K,\infty}(V)$ für viele verschiedene Depolarisationen und Hyperpolarisationen V aus den aufgenommenen Leitfähigkeitskurven, so wie wir das eben für einen Wert von V gemacht haben, und interpolieren die so bestimmten Werte durch eine passende Funktion. Tatsächlich betrachten sie bei der Anpassung nicht direkt $\tau(V)$ und $g_{K,\infty}(V)$, sondern die Parameter $\alpha_n(V)$ und $\beta_n(V)$ aus (8.5), die aber über die Gleichungen

$$n_\infty = \sqrt[4]{\frac{g_{K,\infty}}{\hat{g}_K}}, \quad \alpha_n = \frac{n_\infty}{\tau_n} \quad \text{und} \quad \beta_n = \frac{1 - n_\infty}{\tau_n} \quad (8.8)$$

miteinander zusammenhängen. Zunächst einmal muss ein guter Wert für $\hat{g}_K$ gewählt werden: Er soll einerseits wirklich größer als alle gemessenen Kaliumleitfähigkeiten sein, andererseits auch nicht unnötig groß, weil die Variable n in diesem Fall nicht in die Nähe von 1 käme.

Aufgrund aller Experimente fixieren H & H diese Konstante auf $\hat{g}_K = 36 \text{ m S cm}^{-2}$. Hätten sie hier einen anderen Wert gewählt, würde das an der prinzipiellen Modellierung auch nichts ändern, es handelt sich eigentlich bloß um eine Entdimensionalisierung der Leitfähigkeit. Die aus vielen Messreihen zu unterschiedlichen Spannungen V durch Anpassung bestimmten Werte für α_n und β_n sind in Abb. 8.27 zusammengetragen. H & H interpolieren die Messwerte mit den Kurven

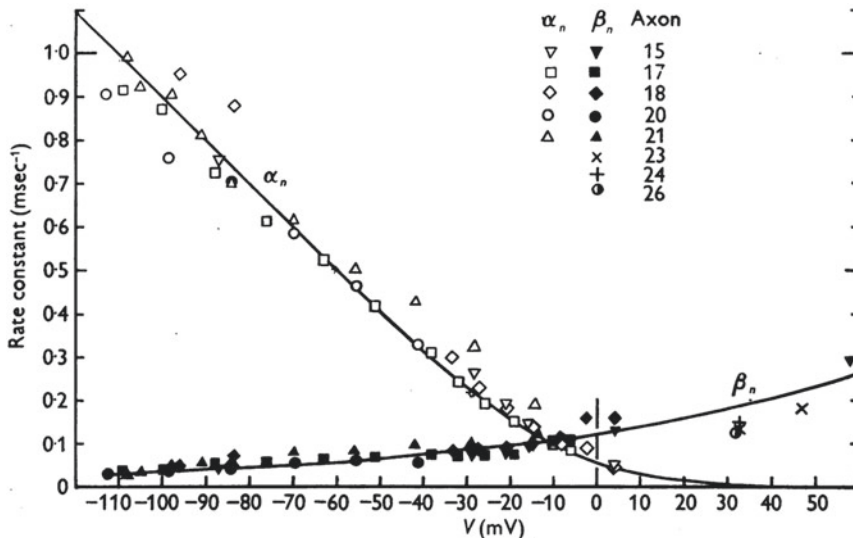


Abb. 8.27 Gemessene Ratenparameter und Interpolationskurve für α_n , β_n (Abbildung aus [43], ©(1952) John Wiley & Sons)

$$\alpha_n(V) = \frac{0.1 \text{ V} + 10 \text{ mV}}{\text{ms}} \frac{1}{10 \text{ mV}} \frac{1}{e^{\frac{V+10 \text{ mV}}{10 \text{ mV}}} - 1} \quad (8.9)$$

$$\beta_n(V) = \frac{0.125}{\text{ms}} e^{\frac{V}{80 \text{ mV}}}.$$

In Aufg. 8.3 soll diese Anpassung nachvollzogen werden. Warum wählen Hodgkin und Huxley genau diesen Ansatz für die angepasste Funktion? Es gibt ja viele andere Funktionen, die die Messpunkte ähnlich gut approximiert hätten. Die Autoren geben drei Gründe an:

- 1) Sie haben keine wesentlich einfachere Funktion gefunden, durch die diese Messwerte dargestellt werden können (Occam's razor).
- 2) Die gefundenen Funktionen α_n und β_n sollten über die Formeln (8.8) die gemessenen Werte von $g_{K,\infty}$ und τ_n für alle betrachteten Spannungen hinreichend genau reproduzieren. Dass das der Fall ist, soll ebenfalls in Aufg. 8.3 nachvollzogen werden.
- 3) Sie geben ein physikalisches Argument an [43, S. 510]:

“... it bears a close resemblance to the equation derived by Goldman (1943) for the movement of charged particle in a constant field. Our equations can therefore be given a qualitative physical basis if it is supposed that the variation of α_n and β_n with membrane potential arises from the effect of the electric field on the movement of a negatively charged particle which rests on the outside of the membrane when V is large and positive and on the inside when it is large and negative.”

Sie stellen diese Begründung aber auch sofort wieder infrage [43, S. 510f.]:

“The analogy cannot be pressed since α and β are not symmetrical about $E = 0$, as they should be if Goldman's theory held in a simple form. Better agreement might be obtained by postulating some asymmetry in the structure of the membrane, but this assumption was regarded as too speculative for profitable consideration.”

Auch hier wird klar, dass die Autoren immer auf der Suche nach inhaltlichen Interpretationen für ihre Ansätze sind, aber richtig einschlägige Interpretationen nicht immer gefunden werden können. Ansätze ohne inhaltliche Interpretation bezeichnet man als *Ad hoc*-Ansätze im Gegensatz zu den Ansätzen aus *First Principles*. Natürlich gibt es viele Abstufungen dazwischen.

8.4.2 Die Modellierung der Natriumleitfähigkeit

Die Modellbildung verläuft in weiten Teilen ähnlich zur Modellierung der Kaliumleitfähigkeit. Die Kurven der Natriumleitfähigkeit weisen jedoch schon ein qualitativ anderes Verhalten auf, siehe Abb. 8.15a. Auf einen steil aufsteigenden Teil folgt ein langsam abfallender Teil, beide Teile enthalten einen Wendepunkt. Klarerweise können solche Kurven

nicht Graphen von Lösungen autonomer Differentialgleichungen erster Ordnung sein, solche Lösungen sind ja stets monoton. Es gibt nun mehrere Möglichkeiten:

- 1) Wir benutzen weiter den Ansatz, dass g_{Na} nur von einer Variablen abhängt; diese muss dann aber eine Differentialgleichung mindestens zweiter Ordnung erfüllen, damit ein nicht-monotones Verhalten ermöglicht wird.
- 2) Wir führen eine zweite Variable ein, beide Variablen lösen eine Differentialgleichung erster Ordnung und die Leitfähigkeit entsteht aus dem Produkt.

H & H folgen der zweiten Variante und wir haben die beiden Variablen bereits kennengelernt, nämlich den m - und den h -Mechanismus. Eine erste Variable m soll für den steilen Anstieg am Anfang sorgen und dann von unten ihren stationären Zustand erreichen, die zweite Variable h startet von ihrem Anfangswert und fällt dann auf ihren stationären Wert, der nah bei null liegt. H & H setzen hier entsprechend zur Kaliumleitfähigkeit an:

$$\begin{aligned} g_{\text{Na}} &= \hat{g}_{\text{Na}} m^3 h \\ \dot{m} &= \alpha_m (1 - m) - \beta_m m \\ \dot{h} &= \alpha_h (1 - h) - \beta_h h \end{aligned} \quad (8.10)$$

Die physiologische Deutung sieht hier also zwei unterschiedliche Teilchensorten vor; drei m -Teilchen und ein h -Teilchen müssen sich in der Position „leitfähig“ befinden. Die Lösungen der beiden letzten Gleichungen aus (8.10) sind dann

$$\begin{aligned} m(t) &= m_{\infty} - (m_{\infty} - m_0) e^{-\frac{t}{\tau_m}} \\ h(t) &= h_{\infty} - (h_{\infty} - h_0) e^{-\frac{t}{\tau_h}}, \end{aligned}$$

wobei die Parameter über die folgenden Gleichungen zusammenhängen:

$$\begin{aligned} m_{\infty} &= \frac{\alpha_m}{\alpha_m + \beta_m}, \quad \tau_m = \frac{1}{\alpha_m + \beta_m}, \\ h_{\infty} &= \frac{\alpha_h}{\alpha_h + \beta_h}, \quad \tau_h = \frac{1}{\alpha_h + \beta_h}. \end{aligned} \quad (8.11)$$

Die theoretische Funktion, deren Parameter an die Messkurven angepasst werden muss, lautet

$$\begin{aligned} g_{\text{Na}}(t) &= \hat{g}_{\text{Na}} (m_{\infty} - (m_{\infty} - m_0) \exp(-t/\tau_m))^3 (h_{\infty} - (h_{\infty} - h_0) e^{-\frac{t}{\tau_h}}) \\ &= g' \left(1 - (1 - \varepsilon) e^{-\frac{t}{\tau_m}} \right)^3 \left((1 - \delta) e^{-\frac{t}{\tau_h}} + \delta \right) \end{aligned} \quad (8.12)$$

mit

$$g' = \hat{g}_{\text{Na}} m_{\infty}^3 h_0, \quad \varepsilon = \frac{m_0}{m_{\infty}} \quad \text{und} \quad \delta = \frac{h_{\infty}}{h_0}.$$

Für jede Messung mit einer festen Spannung V werden die Parameter in Abhängigkeit von V bestimmt. Eine Anpassung, die vom Autor dieses Buches mit GeoGebra vorgenommen wurde, ist in Abb. 8.28 dargestellt. Die aus vielen Experimenten zusammengetragenen Werte sind in Abb. 8.29 dargestellt und werden von H & H durch die folgenden Kurven interpoliert:

$$\begin{aligned} \alpha_m &= \frac{1}{\text{ms}} \frac{(V + 25 \text{ mV})}{10 \text{ mV}} \frac{1}{\exp\left(\frac{V+25 \text{ mV}}{10 \text{ mV}}\right) - 1} & \beta_m &= \frac{4}{\text{ms}} \exp(V/18 \text{ mV}) \\ \alpha_h &= \frac{0,07}{\text{ms}} \exp(V/20 \text{ mV}) & \beta_h &= \frac{1}{\text{ms}} \frac{1}{\exp\left(\frac{V+30 \text{ mV}}{10 \text{ mV}} + 1\right)} \end{aligned} \quad (8.13)$$

Auch hier gelten die zum Abschluss der Modellierung der Kaliumleitfähigkeit gemachten Bemerkungen sinngemäß.

8.4.3 Simulationen mit Hilfe des mathematischen Modells

Wir stellen noch einmal das komplette Modell zusammen: Es besteht aus den vier Variablen V , h , m und n , deren zeitliche Änderungsraten durch die Differentialgleichungen

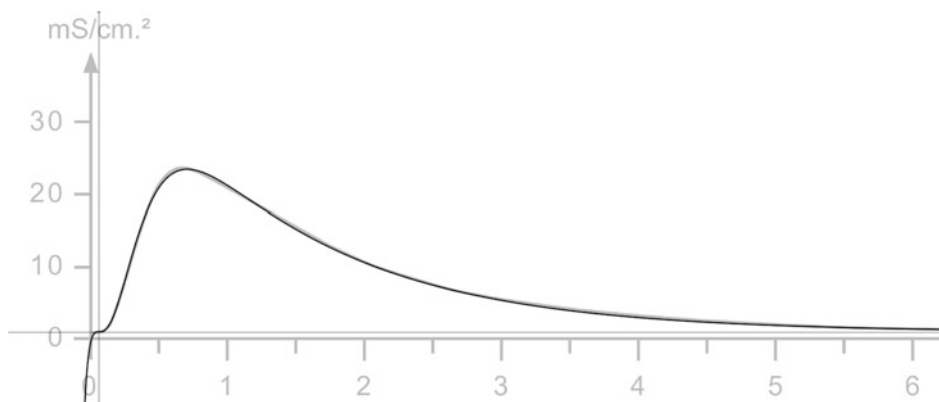


Abb. 8.28 Anpassung der Parameter aus (8.12) an die gemessene Kurve zur Spannung $V = -70 \text{ mV}$ aus Abb. 8.15a. Dafür wurde die gemessene Kurve als Hintergrundbild in GeoGebra eingefügt und an den Achsen ausgerichtet. Die Parameter g' , τ_m , τ_h , ε und δ wurden als Schieberegler eingerichtet und nach Augenmaß eingestellt. Benutzte Parameterwerte: $g' = 43.40 \text{ mS cm}^{-2}$, $\tau_m = 0.448 \text{ ms}$, $\tau_h = 2.61 \text{ ms}$, $\varepsilon = \delta = 0$

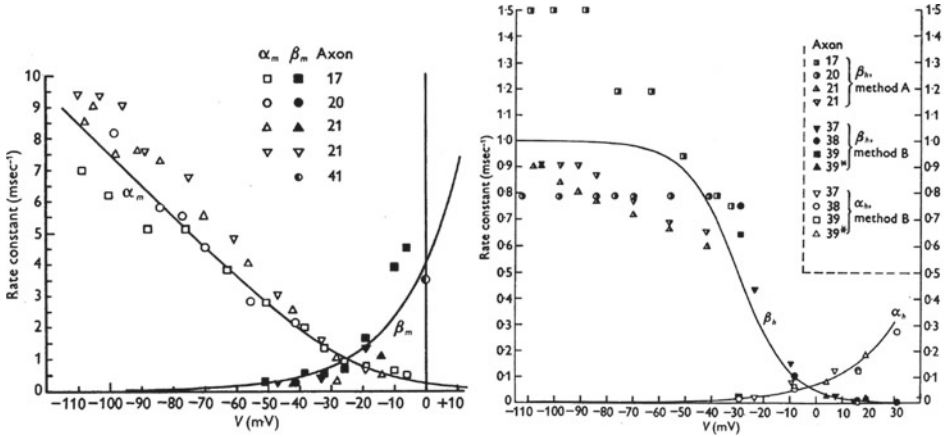


Abb. 8.29 Gemessene Ratenparameter und Interpolationskurve für α_m, β_m (links) und α_h, β_h (rechts) (Abbildungen aus [43, Abb. 4 und Abb. 5], ©(1952), John Wiley & Sons)

$$\begin{aligned}
 \dot{V} &= \frac{1}{C} (-\hat{g}_K n^4 (V - V_K) - \hat{g}_{Na} m^3 h (V - V_{Na}) - \hat{g}_L (V - V_L)) \\
 \dot{h} &= \alpha_h (1 - h) - \beta_h h \\
 \dot{m} &= \alpha_m (1 - m) - \beta_m m \\
 \dot{n} &= \alpha_n (1 - n) - \beta_n n
 \end{aligned} \tag{8.14}$$

bestimmt sind, wobei die Gleichungen für die Ratenfunktionen in (8.13) und (8.9) gegeben sind.

Die Gleichgewichtsspannung V_L für den Leckstrom wird so bestimmt, dass der Zustand $(V, h, m, n) = (0, h_\infty(0), m_\infty(0), n_\infty(0))$ stationär ist. Durch Umstellen der ersten Gleichung in (8.14) führt das auf

$$V_L = \frac{1}{\hat{g}_L} (\hat{g}_K n_\infty(0)^4 V_K + \hat{g}_{Na} m_\infty(0)^3 h_\infty(0) V_{Na}) .$$

Die von H & H benutzten, experimentell bestimmten Parameterwerte sind in Tab. 8.2 zusammengefasst. Die entscheidende Frage ist nun, ob dieses mathematische Modell die Situation an einer Axonmembran tatsächlich wiedergibt, also insbesondere Phänomene wie das Aktionspotential, die Reizschwelle und die Nicht-Erregbarkeit in der Phase der Hyperpolarisation beinhaltet, und das Profil eines berechneten Aktionspotentials, so es denn eines gibt, hinreichend genau mit den gemessenen Aktionspotentialen übereinstimmt.

Wir bemerken zunächst einmal, dass das Gleichungssystem (8.14) zu gegebenen Anfangsdaten (V_0, h_0, m_0, n_0) stets eine eindeutige Lösung hat, weil das System zusammen mit (8.13) und (8.9) die Voraussetzungen des Satzes von Picard-Lindelöf erfüllt. Dazu muss man sich davon überzeugen, dass die vermeintlichen Definitionslücken von α_n und α_m bei

Tab.8.2 Die von H & H für die numerischen Simulationen von Aktionspotentialen benutzten Parameterwerte

C	$\hat{g}_{\text{Na}}$	$\hat{g}_{\text{K}}$	$\hat{g}_L$	V_{Na}	V_{K}
1 F	120 m S cm ⁻²	36 m S cm ⁻²	0.3 m S cm ⁻²	-112 mV	12 mV

$V = 10$ mV bzw. $V = 25$ mV hebbbar sind und die Fortsetzung beliebig oft in 0 differenzierbar ist, siehe (8.13), (8.9) und Aufg. 8.4.

Natürlich ist das System nicht explizit lösbar, und da es vierdimensional ist, können wir nicht wie gewohnt das Phasenporträt analysieren, zumal das auch keine quantitativen Vergleiche ermöglichen würde. Es bleibt also nichts anderes übrig, als numerisch zu lösen. Im Gegensatz zu H & H befinden wir uns in der komfortablen Situation, dass jeder heute übliche Computer über eine vielfach höhere Rechenpower verfügt, als sie 1952 zur Verfügung stand. Wir verwenden ein gemischtes Euler-Verfahren und implementieren dieses in ein Tabellenkalkulationsprogramm. Dazu legen wir eine Schrittweite Δt fest und berechnen aus gegebenen Anfangswerten (t_0, V_0, h_0, m_0, n_0) iterativ die Lösung:

$$\begin{aligned}
 t_{k+1} &= t_k + \Delta t \\
 V_{k+1} &= V_k + \frac{\Delta t}{C} (-\hat{g}_{\text{Na}} m_k^3 h_k (V_k - V_{\text{Na}}) - \hat{g}_{\text{K}} n_k^4 (V_k - V_{\text{K}}) - \hat{g}_L (V_k - V_L)) \\
 h_{k+1} &= h_k + \Delta t (\alpha_h(V_{k+1})(1 - h_k) - \beta_h(V_{k+1})h_k) \\
 m_{k+1} &= m_k + \Delta t (\alpha_m(V_{k+1})(1 - m_k) - \beta_m(V_{k+1})m_k) \\
 n_{k+1} &= n_k + \Delta t (\alpha_n(V_{k+1})(1 - n_k) - \beta_n(V_{k+1})n_k)
 \end{aligned} \tag{8.15}$$

Das Verfahren ist gemischt, weil zur Berechnung der Koeffizienten α und β im $k + 1$ -Schritt bereits die neu berechnete Spannung V_{k+1} verwendet wird. Die Ergebnisse von drei Simulationen sind in Abb. 8.30 wiedergegeben. Zeile 1 zeigt ein ausgelöstes Aktionspotential bei einem Spannungsimpuls von $V_0 = -16$ mV. Die zweite Zeile zeigt den zeitlichen Verlauf für einen Impuls von $V_0 = -5$ mV, der unterhalb der Reizschwelle bleibt. Für die dritte Zeile wurden die Anfangswerte für die Torvariablen h, m und n aus der ersten Simulation nach 5 ms entnommen. Dies spiegelt die Situation wider, in der nach einem anfänglichen Impuls von -16 mV nach 5 ms ein zweiter Impuls der gleichen Amplitude eintrifft, der jedoch zu keinem Aktionspotential führt, sondern abklingt.

Das mathematische Modell enthält also die drei angesprochenen Phänomene, wobei die Reizschwelle zwischen -5.5 mV und -5.6 mV liegt, also etwas niedriger als in den Experimenten von H & H an den Axonen. Weitere eher kleine Unterschiede zwischen den Spannungen und den gemessenen Spannungen werden in [43] diskutiert.

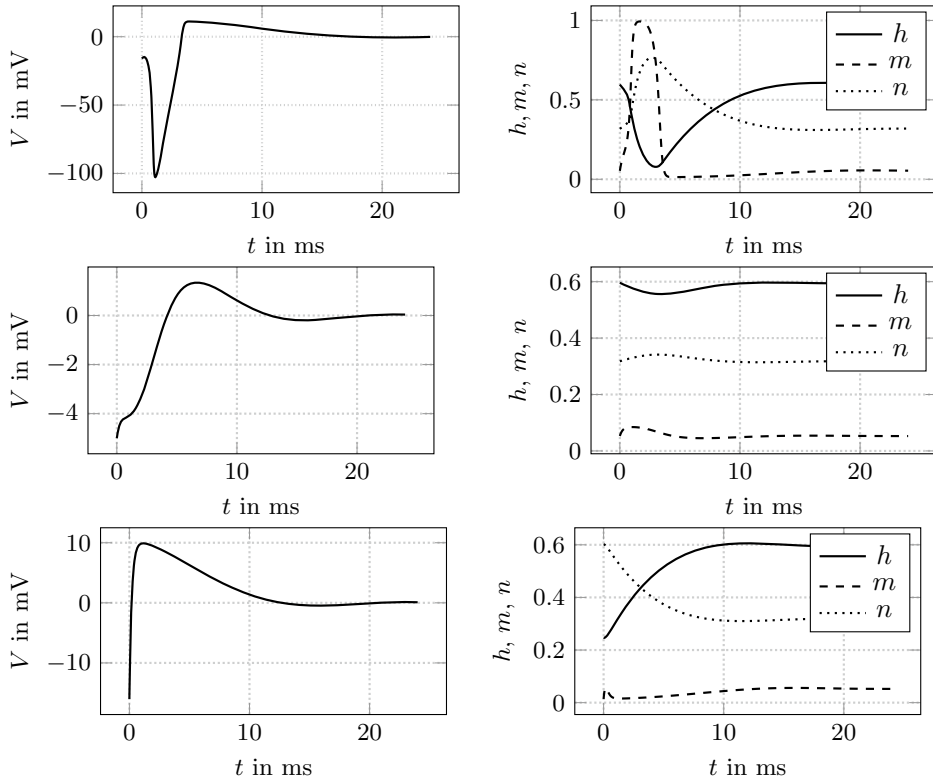


Abb. 8.30 Simulationen mit $\Delta t = 0.002$ ms für die ersten 2000 Iterationsschritte und $\Delta t = 0.02$ ms für die nächsten 1000 Iterationsschritte. Erste Zeile: Ausbildung eines Aktionspotentials; $(V_0, h_0, m_0, n_0) = (-16 \text{ mV}, h_\infty(0), m_\infty(0), n_\infty(0))$. Zweite Zeile: Reiz unterhalb der Reizschwelle; $(V_0, h_0, m_0, n_0) = (-5 \text{ mV}, h_\infty(0), m_\infty(0), n_\infty(0))$. Dritte Zeile: Membran ist inaktiviert; $(V_0, h_0, m_0, n_0) = (-16 \text{ mV}, 0.247, 0.014, 0.602)$. Die Iteration wurde mit dem Tabellenkalkulationsprogramm LibreOffice 4.3.4.1 durchgeführt

8.4.4 Die Erkenntnisse von H & H und das heutige Lehrbuchwissen

Die Arbeiten von H & H von Anfang der 50er Jahre stellten einen Durchbruch für die Elektrophysiologie dar und wurden 1969 mit dem Nobelpreis für Medizin und Physiologie gewürdigt. Interessant ist insbesondere die Frage, ob die Vermutungen bezüglich der genauen Mechanismen der selektiven Membranleitfähigkeit sich durch spätere Experimente bestätigt haben und welche Vermutungen widerlegt werden konnten. Tatsächlich ist es so, dass sich aufgrund verbesserter experimenteller Methoden heute Ströme durch einzelne Moleküle messen lassen (*Patch-Clamp-Technik*, Nobelpreis für Medizin 1991, siehe z. B. [74]) und durch bildgebende Verfahren die Struktur einzelner Moleküle auflösen lässt. Dabei haben sich die Vermutungen von H & H im Wesentlichen bestätigt. Die Leitfähigkeit für die

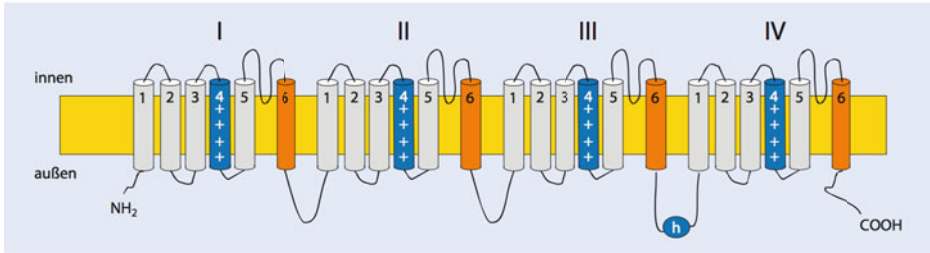


Abb. 8.31 Struktur eines spannungsgesteuerten Na⁺-Kanals bestehend aus einer α-Untereinheit mit vier repeats. Die TMSs 4 jedes einzelnen repeat besitzt eine Reihe von positiv geladenen Gruppen, die als Spannungssensor fungieren, die intrazelluläre Schleife *h* dient als Inhibitor. (Abbildung und Text aus [75], ©(2018), Springer-Verlag GmbH Deutschland)

verschiedenen Ionenströme erfolgt durch spezialisierte Moleküle, die in die Doppellipidschicht der Membran eingelassen sind. Heutzutage kennt man viele verschiedene Moleküle, die z.B. für den spannungsabhängigen Natriumionenstrom verantwortlich sind. In diese Moleküle sind tatsächlich einzelne Bestandteile eingebaut, welche spannungsabhängig die Leitfähigkeit durch ihre Position im Molekül steuern, siehe Abb. 8.31. Auch das mathematische Modell (8.14) findet sich heute noch genauso in aktuellen Lehrbüchern, siehe z. B. [75].

8.5 Toy-Models – das vereinfachte Modell nach FitzHugh-Nagumo

In Aufg. 8.5 soll nachgewiesen werden, dass das Phänomen „Aktionspotential“ auch bei mehr oder weniger geringfügigen Änderungen des H & H-Systems erhalten bleibt. Wünschenswert ist ein möglichst einfaches System mit ähnlichen Eigenschaften, das auch anderen Zugangsweisen als nur der numerischen zugänglich ist. Gesucht ist gewissermaßen ein Minimalsystem, das Aktionspotentiale enthält. Ein besonders bekanntes System ist von Richard FitzHugh aufgestellt worden, siehe [29], das wir in diesem Kapitel studieren wollen. Eine physikalische Realisierung dieses Systems in Form eines Stromkreises wurde von Nagumo et al. vorgeschlagen, siehe [2]. Eine zu diesem System hinführende Heuristik könnte folgendermaßen aussehen: Im H & H-System haben die drei Torvariablen h , m und n unterschiedliche Zeitkonstanten: Die Variable m ist die mit Abstand schnellste Variable, die sich „fast instantan“ auf ihren stationären Wert einstellt. Wenn das passiert ist, gilt $\dot{m} \approx 0$. Die Variable h ist die langsamste Variable; auf einer schnellen Zeitskala gilt also $h \approx \text{const.}$ Es bleibt eine Differentialgleichung für die mittelschnelle Variable n . In einem gewissen Sinne betrachten wir damit den simultanen asymptotischen Grenzwert $\tau_m/\tau_n \rightarrow 0$ und $\tau_n/\tau_h \rightarrow 0$. Aus der Forderung $\dot{m} = 0$ erhalten wir die algebraische Gleichung

$$m = m_\infty(V) = \frac{\alpha_m(V)}{\alpha_m(V) + \beta_m(V)}$$

und damit ein Differentialgleichungssystem für die zwei Variablen V und n :

$$\begin{aligned}\dot{V} &= CV - \hat{g}_{\text{Na}} m_{\infty}(V)^3 h_{\infty}(0) - g_L(V - V_L) - \hat{g}_K n^4(V - V_K) \\ \dot{n} &= \alpha_n(1 - n) - \beta_n n\end{aligned}\quad (8.16)$$

Die komplexen Gleichungen in (8.16) werden nun durch einfachere ersetzt: Wir entwickeln in der ersten Gleichung den nur von V und nicht von n abhängigen Teil um den stationären Punkt $V = 0$ in ein Polynom dritter Ordnung. Wir erhalten dabei einen Ausdruck der Form $AV(B - V)(V - C)$ mit bestimmten Konstanten A , B und C (unterschiedlicher Dimension). In Aufg. 8.5 soll gezeigt werden, dass die vierte Potenz in der Torvariablen n nicht wesentlich für die Herausbildung von Aktionspotentialen ist. Wir vereinfachen auch hier, indem wir n^4 schlicht durch n ersetzen. Ferner nehmen wir der Einfachheit halber an, dass $n_{\infty}(0) = 0$ gilt, dann können wir den quadratischen Term mit nV vernachlässigen. Schlussendlich übernehmen wir von der zweiten Gleichung lediglich die linearen Anteile in V und n (α_n und β_n sind Funktionen von V) und erhalten insgesamt ein vereinfachtes System der Form

$$\begin{aligned}\dot{V} &= AV(B - V)(V - C) - Dn \\ \dot{n} &= EV - Fn\end{aligned}$$

mit bestimmten Konstanten A , B , C , D , E und F . Durch eine geeignete Entdimensionalisierung können wir drei Parameter eliminieren (Skalen für V , n und t können gewählt werden). Die Entdimensionalisierung kann so gewählt werden, dass wir zu dem folgenden System gelangen:

$$\begin{aligned}\varepsilon \dot{v} &= f(v, w) \\ \dot{w} &= g(v, w)\end{aligned}\quad (8.17)$$

mit

$$\begin{aligned}f(v, w) &= v(a - v)(v - 1) - w + I_e \\ g(v, w) &= v - bw.\end{aligned}\quad (8.18)$$

Dabei ist v die entdimensionalisierte Version von V , also die schnelle Spannungsvariable, und w die entdimensionalisierte Version von n , also die im Vergleich damit langsame Torvariable. Die unterschiedlichen Zeitskalen der Variablen v und w spiegeln sich im Parameter ε wider, der klein sein soll. Die anderen dimensionslosen Parameter a und b sind positiv. Außerdem haben wir einen weiteren Parameter I_e in die Differentialgleichung für v hinzugefügt. Dieser Parameter steht für einen zusätzlichen konstanten Strom durch die Membran und wird im weiteren Verlauf eine gewisse Rolle spielen, weil er ein neues Phänomen hervorbringt. Insgesamt konstituieren die Gl. (8.17) und (8.18) das *Toy-Model* von FitzHugh-Nagumo.

Wir beginnen mit ein paar asymptotischen Betrachtungen analog zu Abschn. 7.2.2. Im asymptotischen Limes $\varepsilon = 0$ erhalten wir die Gleichungen

$$\begin{aligned} 0 &= f(v_0, w_0) \\ \dot{w}_0 &= g(v_0, w_0). \end{aligned} \quad (8.19)$$

Wechseln wir von den hier gewählten Variablen τ, v, w auf die Variablen der kleinen Zeitskala $T = \tau/\varepsilon$, $V(T) = v(\varepsilon T)$ und $W(T) = w(\varepsilon T)$, siehe auch 7.31, erhalten wir das System

$$\begin{aligned} \dot{V} &= f(V, W) \\ \dot{W} &= \varepsilon g(V, W) \end{aligned}$$

mit dem asymptotischen Grenzwert

$$\begin{aligned} \dot{V}_0 &= f(V_0, W_0) \\ \dot{W}_0 &= 0. \end{aligned} \quad (8.20)$$

Das zugrunde liegende abstrakte Bild ist das folgende, vergleiche z. B. Abb. 8.32: Wenn wir nicht auf oder in der Nähe einer v -Nullkline sind, ist die Funktion f von der Größenordnung 1 und $\dot{v}$ von der Größenordnung $1/\varepsilon$, also um einen Faktor 1000 oder ähnlich größer als $\dot{w}$. Die Lösung läuft also annähernd horizontal nach rechts oder links, je nach Vorzeichen von f , bis sie auf eine Nullkline von v trifft. In dieser Phase können wir die schnellen Variablen benutzen und der asymptotische Limes $\varepsilon = 0$, $\dot{W}_0 = 0$, also $W_0 = \text{const.}$ und $\dot{V}_0 = f(V_0, \text{const.})$, liefert eine gute Approximation. Ist die Lösung in der Nähe oder

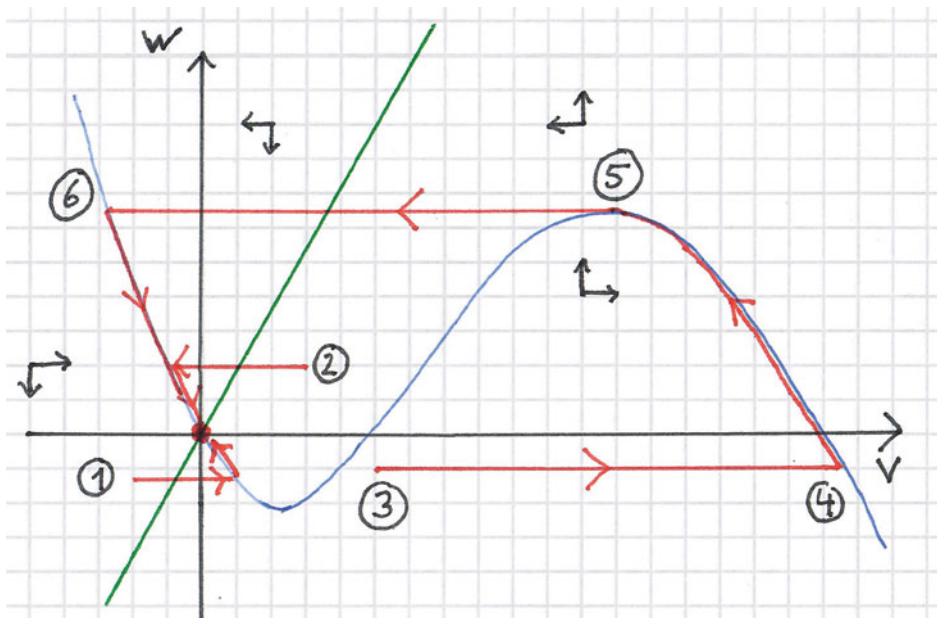


Abb. 8.32 Die Auslenkungen ① und ② kehren direkt zum stationären Punkt zurück, Auslenkung ③ führt zum Durchlaufen eines Aktionspotentials. Blau: v -Nullkline, grün: w -Nullkline, rot: Trajektorien

auf einer Nullkline, dann bleibt sie (wenn die Richtungen das erlauben) „an der Nullkline kleben“. Sie läuft also langsam beinahe auf der Nullkline entlang, solange die Richtungen das erlauben. Wir können jetzt den asymptotischen Limes der langsamen Variablen nutzen, wobei die Differentialgleichung für v durch die algebraische Gleichung $f(v_0, w_0) = 0$ ersetzt wird. Die nullte Approximation verläuft also buchstäblich in der v -Nullkline. Gelangt die Lösung nun an einen Punkt, an dem sie die v -Nullkline verlassen muss, siehe z. B. Punkt ⑤ in Abb. 8.32, ist die v -Ableitung wieder dominierend, die Lösung verläuft so gut wie waagrecht im Phasenporträt und lässt sich am besten in den schnellen Variablen beschreiben.

Wir kehren zum konkreten System (8.18) von FitzHugh zurück und werden sehen, dass für verschiedene Werte von I_e sehr vielseitige Phänomene in diesem einfachen Modell verborgen sind.

8.5.1 Ein Aktionspotential

Wir beginnen mit $I_e = 0$. Für die Nullklinen haben wir

$$\begin{aligned}\dot{v} = 0 &\Leftrightarrow w = v(a - v)(v - 1), \\ \dot{w} = 0 &\Leftrightarrow w = \frac{1}{b}v.\end{aligned}$$

Wir betrachten hier den Fall, dass es nur einen stationären Punkt gibt, nämlich $P_0^* = (0, 0)$. Das ist genau dann der Fall, wenn die quadratische Gleichung $(a - u)(u - 1) - \frac{1}{b} = 0$ keine reellen Lösungen hat. Dieser Fall wird unser Spielzeugmodell für die Ausbildung von Aktionspotentialen sein. Dabei ist wie angesprochen v vergleichbar mit der Spannungsvariablen und w mit den Torvariablen. Für die Signaturmatrix ergibt sich

$$J(P_0^*) \begin{pmatrix} - & - \\ + & - \end{pmatrix},$$

also $\det J(P_0^*) > 0$ und $\text{tr} J(P_0^*) < 0$. Damit ist P_0^* ein stabiler Knoten oder Strudel. Aus den Richtungen ergibt sich, dass P_0^* ein stabiler Knoten sein muss, die Trajektorien können sich nah am stationären Punkt nicht um den stationären Punkt drehen. Wir betrachten jetzt bestimmte Auslenkungen aus dem stationären Punkt und die entstehenden Lösungen und Trajektorien: Zunächst betrachten wir eine kleine Auslenkung, siehe ① und ② in Abb. 8.32. Der Startpunkt ① habe die Koordinaten $(\hat{v}, \hat{w})$. Er befindet sich nicht auf der v -Nullkline, also nutzen wir die schnellen Variablen T, U, V . Das U -Wachstum ist positiv, folglich läuft die Lösung im Phasenporträt nach rechts, bis sie auf die Nullkline trifft. Die Lösung lässt sich in dieser Phase gut durch die Lösung der Differentialgleichung $\dot{V}_0 = V_0(a - V_0)(V_0 - 1) - \hat{w}$ mit dem Anfangswert $\hat{v}$ darstellen. Nach einer kurzen Zeitspanne ist die Lösung in der Nähe der Nullkline. Die w -Richtung ist in diesem Bereich positiv, folglich läuft die Lösung annähernd auf der v -Nullkline nach oben zum stationären Punkt P_0^* . In dieser Phase ist die

Lösung gut approximierbar durch eine geeignete Lösung des Systems

$$\begin{aligned}0 &= v(a - v)(v - 1) - w \\ \dot{w} &= v - bw.\end{aligned}$$

Für die Lösung zur Auslenkung ② gilt Entsprechendes.

Nun betrachten wir die Auslenkung ③. Die Lösung durchläuft vier Phasen:

- (I) Die Lösung entwickelt sich schnell zum Punkt ④. Die Spannungsvariable v wächst dabei auf einen Wert nah an 1 an, während die Torvariable w annähernd konstant bleibt. Dies entspricht der Upstroke-Phase.
- (II) Die Lösung läuft im Phasenporträt langsam und annähernd auf der v -Nullkline vom Punkt ④ zum lokalen Hochpunkt der v -Nullkline, dem Punkt ⑤.
- (III) Die Lösung muss sich an diesem Punkt aufgrund der w -Richtung von der Nullkline lösen. Sie entwickelt sich schnell und nahezu waagerecht nach rechts, bis sie am Punkt ⑥ wieder auf die Nullkline trifft. Die Spannungsvariable v fällt unter null ab. Die Phasen (II) und (III) entsprechen der Repolarisationsphase.
- (IV) Die Lösung läuft langsam auf der Nullkline in den stationären Punkt. Das entspricht der Phase der Hyperpolarisation. Auch in diesem *Toy-Model* ist das System in dieser Phase nur sehr schlecht erregbar: Die Störung müsste den Zustand ja über den aufsteigenden Teil der v -Nullkline auslenken, die aufgrund der kubischen Form einen größeren Abstand als in der Umgebung des stationären Punkts P_0^* hat.

Eine numerisch berechnete Trajektorie mit zugehörigem Spannungsverlauf ist in Abb. 8.33 wiedergegeben. Es ist bemerkenswert, dass der Lösungsverlauf ab dem Erreichen von Punkt ④ immer annähernd gleich aussieht: Unabhängig von der genauen Anfangsauslenkung wird immer dasselbe Aktionspotential durchlaufen. Ursache dafür ist $\varepsilon \ll 1$. Das sollen Sie in Aufg. 8.6 auch numerisch nachvollziehen.

8.5.2 Ein biochemischer Schalter

Wir betrachten jetzt den Fall $I_e = 0$ mit drei stationären Punkten, siehe Abb. 8.34. Für die Signaturmatrizen erhalten wir

$$J(P_0^*) = \begin{pmatrix} - & - \\ + & - \end{pmatrix}, \quad J(P_1^*) = \begin{pmatrix} + & - \\ + & - \end{pmatrix} \quad \text{und} \quad J(P_2^*) = \begin{pmatrix} - & - \\ + & - \end{pmatrix}.$$

Bezüglich P_1^* können wir nicht allein aus der Signaturmatrix heraus entscheiden, ob es sich um einen stabilen oder instabilen Punkt handelt. Aus der asymptotischen Betrachtung (waagerechter Verlauf bzw. Verlauf in der v -Nullkline) können wir aber schließen (rational, nicht deduktiv), dass für hinreichend kleine ε P_1^* instabil ist. Dieses System verhält sich

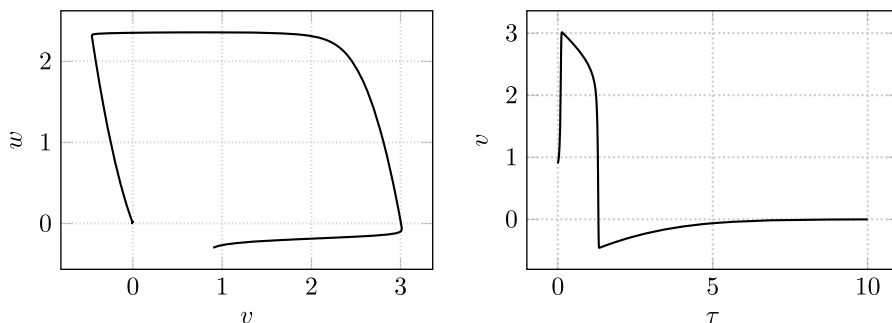


Abb. 8.33 Numerisch berechnetes Aktionspotential mit vier Phasen. Parameterwerte: $a = 3$, $b = 0.4$, $I_e = 0$. Verhältnis der Zeitskalen: $\varepsilon = 0.05$, explizites Euler-Verfahren mit Schrittweite $\Delta\tau = 0.005$ und Anfangswert $v_0 = 0.9$, $w_0 = -0.3$. Links: die Lösungstrajektorie im Phasenporträt. Rechts: die Lösung im τ - v -Diagramm. Die Phasen II-IV laufen immer annähernd gleich ab, unabhängig von der (hinreichend großen) Anfangsauslenkung. Die Iteration wurde mit dem Tabellenkalkulationsprogramm LibreOffice 4.3.4.1 durchgeführt

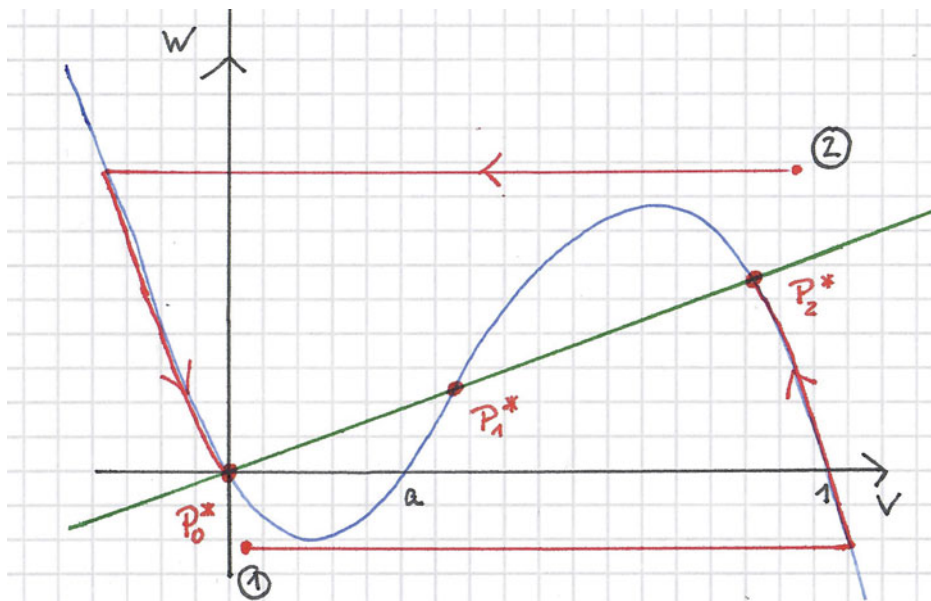


Abb. 8.34 Drei stationäre Punkte, ein biochemischer Schalter. Die Auslenkung ① aus dem Zustand P_0^* entwickelt sich in den Zustand P_2^* . Die Auslenkung ② aus dem Zustand P_2^* entwickelt sich in den Zustand P_0^*

wie ein *biochemischer Schalter*: Es gibt zwei stabile stationäre Zustände, nämlich P_0^* und P_2^* . Wird aus diesen Zuständen über eine Schwelle hinaus ausgelenkt, so entwickelt sich das System in den jeweiligen anderen Zustand. Numerische Berechnungen sind in Abb. 8.35 dargestellt.

8.5.3 Eine Folge von Aktionspotentialen: das neuronale Feuern

Nun betrachten wir die Situation mit $I_e > 0$ und a, b so, dass die Nullklinen nur einen Schnittpunkt haben und dieser im aufsteigenden Teil der kubischen Funktion liegt (siehe Abb. 8.36). Aus der asymptotischen Situation heraus schließen wir, dass P_0^* instabil ist. Außerdem enthält der in Abb. 8.36 schraffierte Bereich eine Halbtrajektorie. Der stationäre Punkt im Inneren lässt sich ausschneiden und damit erfüllt das System die Voraussetzungen des Satzes 6.3 von Poincaré-Bendixson. Also konvergieren alle nichtstationären Lösungen gegen einen Grenzzyklus und nähern sich einer periodischen Lösung an. Bei diesen periodischen Lösungen, die sich als eine gleichmäßige Abfolge von Aktionspotentialen deuten lassen, spricht man vom *neuronalen Feuern*. Im Fall $\varepsilon \ll 1$ hat die periodische Grenzlösung vier unterscheidbare Phasen in denen sich schnelle und langsame Entwicklungen abwechseln, siehe die Ergebnisse numerischer Berechnungen in Abb. 8.37.

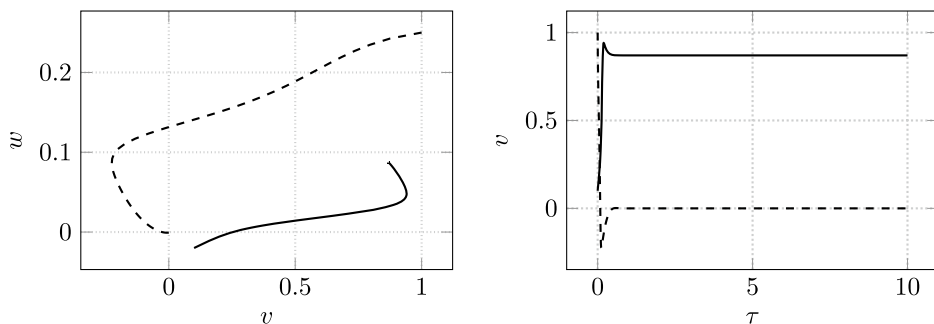


Abb. 8.35 Numerisch berechneter Schalter. Parameterwerte: $a = 0,1, b = 10, I_e = 0$. Verhältnis der Zeitskalen: $\varepsilon = 0,01$. Explizites Euler-Verfahren mit Schrittweite $\Delta\tau = 0,005$ und Anfangswerten $(v_0, w_0) = (0, 1; -0,02)$ und $(v_1, w_1) = (1; 0,25)$. Links: die Lösungstrajektorien im Phasenporträt. Rechts: die Lösung im τ - v -Diagramm. Die Iteration wurde mit dem Tabellenkalkulationsprogramm LibreOffice 4.3.4.1 durchgeführt

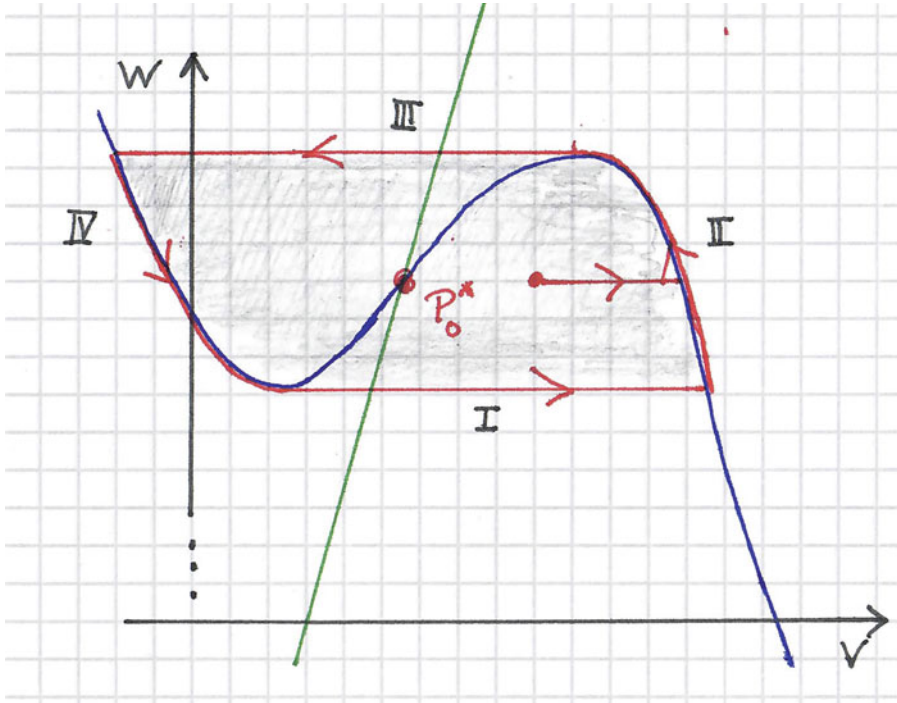


Abb. 8.36 Das FitzHugh-Nagumo-Modell mit einem instabilen stationären Punkt und einem stabilen Grenzzyklus: neuronales Feuern. Blau: v -Nullkline, grün: w -Nullkline, rot: Trajektorien

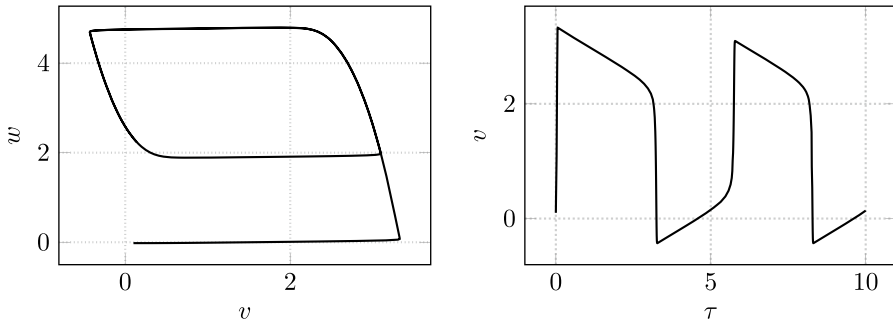


Abb. 8.37 Das periodische *Feuern* eines Aktionspotentials. Numerische Approximation zu (8.17), (8.18) mit Parameterwerten $\varepsilon = 0.05$, $a = 3$, $b = 0.4$, $I_e = 2.6$. Explizites Euler-Verfahren mit Schrittweite $\Delta\tau = 0.005$, Startwert $v_0 = 0.1$, $w_0 = -0.02$. Links: die Lösungstrajektorie im Phasenporträt. Rechts: die Lösung im τ - v -Diagramm. Die Iteration wurde mit dem Tabellenkalkulationsprogramm LibreOffice 4.3.4.1 durchgeführt

8.6 Vom Spielzeugmodell zum großen Modell: neuronales Feuern im Modell von Hodgkin und Huxley

Gibt es auch im H & H-Modell das Phänomen des neuronalen Feuerns, wenn wir einen zeitlich unabhängigen Strom I_e durch die Zellmembran zulassen?

Um dieser Frage nachzugehen, ersetzen wir die erste Gleichung im H & H-Modell (8.14) durch

$$\dot{V} = \frac{1}{C} (-\hat{g}_K n^4 (V - V_K) - \hat{g}_{Na} m^3 h (V - V_{Na}) - \hat{g}_L (V - V_L) + I_e) \quad (8.21)$$

und ersetzen dementsprechend im Euler-Verfahren (8.15) die zweite Zeile durch

$$V_{k+1} = V_k + \frac{\Delta t}{C} (-\hat{g}_{Na} m_k^3 h_k (V_k - V_{Na}) - \hat{g}_K n_k^4 (V_k - V_K) - \hat{g}_L (V_k - V_L) + I_e). \quad (8.22)$$

Wir testen numerisch unterschiedliche Werte von I_e , ausgehend von dem stationären Punkt $(V_0, h_0, m_0, n_0) = (0, h_\infty(0), m_\infty(0), n_\infty(0))$ des ungestörten H & H-Modells, und bemerken, dass für $I_e > 7$ mA dieses Phänomen tatsächlich auftritt, für Ströme mit kleineren und negativen Stromstärken jedoch nicht (siehe Abb. 8.38): Für die Störung $I_e = 5$ mA erfolgt ein einmaliges Aktionspotential, das danach zu einem neuen stationären Punkt mit einer Spannung von ca. $V = -3$ mV abklingt. Mit der Störung $I_e = 10$ mA setzt ein periodisches Feuern ein, wobei das erste Aktionspotential eine größere Amplitude hat als die folgenden. In Aufg. 8.7 sollen diese Phänomene numerisch erkundet werden, insbesondere ob und ggf. wie die Amplitude und die Frequenz des Feuerns von der angelegten Stromstärke I_e abhängen.

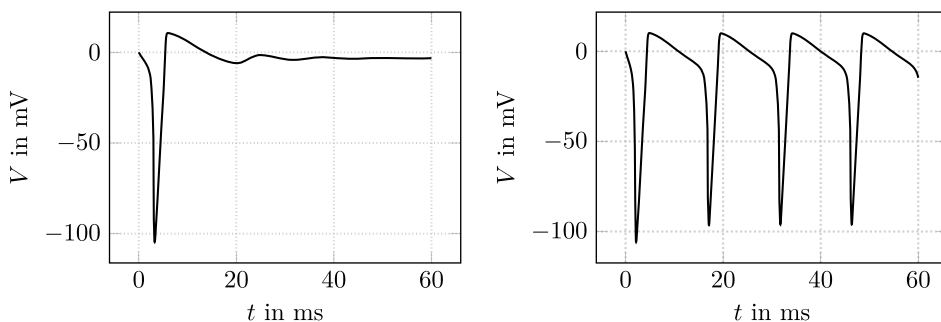


Abb. 8.38 Das H & H-Modell mit zusätzlichem Strom I_e . Parameterwerte wie in Tab. 8.2, gemischtes Euler-Verfahren mit Schrittweite $\Delta t = 0.02$ ms. Anfangswert $(V_0, h_0, m_0, n_0) = (0 \text{ mV}, h_\infty(0), m_\infty(0), n_\infty(0))$. Links: $I_e = 5$ mA. Rechts: $I_e = 10$ mA. Die Iteration wurde mit dem Tabellenkalkulationsprogramm LibreOffice 4.3.4.1 durchgeführt

An diesem Beispiel lässt sich der Nutzen von Spielzeugmodellen belegen: Durch die Einfachheit des Modells lassen sich leichter Phänomene entdecken, die man dann umgekehrt wieder in den großen Modellen suchen kann. Von der medizinischen Seite her gesehen ließe sich nun die Frage aufwerfen, ob sich bestimmte Störungen wie nervöse Ticks oder nervöse Anfallsleiden durch das Auftreten solcher zusätzlichen Ströme durch die Zellmembran von Nervenzellen physiologisch deuten lassen.

Aufgaben

8.1 Berechnung der Natrium-Gleichgewichtsspannung mit Hilfe der Nernst-Gleichung

Berechnen Sie den Wert der Natrium-Gleichgewichtsspannung U_{Na} bei einer Temperatur von 5°C , wenn die Umgebung eine 8-mal höhere Natriumkonzentration aufweist als das Axoninnere, mit Hilfe der Nernst-Gleichung (8.1). Vergleichen Sie den Wert mit den von H & H mittels anderer Methoden gefundenen Wert von $V_{\text{Na}} = -112\text{ mV}$, wobei das Ruhepotential des Axons mit $U_r = 60\text{ mV}$ angegeben ist.

Berechnen Sie das Natrium-Konzentrationsverhältnis, das zwischen dem Inneren des Axons und seiner Umgebung herrschen muss, damit bei 6°C die Natrium-Gleichgewichtsspannung bei -112 mV liegt.

8.2 Funktionsanpassung an eine Messkurve der Kaliumleitfähigkeit

Interpolieren Sie für $l = 1, 2, 3, 4$ so gut wie möglich die aus einer t - $g_K(t)$ -Messung bei einer Depolarisation von -109 mV von H und H stammenden Daten in Tab. 8.3 mit Hilfe der Modellfunktion aus (8.7). Laden Sie dazu die Messwerte in ein TK oder DGS und variieren Sie die Parameter der Modellfunktionen mit Hilfe von Schiebereglern nach Augenmaß. H und H finden für $l = 4$ die Werte $g_{K,\infty} = 20.07\text{ m S cm}^{-2}$ und $\tau_n = 1.05\text{ ms}$. Vergleichen Sie mit ihrem Ergebnis.

8.3 Schätzung der Ratenparameter für die Kaliumleitfähigkeit

In dieser Aufgabe soll die Schätzung der Ratenparameter α_n und β_n für die Kaliumleitfähigkeit des Riesenaxons nachvollzogen werden. In Tab. 8.4 sind die Messwerte aus den Experimenten von H & H bezüglich der Spannungsabhängigkeit der Kaliumleitfähigkeit angegeben.

Versuchen Sie diese Messwerte mit unterschiedlichen Modellfunktionen zu interpolieren. Gehen Sie dazu vor wie in Aufg. 8.2. Vergleichen Sie Ihre Ergebnisse mit der Interpolation von H & H, siehe (8.9).

Tab. 8.3 Messwerte einer $t - g_K(t)$ -Kurve aus [43, Abb. 3]

$\frac{t}{\text{ms}}$	0.16	0.34	0.55	0.74	1.10	1.52	2.01	2.82	4.16	6.41	8.77
$\frac{g_{K,\infty}\text{cm}^2}{\text{m S}}$	0.03	0.04	0.15	0.31	0.73	1.13	1.39	1.71	1.91	2.04	2.09

Tab. 8.4 Messwerte für die Parameter der Kaliumleitfähigkeit aus [43, Table 1]: Axon 17, Temperatur 6 °C

$\frac{V}{mV}$	$\frac{g_{K,\infty} cm^2}{mS}$	$\frac{\tau}{ms}$	n_∞	$\frac{\alpha_n}{ms}$	$\frac{\beta_n}{ms}$
-109	20.70	1.05	0.961	0.915	0.037
-100	20.00	1.10	0.953	0.866	0.043
-88	18.60	0.935	0.935	0.748	0.052
-76	17.00	1.50	0.915	0.610	0.057
-63	15.30	1.70	0.891	0.524	0.064
-51	13.27	2.05	0.859	0.419	0.069
-38	10.29	2.60	0.806	0.310	0.075
-32	8.62	3.20	0.772	0.241	0.071
-26	6.84	3.80	0.728	0.192	0.072
-19	5.00	4.50	0.674	0.150	0.072
-10	1.47	5.25	0.496	0.095	0.096
-6	0.98	5.25	0.448	0.085	0.105

8.4 Das H & H-System und die Voraussetzungen des Satzes von Picard-Lindelöf

Zeigen Sie, dass die Definitionslücken in den Ratenfunktionen α_n und α_m hebbbar sind, siehe (8.13) und (8.9), und dass die Fortsetzungen beliebig oft differenzierbar sind.

8.5 Numerik zum H & H-System

- Implementieren Sie das Euler-Verfahren (8.15) für das H & H-Modell (8.14) in einem Tabellenkalkulationsprogramm. Achten Sie dabei darauf, das Programm so zu gestalten, dass Sie Parameterwerte des Modells wie V_{Na} , Anfangswerte, und Schrittweiten bequem variieren können.
- Ermitteln Sie die Reizschwelle in Ihrer numerischen Implementierung möglichst genau.
- Überlegen Sie, wie sich die Aktionspotentiale und Reizschwellen ändern, wenn Sie die Ruhespannungen V_{Na} , V_K und V_L variieren. Überprüfen Sie Ihre Vermutung numerisch.
- Experimentieren Sie, wie weit Sie das Modell ändern können, ohne dass die Aktionspotentiale verschwinden. Sie können z. B.
 - in den Definitionen von g_{Na} und g_K andere Potenzen der Torvariablen h , m und n verwenden.
 - die Torvariablen h , m und n „versklaven“, indem diese nicht mehr eine DGL lösen, sondern in der Form $h = h_\infty(V)$ usw. dem Spannungsverlauf instantan folgen.
- Lassen Sie sich weitere Variationen einfallen und probieren Sie sie numerisch aus.

8.6 Numerik zum FitzHugh-Nagumo-System und zu seinem asymptotischen Grenzwert

Implementieren Sie das explizite Euler-Verfahren zu dem System (8.17), (8.18). Achten Sie dabei darauf, dass Sie die Werte für ε , a , b , I_e und Δt sowie die Startwerte bequem variieren können.

- a) Berechnen Sie numerisch Lösungen für unterschiedliche Werte a , b , ε , I_e und unterschiedliche Anfangsauslenkungen, so dass Sie die drei angesprochenen Phänomene *Aktionspotential*, biochemischer Schalter und neuronales Feuern beobachten können.
- b) Beschreiben Sie, welche numerischen Probleme auftreten, wenn $\varepsilon \ll 1$ gilt. Überlegen Sie, woran das liegt.
- c) Implementieren Sie für $\varepsilon \ll 1$ das System (8.17), (8.18) numerisch abwechselnd mit Hilfe des asymptotischen Limes auf der langsamen Skala (8.19) und des asymptotischen Limes auf der schnellen Skala (8.19). Für die Spannungsvariable v_0 auf der langsamen Skala müssen Sie eine kubische Gleichung numerisch lösen. Verwenden Sie dazu den Newton-Algorithmus.
- d) Vergleichen Sie die numerischen Lösungen des vollen Systems mit den numerischen Lösungen des asymptotischen Systems für unterschiedliche Werte von ε .

8.7 Neuronales Feuern im H & H-System?

Implementieren Sie das Euler-Verfahren (8.22) für das durch einen zusätzlichen Strom I_e gestörte System von H & H, (8.21), mit Hilfe eines TK.

Erkunden Sie, für welche Werte von I_e das Phänomen des neuronalen Feuerns auftritt und beschreiben Sie, wie die Amplitude und die Frequenz vom angelegten Strom I_e abhängen.

Bisher erschienene Bände der Reihe Mathematik Primarstufe und Sekundarstufe I + II

Herausgegeben von

Prof. Dr. Friedhelm Padberg, Universität Bielefeld

Prof. Dr. Andreas Büchter, Universität Duisburg-Essen

Bisher erschienene Bände (Auswahl):

Didaktik der Mathematik

- T. Bardy/P. Bardy: Mathematisch begabte Kinder und Jugendliche (P)
C. Benz/A. Peter-Koop/M. Grüßing: Frühe mathematische Bildung (P)
M. Franke/S. Reinhold: Didaktik der Geometrie (P)
M. Franke/S. Ruwisch: Didaktik des Sachrechnens in der Grundschule (P)
K. Hasemann/H. Gasteiger: Anfangsunterricht Mathematik (P)
K. Heckmann/F. Padberg: Unterrichtsentwürfe Mathematik Primarstufe, Band 1 (P)
K. Heckmann/F. Padberg: Unterrichtsentwürfe Mathematik Primarstufe, Band 2 (P)
F. Käpnick: Mathematiklernen in der Grundschule (P)
G. Krauthausen: Digitale Medien im Mathematikunterricht der Grundschule (P)
G. Krauthausen: Einführung in die Mathematikdidaktik (P)
G. Krummheuer/M. Fetzner: Der Alltag im Mathematikunterricht (P)
F. Padberg/C. Benz: Didaktik der Arithmetik (P)
E. Rathgeb-Schnierer/C. Rechtsteiner: Rechnen lernen und Flexibilität entwickeln (P)
P. Scherer/E. Moser Opitz: Fördern im Mathematikunterricht der Primarstufe (P)
H.-D. Sill/G. Kurtzmann: Didaktik der Stochastik in der Primarstufe (P)
A.-S. Steinweg: Algebra in der Grundschule (P)
G. Hinrichs: Modellierung im Mathematikunterricht (P/S)
A. Pallack: Digitale Medien im Mathematikunterricht der Sekundarstufen I + II (P/S)
R. Danckwerts/D. Vogel: Analysis verständlich unterrichten (S)
C. Geldermann/F. Padberg/U. Sprekelmeyer: Unterrichtsentwürfe Mathematik Sekundarstufe II (S)
G. Greefrath: Didaktik des Sachrechnens in der Sekundarstufe (S)

- G. Greefrath: Anwendungen und Modellieren im Mathematikunterricht (S)
G. Greefrath/R. Oldenburg/H.-S. Siller/V. Ulm/H.-G. Weigand: Didaktik der Analysis für die Sekundarstufe II (S)
K. Heckmann/F. Padberg: Unterrichtsentwürfe Mathematik Sekundarstufe I (S)
K. Krüger/H.-D. Sill/C. Sikora: Didaktik der Stochastik in der Sekundarstufe (S)
F. Padberg/S. Wartha: Didaktik der Bruchrechnung (S)
V. Ulm/M. Zehnder, Mathematische Begabung in der Sekundarstufe (S)
H.-J. Vollrath/H.-G. Weigand: Algebra in der Sekundarstufe (S)
H.-J. Vollrath/J. Roth: Grundlagen des Mathematikunterrichts in der Sekundarstufe (S)
H.-G. Weigand/T. Weth: Computer im Mathematikunterricht (S)
H.-G. Weigand et al.: Didaktik der Geometrie für die Sekundarstufe I (S)

Mathematik

- M. Helmerich/K. Lengnink: Einführung Mathematik Primarstufe – Geometrie (P)
A. Büchter/F. Padberg: Arithmetik und Zahlentheorie (P/S)
A. Büchter/F. Padberg: Einführung in die Arithmetik (P/S)
K. Appell/J. Appell: Mengen – Zahlen – Zahlbereiche (P/S)
A. Filler: Elementare Lineare Algebra (P/S)
H. Humenberger/B. Schuppar: Mit Funktionen Zusammenhänge und Veränderungen beschreiben (P/S)
S. Krauter/C. Bescherer: Erlebnis Elementargeometrie (P/S)
H. Kütting/M. Sauer: Elementare Stochastik (P/S)
T. Leuders: Erlebnis Algebra (P/S)
T. Leuders: Erlebnis Arithmetik (P/S)
F. Padberg/A. Büchter: Elementare Zahlentheorie (P/S)
F. Padberg/R. Danckwerts/M. Stein: Zahlbereiche (P/S)
H. Albrecht: Elementare Koordinatengeometrie(S)
A. Büchter/H.-W. Henn: Elementare Analysis (S)
B. Schuppar: Geometrie auf der Kugel – Alltägliche Phänomene rund um Erde und Himmel (S)
B. Schuppar/H. Humenberger: Elementare Numerik für die Sekundarstufe (S)
G. Wittmann: Elementare Funktionen und ihre Anwendungen (S)
S. Bauer, Mathematisches Modellieren (S)

P: Schwerpunkt Primarstufe

S: Schwerpunkt Sekundarstufe

Literatur

1. Adelmeyer, M.: Theorem von Sarkovskii. Berichte über Mathematik und Unterricht 89–04. ETH, Zürich (1990)
2. Arimoto, S., Nagumo, J., Yoshizawa, S.: An active pulse transmission line simulating nerve axon. Proc. IRE **50**, 2061–2070 (1962)
3. Arnol'd, V.I.: Gewöhnliche Differentialgleichungen. Springer, Berlin (1980). Translated from the Russian by Brigitte Mai
4. Arnol'd, V.I.: Catastrophe Theory, third Aufl. Springer, Berlin (1992). Translated from the Russian by G. S. Wassermann, Based on a translation by R. K. Thomas
5. Bauer, S.: Modellieren mit Differentialgleichungen. In Beiträge zum Mathematikunterricht. S. 59–62 (2017)
6. Bauer, S.: Modellieren mit Differentialgleichungen über das logistische Modell hinaus. Der Mathematikunterricht **63**(1), 28–36 (2017)
7. Bendixson, I.: Sur les courbes définies par des équations différentielles. Acta. Math. **24**(1), 1–88 (1901)
8. Bernstein, J.: Untersuchungen zur Thermodynamik der bioelektrischen Ströme. Erster Teil. Pflügers Arch. **92**, 521–562 (1902)
9. Bevington, P.R., Robinson, D.K.: Data Reduction and Error Analysis for the Physical Sciences, 3. Aufl. McGraw-Hill, New York (2003)
10. Blum, W., Löding, W.: Diskussion: Die Hamburger Abituraufgaben im Fach Mathematik. Mitt. Dtsch. Math.-Ver. **22**(3), 134–135 (2014)
11. Böge, A., Böge, W.: Technische Mechanik. Springer Vieweg, Wiesbaden (2015)
12. Bradley, D.M.: Verhulst's logistic curve. Coll. Math. J. **32**, 94–98 (2001)
13. Brandes, R., Lang, F., Schmidt, R.F. (Hrsg.): Physiologie des Menschen mit Pathophysiologie, 32. Aufl. Springer, Berlin (2019). (Springer Lehrbuch)
14. Britton, N.F.: Essential mathematical biology. Springer Undergraduate Mathematics Series. Springer, London (2003)
15. Büchter, A., Henn, H.W.: Handbuch der Mathematikdidaktik, chapter Schulmathematik und Realität. Springer Spektrum, Heidelberg (2015)
16. Natural Resources Canada. Spruce budworm. <https://www.nrcan.gc.ca/our-natural-resources/forests-forestry/wildland-fires-insects-disturban/top-forest-insects-diseases-cana/spruce-budworm/13383>. Zugegriffen: 10. März 2020

17. Carpenter, S.R., Ludwig, D., Brock, W.A.: Managment of eutrophication for lakes subject to potentially irreversible change. *Ecol. Appl.* **3**, 751–771 (1999)
18. Cole, K.S.: Dynamic electrical characteristics of the squid axon membran. *Arch. Sci. physiol.* **3**, 253–258 (1949)
19. Deuffhard, P., Bornemann, F.: Numerische Mathematik 2, 3. Aufl. de Gruyter, Berlin (2008)
20. Doktor, J.: Über die Authentizität von Modellierungsaufgaben zu zeitabhängigen Größen im Zentralabitur Mathematik. Master's thesis, Universität Duisburg-Essen, Deutschland (2018)
21. Eck, C., Garcke, H., Knabner, P.: Mathematische Modellierung, 2. Aufl. Springer, Berlin (2011)
22. Eilers, J.: Nervenzellen. In: Schmidt, R.F. (Hrsg.) *Physiologie des Menschen*, Springer-Lehrbuch. Springer, Berlin (2019)
23. Elton, C., Nicholson, M.: The ten-year cycle in numbers of the lynx in canada. *J. Anim. Ecol.* **11**(2), 215–244 (1942)
24. ADAC e. V.: Berechnung des Anhalteweges. https://www.adac.de/_mmm/pdf/verkehr_und_mathe_anhalteweg_45164.pdf. Zugegriffen: 20. März 2020
25. Fakler, B., Eilers, J.: Ruhemembranpotenzial und Aktionspotenzial. In: Brandes, R., Lang, F., Schmidt, R.F. (Hrsg.) *Physiologie des Menschen*. Springer-Lehrbuch. Springer, Berlin (2019)
26. Feigenbaum, M.J.: The universal metric properties of nonlinear transformations. *J. Statist. Phys.* **21**, 669–706 (1979)
27. Borromeo Ferri, R., Greefrath, G., Kaiser, G. (Hrsg.): *Mathematisches Modellieren für Schule und Hochschule. Theoretische und didaktische Hintergründe*. Springer Spektrum, Wiesbaden (2013)
28. Fischer, G.: *Lineare Algebra*. Rowohlt Taschenbuch Verlag GmbH & Verlag Vieweg, Reinbek bei Hamburg (1975). Unter Mitarbeit von Richard Schimpl, *Mathematik Grundkurs*, 17
29. FitzHugh, R.: Impulses and physiological states in theoretical models of nerve membrane. *Biophys. J.* **1**, 445–466 (1961)
30. Forster, O.: *Analysis 2*. Vieweg+Teubner Verlag, Wiesbaden (2011)
31. Fowkes, N.D., Mahony, J.J.: *Einführung in die Mathematische Modellierung*. Spektrum Akademischer Verlag GmbH, Heidelberg (1996). Translated from the 1994 English original by Ulrich Mitreuter
32. Fowler, A.C.: *Mathematical models in the applied sciences*. Cambridge Texts in Applied Mathematics. Cambridge University Press, Cambridge (1997)
33. Galileo, G.: *Discorsi e dimostrazioni matematiche, intorno a due nove scienze*. Deutsch, Thun, Frankfurt (1638, 1995). Übersetzung von A. von Oettingen, Nachdruck
34. Gilpin, M.E.: Do hares eat lynx? *Amer. Nat.* **107**, 727–730 (1973)
35. Grobman, D.M.: Homeomorphism of systems of differential equations. *Dokl. Akad. Nauk SSSR* **128**, 880–881 (1959)
36. Großmann, S., Thomae, S.: Invariant distributions and stationary correlation functions of one-dimensional discrete processes. *Z. Naturforsch. A* **32**, 1353–1363 (1977)
37. Haas, N., Morath, H.J.: *Anwendungsorientierte Aufgaben für die Sekundarstufe II. Mathematik*. Schroedel, Braunschweig (2005)
38. Hartman, P.: A lemma in the theory of structural stability of differential equations. *Proc. Amer. Math. Soc.* **11**, 610–620 (1960)
39. Hertz, H.: Die Prinzipien der Mechanik in neuem Zusammenhange dargestellt. J.A. Barth, Leipzig (1894)
40. Hodgkin, A.L., Huxley, A.F.: The components of membrane conductance in the giant axon of loligo. *J. Physiol.* **116**, 473–496 (1952)
41. Hodgkin, A.L., Huxley, A.F.: Currents carried by sodium and potassium ions through the membrane of the giant axon of Lologo. *J. Physiol.* **116**, 449–472 (1952)

42. Hodgkin, A.L., Huxley, A.F.: The dual effect of membrane potential on sodium conductance in the giant axon of loligo. *J. Physiol.* **116**, 497–506 (1952)
43. Hodgkin, A.L., Huxley, A.F.: A quantitative description of membrane current and its application to conduction and excitation in nerve. *J. Physiol.* **117**, 500–544 (1952)
44. Hodgkin, A.L., Huxley, A.F., Katz, B.: Ionic currents underlying activity in the giant axon of the squid. *Arch. Sci. physiol.* **3**, 129–150 (1949)
45. Hodgkin, A.L., Huxley, A.F., Katz, B.: Measurement of current-voltage relations in the membrane of the giant axon of loligo. *J. Physiol. (Lond.)* **116**, 424–484 (1952)
46. Holling, C.S.: The components of predation as revealed by a study of small mammal predation on the European pine sawfly. *Can. Entomol.* **91**, 293–320 (1959)
47. Huckenbeck, W.: *Arzt und Alkohol, Forensische Alkoholologie (1)*. Institut für Rechtsmedizin der Heinrich-Heine-Universität, Düsseldorf (1999)
48. Hussell, M.P., May, R.M., Lawton, J.: Pattern of dynamic behaviour in single species populations. *J. Anim. Ecol.* **45**, 471–486 (1976)
49. Jahnke, T., Klein, H.P., Kühnel, W., Sonar, T., Spindler, M.: Die Hamburger Abituraufgaben im Fach Mathematik. Entwicklung von 2005 bis 2013. *Mitt. Dtsch. Math.-Ver.* **22**(2), 115–122 (2014)
50. Jones, D.D.: The budworm site model. In: Norton, C.A., Holling, C.S. (Hrsg.) *Pest managment. proceeding of an international conference*, S. 91–155. Pergamon Press, Oxford (1979)
51. Kaiser, G., Blum, W., Borromeo Ferri, R., Greefrath, G.: *Handbuch der Mathematikdidaktik, chapter Anwendungen und Modellieren*. Springer Spektrum, Berlin (2015)
52. Kelley, W.G., Peterson, A.C.: *Difference Equations*, 2. Aufl. Academic Press, San Diego (2001)
53. Klein, F.: *Elementarmathematik vom höheren Standpunkte aus. Arithmetik, Algebra, Analysis*. Vierte Auflage, Erster Bd. Springer, Berlin (1968)
54. Kochanek, M., Michels, G.: *Repetitorium Internistische Intensivmedizin*, 3. Aufl. Springer, Berlin (2017)
55. Krabs, W.: *Mathematische Modellierung*. Teubner, Stuttgart (1997)
56. Kühnel, W.: *Differentialgeometrie*. Vieweg, Braunschweig (1999)
57. Kühnel, W.: Modellierungskompetenz und Problemlösekompetenz im Hamburger Zentralabitur zur Mathematik. *Math. Semesterber.* **62**(1), 69–82 (2015)
58. Levin, S.A. (Hrsg.): *Frontiers in mathematical biology. Lecture Notes in Biomathematics*, Bd. 100. Springer, Berlin (1994)
59. Lotka, A.J.: *Elements of Physical Biology*. William and Wilkins Company, Philadelphia (1925)
60. Ludwig, D., Jones, D.D., Holling, C.S.: Qualitative analysis of insect outbreak systems: The spruce budworm and forest. *J. Anim. Ecol.* **47**(1), 315–323 (1978)
61. Malthus, T.R.: *An Essay on the Principle of Population*. Penguin Books, London (1970). Erstveröffentlichung im Jahr 1798
62. Marmont, G.: Studies on the axon membran. *J. cell. comp. Physiol.* **34**, 351–382 (1949)
63. Michaelis, L., Menten, M.I.: Die Kinetik der Invertinwirkung. *Biochem. Z.* **49**, 333–369 (1913)
64. Morris, R.F. (Hrsg.): *The dynamics of epidemic Spruce budworm populations*, Bd. 21. *Memoirs of the Entomological Society of Canada* (1963)
65. Murray, J.D.: *Mathematical biology. I*, Bd. 17 of *Interdisciplinary Applied Mathematics*, third Aufl. Springer, New York (2002). An introduction
66. Nagel, K., Schreckenberg, M.: A cellular automaton model for freeway traffic. *J. Phys. I France* **2**, 2221–2229 (1992)
67. Nicholson, A.J.: An outline of the behaviour of the dynamics of animal populations. *Australian J. Zool.* **2**, 9–65 (1954)
68. Nisbet, N., Gurney, W.S.C.: *Modelling Fluctuating Populations*. Wiley, New York (1982)

69. Ortlieb, C.P.: Heinrich Hertz und das Konzept des Mathematischen Modells. In: Wolfschmidt, G. (Hrsg.) *Heinrich Hertz (1857–1894) and the Development of Communication - Proceedings of the International Scientific Symposium in Hamburg, 2007*. Books on Demand, Norderstedt (2008)
70. Ortlieb, C.P., von Dresky, C., Gasser, I., Günzel, S.: *Mathematische Modellierung*. Springer Spektrum, Wiesbaden (2009)
71. Perko, L.: *Differential equations and dynamical systems*, Bd. 7 of Texts in Applied Mathematics, third Aufl. Springer, New York (2001)
72. Poincaré, H.: Sur les courbes définies par une équation différentielle. *Compt. Rend.* **90**, 673–675 (1882)
73. Poland, H.: Fur-bearing animals in nature and commerce. *Nature* **46**, 605–606 (1892)
74. Rettinger, J., Schwarz, S., Schwarz, W.: *Methodische Grundlagen*. In *Elektrophysiologie*. Springer Spektrum, Berlin (2018a)
75. Rettinger, J., Schwarz, S., Schwarz, W.: *Theorie der Erregbarkeit*. In *Elektrophysiologie*. Springer Spektrum, Berlin (2018b)
76. Richter, T.: *Planung von Autobahn und Landstraße*. Springer Vieweg, Wiesbaden (2016)
77. Rothe, F.: The periods of the Volterra-Lotka system. *J. Reine Angew. Math.* **355**, 129–138 (1985)
78. Sandefur, J.T.: *Discrete Dynamical Systems. Theory and Applications*. Oxford University Press, London (1990)
79. Scheid, H., Schwarz, W.: *Elemente der Arithmetik und Algebra*, 6. Aufl. Springer Spektrum, Berlin (2016)
80. Schuppar, B., Humenberger, H.: *Elementare Numerik für die Sekundarstufe*. Springer Spektrum, Berlin (2015)
81. Segel, L.A., Handelman, G.H.: *Mathematics Applied to Continuum Mechanics*. Macmillan, New York (1977)
82. Siller, H.S.: Blockabfertigung im (Tauern)-Tunnel. In: Borromeo Ferri, R., Greefrath, G., Kaiser, G. (Hrsg.) *Mathematisches Modellieren für Schule und Hochschule*. Springer Spektrum, Wiesbaden (2013)
83. Singer, M.V.: *Alkohol und Alkoholfolgekrankheiten : Grundlagen – Diagnostik – Therapie*, 2., vollständig überarbeitete und aktualisierte Aufl. Springer Medizin Verlag, Heidelberg (2005)
84. Sonar, T.: *Angewandte Mathematik, Modellbildung und Informatik*. Vieweg, Braunschweig (2001)
85. Sonar, T.: *3000 Jahre Analysis. Vom Zählstein zum Computer. [From Counting Pebbles to the Computer]*. Springer, Heidelberg (2011). *Geschichte, Kulturen, Menschen. [History, cultures, people]*, With a foreword by Heinz-Wilhelm Alten
86. Sussmann, H.J., Zahler, R.S.: Catastrophe theory as applied to the social and biological sciences: a critique. *Synthese* **37**(2), 117–216 (1978). (Mathematical methods in the social sciences, III)
87. Teschl, G.: *Ordinary differential equations and dynamical systems*. Graduate Studies in Mathematics, Bd. 140. American Mathematical Society, Providence, RI (2012)
88. Verhulst, P.: Notice sur la loi que la population poursuit dans son accroissement. *Corresp. Math. Phys* **10**, 113–121 (1838)
89. Volterra, V.: Variazioni e fluttuazioni del numero d'individui in specie animali conviventi. *Mem. R. Accad. Naz. dei Lincei. Ser. VI* **2**, 31–113 (1926)
90. Šarkovskii, O.M.: Co-existence of cycles of a continuous mapping of the line into itself. *Ukrain. Mat. Ž.* **16**, 61–71 (1964)
91. Waldvogel, J.: The period in the Lotka-Volterra system is monotonic. *J. Math. Anal. Appl.* **114**(1), 178–184 (1986)
92. Walter, W.: *Gewöhnliche Differentialgleichungen: eine Einführung*, 7. Aufl. Springer, Berlin (2000)

-
93. Wigner, E.: The unreasonable effectiveness of mathematics in the natural sciences. *Comm. Pure Appl. Math.* **13**, 1–14 (1960)
 94. Wu, N.: Verkehr auf Schnellstraßen im Fundamentaldiagramm - Ein neues Modell und seine Anwendungen. *Straßenverkehrstechnik* **44**(8), 378–403 (2000)
 95. Zeeman, E.C.: Levels of structure in catastrophe theory illustrated by applications in the social and biological sciences. In *Proceedings of the International Congress of Mathematicians* (Vancouver, B. C., 1974), Bd. 2, S. 533–546 (1975)

Stichwortverzeichnis

A

- Abbau, [148](#), [149](#)
- Abbauprodukt, [149](#)
- Abbaurrate, [149](#)
 - maximale, [150](#)
- Abfluss, [157](#)
- Abies balsamea, [124](#)
- Ableitung
 - innere, [29](#)
 - partielle, [27](#)
- Ableitungsmatrix, *siehe* Jacobi-Matrix
- Ableitungsvektor, dimensionsloser, [18](#)
- Absammelstrategie, [108](#)
- Absorption, [157](#), [160](#)
- Absorptionsrate, [160](#)
- Abwurfgeschwindigkeit, [225](#)
- Ad-hoc-Ansatz, [303](#)
- Ad-hoc-Formulierung, [8](#)
- Ähnlichkeitsansatz, [141](#)
- Aktionspotential, [271](#), [320](#)
- Alkoholhydrogenase, [149](#)
- Alkoholkonzentration, [146](#), [164](#)
- Alkoholmoleküle, [148](#)
- Amplitude, [27](#), [174](#), [317](#)
- Anfangsdatum, *siehe* Anfangswert
- Anfangsgeschwindigkeit, [228](#)
- Anfangswert, [68](#), [102](#)
- Anfangswertproblem, [111](#)
 - Existenz und Eindeutigkeit der Lösungen, [115](#)
- Anhalteweg, [40](#), [57](#)
- Annahme, [2](#)
- Ansatz, [188](#)
- Antrieb, elektrochemischer, [274](#), [277](#), [298](#)
- Approximation
 - der Bogenlänge, [31](#)
 - der Ordnung k , [233](#)
 - k -ter Ordnung, [230](#)
 - lineare, [16](#), [144](#)
 - nullter Ordnung, [230](#), [312](#)
 - numerische, [33](#)
- Äquivalenzklasse, [17](#)
 - parametrisierter Kurven, [17](#)
- Areakosinushyperbolicus, *siehe* Funktion, hyperbolische
- Areasinushyperbolicus, *siehe* Funktion, hyperbolische
- Atom, [235](#)
- Aufprallgeschwindigkeit, [161](#)
- Aufpunkt, [16](#)
- Aufstrich, [273](#)
- Aufwand, konstanter, [119](#)
- Ausgleichsgerade, [287](#)
- Auslenkung, [26](#)
- Auslenkungswinkel, [216](#)
- Außenpotential, *siehe* Potential
- Auszahlung, [66](#)
- Auto, [40](#)
- Autoverkehr, [39](#)
- AWP, *siehe* Anfangswertproblem
- Axon, [272](#)
- Axonmembran, [275](#)

B

Bakterien, 161
Banach'scher Fixpunktsatz, 54
Basis, 187
Bejagungsdruck, 207
Bejagungsschwelle pro Zweig, 248
Benefit, 204
Benzol-L-Argine-Ethyl-Ester, 238
Bernoulli'sche Differentialgleichung, 143
Berühren
 in erster Ordnung, 34
 in zweiter Ordnung, 34
Berührkreis, 34
Beschleunigung, 27
Beschleunigungskonstante, 57
Beute, 167
Bewegung, gleichmäßig beschleunigte, 41
Beweis, 233
Bifurkationsdiagramm, 91
Bifurkationsparameter, 91
Binnensee, 156
Blow-up, 137
Blutalkoholkonzentration, *siehe* Alkoholkonzentration
Blutkreislauf, 148
Bogenlänge
 minimale, 61
Bogenlängenfunktional, 61
Bogenlänge, 7, 18, 31
 Parametrisierung nach, 33
Bogenlängenparameter, 18
Brachistochrone, 61
Bremsbeschleunigung, 57
Bremsverzögerung, 41
Bremsweg, 40, 57
Bündel, 178

C

Chaos, 84
 deterministisches, 84
Cholin, 283
Choristoneura fumiferana, 124
Cobwebbing, 68
Computer-Algebra-System, 90
Coulomb, 279
Cusp-Katastrophe, 131

D

Definitionsbereich, 45
Definitionslücke, 319
Dendrit, 272
Depolarisation, 277
Determinante, 188
DGL, *siehe* Differentialgleichung
DGLS, *siehe* Differentialgleichungssystem
Dichte, 28
Differentialgleichung, 111
 autonome, 132
 Bernoulli'sche, 162
 lineare homogene, 141
 lineare inhomogene, 141
 logistische, 112, 162
 nichtlineare, 300
 partielle, 27, 185
 zweiter Ordnung, 48
Differentialgleichungssystem, 150, 169
 lineares, 180
 zeitautonomes, 216
Differentialrechnung, 45
 mehrdimensionale, 45
Differenzengleichung
 mit konstanten Koeffizienten, 68
 autonome, 68
 lineare, 65
 logistische, 82, 105
Differenzenquotient, 28, 45
Diffusion, 150, 160
Diffusionskonstante, 152
Diffusionsrate, 160
Dimension
 dimensionslose Variable, 5
 Geschwindigkeit, 16
 Länge, 4
Dimensionsanalyse, 26
Diskriminante, 197
Drehmatrix, 25
Drehung, 22
Dreieck, 24
Druck, 37
Durchlässigkeit, *siehe* Leitfähigkeit

E

Eigenraum, 191
Eigenvektor, 188
Eigenwert, 188

- doppelter, 188
- Einlage, 66
- Einlagerung, 159
- Einspannbedingung, 44
- Einzugsbereich, 102, 103
- Elastizitätsmodul, Young'scher, 28
- Elektrode, 275
- Elektrophysiologie, 271
- Ellipse, 193
- Emigrationsrate, 110
- Energie, 44, 217, 225
 - kinetische, 63
 - potentielle, 44, 63
- Energiefunktional, 45
- Entdimensionalisierung, 4, 5, 26, 28, 162, 184
- Entfernung, horizontale, 33
- Entwicklung, asymptotische, 230
- Enzym, 149, 237
- Epidemie, 124
- Erdbeben, 26
- Erdbeschleunigung, 44, 224
- Erdoberfläche, 228
- Erdradius, 226
- Erhaltungsgroße, 159
- Erzeugendensystem, 187
- Euler'sche Formel, 189
- Euler-Lagrange-Gleichungen, 48
- Euler-Verfahren, 144, 320
 - explizites, 144, 164, 211
 - im Mehrdimensionalen, 175
 - gemischtes, 307
 - implizites, 144
- Evidenz, numerische, 233
- Existenz und Eindeutigkeit, 132
- Existenzintervall, 137
 - maximales, 195
- Experiment, 224
- Exponentialfunktion, komplexe, 189
- F**
- Fahnenmast, 26
- Fahrspur, 40
- Faktor, integrierender, 177
- Fall, freier, 224
- Fallrinne, 224
- Fallschirmspringer, 161
- Fangquote, 118
- Fangrate
 - konstante, 118
- Fangrate, maximale, 120, 124
- Fangstrategie, 117
 - konstante Fangrate, 118
 - konstanter Aufwand, 119
- Farad, 279
- Faraday-Konstante, 279
- Federhärte, 61
- Federsee, 158
- Feedback-Mechanismus, 276
- Fehler, 154
 - absoluter, 154, 165
 - relativer, 154, 165
- Feigenbaum-Konstante, 93
- Feuern, neuronales, 272, 315, 317, 320
- First Principles, 8, 56
- Fischbestand, 117, 156
- FitzHugh-Nagumo-Modell, 310
- FitzHugh-System, 320
- Fixpunktgleichung, 69
- Fixstelle, 118
- Fluchtgeschwindigkeit, 226
- Fluss
 - durchschnittlicher, 40
 - eines dynamischen Systems, 195
 - topologischkonjugiert, 196
- Flüssigkeit, 36
- Fluxionsrechnung, 202
- Formel, explizite, 66
- Formparameter, 22, 25
- Frenet'sche Gleichungen, 20
- Frenet'sches Zweibein, 19
- Frequenz, 27, 317
- Fundamentallemma der Variationsrechnung, 47, 60
- Fundamentallösung, 188
- Fundamentalsystem, 188
- Funktion
 - dimensionslose, 29
 - hyperbolische, 7, 31
 - Kosinushyperbolicus, 8
 - lokal sublinear, *siehe* Lipschitz-Stetigkeit,
 - lokale
 - mesotrophe, 156
 - oligotrophe, 156
 - trigonometrische, 23
- Funktional, 45
- funktionaler Zusammenhang, 40

G

Galileo-System, 231
 Gaskonstante, universelle, 279
 Gebrauch des Gleichheitszeichens, 3
 Geburtenrate, 110
 Geburtskonstante, 82
 Gegenbeispiel zur eindeutigen Lösbarkeit, 133
 Generationenfolge, nichtueberlappende, 93
 Genprodukt, 158
 GeoGebra, 33, 165
 Geometrie
 dynamische Software, 318
 Geometrie, analytische, 16
 Gerade, 35
 Geschwindigkeit, 39, 224
 optimale, 39, 41, 57
 Geschwindigkeitsskala, 228
 Geschwindigkeitsvektor, 16, 36
 Gesetz, Hooke'sches, 61
 Gewichtsanteil, 148
 Gewinnfunktion, 58
 Glaskapillare, 274
 Gleichgewichtsspannung, 286, 299
 Kalium, 287
 Natrium, 287, 318
 Gleichheitszeichen, 3
 Gleichung
 algebraische, 239, 309
 implizite, 144
 quadratische, 90
 vierter Ordnung, 90
 Gleichungssystem, 33
 lineares, 187
 Gompertz-Modell, 106
 Gradient, 60
 Gravitation, Galilei'sche, 227
 Gravitationsgesetz, Newton'sches, 225
 Gravitationskonstante, 37, 226
 Gravitationskraft, 225
 Grenzwert, asymptotischer, 231, 272, 309, 320
 Grenzyklus, 212
 stabiler, 212, 316
 Größenbereich, 45
 Gültigkeitsbereich, 227
 Guthaben, 65

H

H & H-System, 320

h-Mechanismus, 293, 304
 Halber-Tacho-Regel, 40
 Halbtrajektorie, 315
 Halbwertszeit, 66
 Hassell-Modell
 mit exakter Kompensation, 94
 mit Überkompensation, 97
 mit Unterkompensation, 97
 Häufungspunkt, 84
 Herstellungskosten, 58
 Heuristik, 309
 Höchstgeschwindigkeit, 35
 Höhenlinie, 60
 Höhenunterschied, vertikaler, 33
 Hohlzylinder, 27, 28, 31
 Homöomorphismus, 196
 Hooke'sches Gesetz, 61
 Hudson's Bay Company, 182
 Hydrolyse, 238
 Hyperbel, 32, 95
 Hyperpolarisation, 273, 277, 313
 hypertroph, 156
 Hysterese, 103, 124

I

Imaginärteil, 192
 Immigrationsrate, 110
 Immobilienkredit, 66
 Inaktivierung, 297
 Induktion, vollständige, 66
 Initiationsphase, 273
 Innenpotential, *siehe* Potential
 Insekt
 fertiles, 100, 221
 steriles, 99, 221
 Insektenpopulation, 93
 instabiles Gleichgewicht, *siehe* instabiler
 stationärer Punkt
 Instabilität, 70, 86
 monotone, 70
 oszillierende, 70
 Instantan, 42
 Integral, erstes, 176, 185, 225
 Integration, partielle, 47
 Interaktionsmodell, 183
 Interaktionsterm, 168, 204
 Interpolation, 318
 Interpolationskurven, 295

Ion, 278
Ionenstrom, 279, 282
Iontor, 282
Ionttransport
 aktiver, 286
 passiver, 286
Isotherme, 51
Iteration, 144, 307
 für η , 33

J

Jacobi-Matrix, 180, 186, 194
Jäger, *siehe* Räuber

K

Kabel, elastisches, 61
Kaliumionen, 274
Kaliumstrom, 284
Kapazität, 114, 279
 pro Zweig, 248
 κ , 21
Kartellabsprachen, 59
Katastrophentheorie, 131
Kette, 44
Kettenregel, 29
Kippunkt, 121
Klothoide, 22, 25
Klothoidenparameter, 24
Knoten
 instabiler, 191
 stabiler, 191, 312
Koeffizienten der asymptotischen Entwicklung, 230
Körperflüssigkeit, 148
Koexistenz, 198
Koexistenzzustand, 172
Komplex, 149
Kondensator, 279
 elektrischer, 279
Konditionierungsspannung, 294
Konjugation, komplexe, 189
Konkurrenz, 113
 interspezifische, 220
 intraspezifische, 184, 198, 207, 218
Konvergenz, 33
Konzentrationsänderung, schnellste, 147
Konzentrationsverhältnis, 318

Koordinatensystem, 6, 8, 14, 23, 26, 33, 35, 224
Kosinusfunktion, 190
Kosinushyperbolicus, *siehe* Funktion, hyperbolische
 hyperbolische Funktionen, 44, 61
Kraft, 28
 elastische, 27
Kraftmoment, 216
Krebse, 161
Kreisbogen, 15, 25, 34
Krümmung, 15, 19, 25
 konstante, 21
Krümmungsradius, 35
Kurve
 geschlossene, 178
 kritische, 130, 254
 mit konstanter Krümmung, 34
 regulär parametrisierte, 33
 regulär orientierte, 17
 regulär parametrisierte, 16
Kurvenintegral, 34
 wohldefiniertes, 34
Kurventheorie, 15

L

Lachs, 107
Ladung, elektrische, 279
Ladungsträger, 279
Lageenergie, *siehe* Energie, potentielle
Lagrange'scher Multiplikator, 50, 60
Lagrange-Dichte, 46
Lagrange-Restglied, 46
Landstraße Richtlinien für die Anlage von, 24
Länge
 entdimensionalisierte, 33
 Planck-Länge, 37
Längenskala, 5, 228
Langzeitverhalten, 95, 116
Laplace-Wahrscheinlichkeit, 100
Läsungsfolge, oszillierende, 70
Lebensdauer, *siehe* Existenzintervall
Leckstrom, 297
Leitfähigkeit, 286, 298, 299
 Kaliumleitfähigkeit, 318
 selektive, 274
Lenkbewegung, 14
Lenkradeinschlag, 14
Lichtgeschwindigkeit, 37

Limes, 28
Limes, asymptotischer, *siehe* Grenzwert, asymptotischer
Linearfaktor, 90
Linearisierung, 87, 119, 180, 186
Linearkombination, 187
Linkskurve, 35
Lipschitz-Konstante, 139
Lipschitz-Stetigkeit, 138
Lipschitz-stetigkeit
 lokal, 138, 162, 170
Lösung
 algebraische, 90
logistische Differentialgleichung, 112
Logligo ferosi, *siehe* Tintenfisch
lokal Lipschitz-stetig, 163
Lösen, qualitative, 115
Lösung
 algebraisch-analytische, 33
 eindeutige, 133, 170
 globale, 133, 170, 195
 lokale, 133
 maximale, 170, 216
 numerische, 32, 239
 partikuläre, 141
 periodische, 179, 315
Lösungsschema, lineares, 141
Lotka-Volterra-Regeln, 180
Luchs, 182
Luftwiderstand, 161, 224
Lyapunov-Funktion, 176

M

m-Mechanismus, 292, 304
Magen-Darm-Trakt, 147
Markt mit zwei Anbietern, 58
Masse, 28, 150, 225
 der Erde, 226
 Planck-Masse, 37
Massendichte, 36, 45
Massenpunkt, 216
Maßzahl, 40
Material, 28
Medikamentenkonzentration, 165
Meerwasser, 283
Membran, 271
Menge, kompakte, 216
Michaelis-Konstante, 240

Michaelis-Menten-Reaktion, 237
Minimalstelle, 45
Minimalsystem, 309
Minimierer, 46
 unter Nebenbedingungen, 49
Minimierungsprinzip, 44
Mittelwert, 179
Mittelwertsatz der Differentialrechnung, 88
Modell, 2
 des exponentiellen Wachstums, 110
 lineares, 81
 logistisches, 82
 normatives, 65
Modellierung, deskriptive, 154
mol, 236
Molekül, 235

N

Nachfrage, 58
Nachfrage-Preis-Modell, 78
Nachfragefunktion, 58
Näherung, lineare, 119
Nash-Gleichgewicht, 59, 79
Natriumionen, 274
Natriumleitfähigkeit, 303
Natriumstrom, 283
Naturkonstante, 37, 226
Navier-Stokes-Gleichungen, 36
Nernst-Gleichung, 278, 299, 318
Nervenzelle, 271
Netz, neuronales, 272
Newton'sche Gleichung, 27
Newton-Algorithmus, *siehe* Newton-Verfahren
Newton-System, 231
Newton-Verfahren, 33, 144, 153, 320
Nicht-Erregbarkeit, 306
Niveaulinie, 217, 226
Niveaumenge, 176
Nobelpreis, 271
Normaleneinheitsvektor, 20
Nuklid, 66
Nullkline, 173, 238, 311

O

Oberlösung, 138
Oberlösung, 164
Öffnungsfläche, 161

ω -Limesmenge, 214

Orbit, *siehe* Trajektorie

Ortsabhängigkeit, 148

Ortsvektor, 35

Ozeanriese, 37

P

p - q -Formel, 188

Parallelschaltung, 282

Parameter, 27

dimensionsloser, 37, 165, 184

Parametertransformation, 17

Parametrisierung, 16

nach Bogenlänge, 18, 33

Parasitenpopulation, 167

Passivitätshypothese, 288

Pendel, mathematisches, 216

percapita reproduction rate, *siehe* Pro-Kopf-Reproduktionsrate

Periode, 27, 179, 193

Pharmakokinetik Grundannahme der, 148

Phasenporträt, 171, 185, 225, 238, 272

Phosphatgehalt, 156

Planck'sches Wirkungsquantum reduziertes, 37

Planck-Größen, 37

Planck-Länge, 37

Planck-Masse, 37

Planck-Zeit, 37

Planung, 33

Polarisation, 274

Polarkoordinaten, 35

Polonium 218, 67

Polynom, charakteristisches, 188

Polynomdivision, 90

Population, 81

Populationsgröße, 81

mittlere, 179

Potential, 275

chemisches, 278

elektrisches, 278

elektrochemisches, 277

Potentialgefälle, 278

potentielle Energie, 63

Potenzreihe, 23

Preisbindung, 59

Prinzip der minimalen Wirkung, 63

Pro-Kopf-Geburtenrate maximale, 207

Pro-Kopf-Reproduktionsrate, 82, 94, 168

Pro-Kopf-Sterberate, 207

Produkt, 237

Prozess

erster Ordnung, 235

zweiter Ordnung, 235

Punkt

hyperbolischer, 194

kritischer, 61

periodischer, 89, 106

Punkt, stationärer, 70, 86, 185, 312

instabiler, 316

Punkt, stationärer, asymptotisch stabiler, 86

monotoner, 70, 86

oszillierender, 70, 86

Q

Quadranten, 173

Quasi-Steady-State-Hypothese, 239

Querschnitt, 28

R

Radium 226, 66

Radius, 15, 25, 35

äußerer, 27

innerer, 28

Radon 222, 66

Rand zu Rand, 170

Randbedingung, 27

Randterm, 47

Randwertproblem periodisches, 28

Räuber, 167

Räuber-Beute-System

blutrünstiges, 219

mit Sättigung, 205, 219

mit Schwelle und Sättigung, 219

Räuber-Beute-Interaktion, 168

Räuber-Beute-System

nach Lotka-Volterra, 168

Raum, affin-linearer, 50

Raubereich, 36

Reaktion

autokatalytische, 158

chemische, 235

enzymatische, 149, 237

Reaktionsgeschwindigkeit, 236

Reaktionskette, 236

Reaktionsweg, 40, 42

- Realsituation, 2
 Realteil, 192
 Recyclingrate, 157
 Reibung, 62, 224
 Reibungskraft, 217
 Reizschwelle, 306, 319
 Rekonvaleszenzzeit, 119
 Rekursion, 66
 Repolarisation, 273, 313
 Richtungsableitung, 45
 Richtungsfeld, 144
 Riesenaxon, 274
 Rückrichtung, 237
 Ruhelage, 26
 Ruhespannung, 275
 Ruhezustand, 273
- S**
- Sarkovskii-Ordnung, 92
 Sattelpunkt, 192
 Sättigung, 205, 220, 221
 Sättigungsgrenze, 121
 Sättigungswert, 150
 Satz
 - über implizite Funktionen, 53
 - von Cayleigh-Hamilton, 189
 - von der linearisierten Stabilität, 87
 - von Grobman-Hartman, 194
 - von Peano, 169
 - von Picard-Lindelöf, 169
 - von Poincaré-Bendixson, 214
 - von Poincaré-Bendixson, 315
 - von Sarkovskii, 91, 106
- Schalter
 - biochemischer, 158, 313
- Scheitelpunkt, 207
 Schieberegler, 33, 318
 Schiffsrumpf, 36
 Schneeschuhhase, 182
 Schnellstraße, 35
 Schrecksekunde, 41, 57
 Schrittweite, 211, 307
 Schwankung, zufällige, 119
 Schwelle, 315
 Schwelleneffekt, 122, 220, 221
 Schwellenwert eines Aktionspotentials, 276
 Schwerkraft, 61
 Sedimentation, 157
- See, 156
 Sekantenstücke, 31
 Separatrix, 192
 Sicherheitsabstand, 39
 Signal, 272
 Signaturmatrix, 208, 219, 312
 Simulation, 5, 305
 Sinusfunktion, 190
 Sinushyperbolicus, *siehe* Funktion, hyperbolische
 Skala, 29, 228
 Skalarprodukt, 19, 60
 Skalierung, 5
 Spannung, elektrische, 272, 275
 Spannungsmessung Orientierung der, 275
 Spannungsprofil, stufenförmiges, 279
 Spannungsquelle, 298
 Spiegelsymmetrie, 2
 Spinkurve, *siehe* Klothoide
 Spirale, logarithmische, 35
 Spur einer Matrix, 188
 Stab, 216
 Stabilität, 70, 86, 89
 Stahl, 28
 Standardklothoide, *siehe* Klothoide
 Startwert, *siehe* Anfangswert
 stationärer Zustand, 44
 Stau aus dem Nichts, 43
 Steckbriefaufgabe, 14, 33
 Sterberate, 110
 Sterbewahrscheinlichkeit, 82
 Steril Insect Release Method (SIRM), 99, 221
 Stoff, chemischer, 235
 Stoffkonzentration, 236
 Störung, 119
 Straßenverlauf, 35
 Strecke, 21, 34
 Strom
 - Ionenstrom, 282
 - Kaliumstrom, 284
 - kapazitiver, 279, 280
 - Leckstrom, 297
 - Natriumstrom, 283
- Strudel
 - instabiler, 193
 - stabiler, 193, 312
- Substitution, 143
 Substrat, 237
 Substratenzymkomplex, 237

Summenformel, geometrische, 66
Symbiose, 220
System, asymptotisches, *siehe* Grenzwert,
asymptotischer

T

Tabellenkalkulation, 33, 105, 153, 164, 165,
175, 318–320
Tabellenkalkulationsprogramm, 83, 307
Tangente, 24, 34, 35
Tangenteneinheitsvektor, 19, 24
Tangentenwinkel, 21
Tangential, 16, 51
Tannenwickler, 124
Taylor-Entwicklung, 230
mehrdimensionale, 46
Taylor-Formel, mehrdimensionale, 59
Teilchenstrom, 279
Temperatur absolute, 279
Tempo, 17
Tempolimit, 43
Testfunktion, 45
Tintenfisch, 271, 274
Torvariable, 309
langsame, 310
schnelle, 310
Toy-Model, 272
Trägerteilchen, 274
Trägheitsradius, 28
Trajektorie, 171, 238, 312
in den Achsen, 173
punktförmige, 171
Trassierung, 14, 21
Trassierungselement, 22
Trennfläche, 148
Trennung der Variablen, 132, 140, 161, 162
Tryptin, 239

U

Umkehrfunktion, 18, 135
Umlaufrichtung, 179
Unstetigkeit, 280
Unterlösung, 138
Unterlösung, 164
Upstroke, 273, 274, 313

V

Vakuum, 37
Variable, dimensionslose, 7, 29
Variation, 45
der Konstanten, 142
eines Funktionals, 45
Vektorraum, 187
Verbreitung
von Gerüchten, 155
von Krankheiten, 155
Vereinfachung, 2
Vergleichssatz, 139, 163
Verhalten, chaotisches, 84
Verhältnis dimensionsloses, 229
Verkehrsdichte, 43
Verkehrsfluss, 39, 40, 57
empirischer, 43
Verkehrsphysik, 40
Verschiebung, 16
Verschiebungsvektor, 16, 25
Verteilungsvolumen, 148
Vielfaches, lineares, 20
Vielfachheit
algebraische, 188
geometrische, 188
Viskosität, dynamische, 36
Voltage-Clamp, 274, 276
Volumen, 28
Vorkonditionierung, 296

W

Wachstum
exponentielles, 111
geometrisches, 112
logistisches, 113
Winkel, 35
Winkelbeschleunigung, 216
Wirkstoffdosierung, 165
Wirkstoffkonzentration, 165
Wirkung, 63
Wirkungsfunktional, 63
Wirtspopulation, 167

Z

Zähflüssigkeit, 37
Zeit, 40
Planck-Zeit, 37

- typische, 301
- zu minimierende, 40
- Zeitableitung, 29
- Zeitfunktional, 62
- Zeitkonstante, 309
- Zeitskala, 228, 280
- Zeittranslation, 115
- Zellkörper, 272
- Zellmembran, 272
- Zentralabitur, 165
- Zentrum, 193
- Zerfall
 - spontaner, 149, 158
- Zerfallsfaktor, 67
- Ziffer, signifikante, 3
- Zinssatz, 65
- Zuflussrate, 157
- Zugspannung, 61
- Zusammenhang
 - linearer, 40
- Zustand, 44
 - endemischer, 255
 - epidemischer, 255
 - eutropher, 156
 - hypertropher, 158
 - mesotropher, 158
 - stationärer, 44
 - zugänglicher, 44
- Zwischenwertsatz, 92, 134
- Zykel, 193



Willkommen zu den Springer Alerts

Unser Neuerscheinungs-Service für Sie:
aktuell | kostenlos | passgenau | flexibel

Mit dem Springer Alert-Service informieren wir Sie individuell und kostenlos über aktuelle Entwicklungen in Ihren Fachgebieten.

Abonnieren Sie unseren Service und erhalten Sie per E-Mail frühzeitig Meldungen zu neuen Zeitschrifteninhalten, bevorstehenden Buchveröffentlichungen und speziellen Angeboten.

Sie können Ihr Springer Alerts-Profil individuell an Ihre Bedürfnisse anpassen. Wählen Sie aus über 500 Fachgebieten Ihre Interessensgebiete aus.

Bleiben Sie informiert mit den Springer Alerts.

Jetzt
anmelden!

Mehr Infos unter: springer.com/alert

Part of **SPRINGER NATURE**