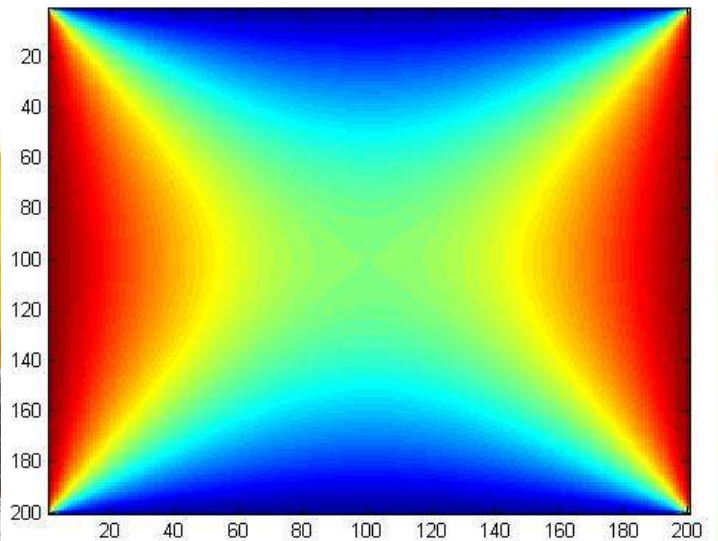


Woche der Modellierung mit Mathematik

Dokumentationsbroschüre
10.02.-16.02.2013



WOCHE DER MODELLIERUNG MIT MATHEMATIK



PÖLLAU BEI HARTBERG, 10.02.-16.02.2013

WEITERE INFORMATIONEN:

[HTTP://MATH.UNI-GRAZ.AT/MODELLWOCHE/2013/](http://math.uni-graz.at/modellwoche/2013/)

ORGANISATOREN UND SPONSOREN



Comfortplan
Online - Versicherungsmakler



Vorwort

Viele Wissenschaften erleben zurzeit einen ungeheuren Schub der Mathematisierung. Mathematische Modelle, die vor wenigen Jahrzehnten noch rein akademischen Wert hatten, können heute mit Hilfe von Computern vollständig durchgerechnet werden und liefern praktische Vorhersagen, die helfen, Phänomene zu verstehen, Vorgänge zu planen, Kosten einzusparen. Damit unsere Gesellschaft auch in Zukunft mit der technologischen Entwicklung schritthält, ist es wichtig, bereits junge Leute für diese Art mathematischen Denkens zu begeistern und in der Gesellschaft das Bewusstsein für den Nutzen angewandter Mathematik zu heben. Dies war für uns einer der Gründe, die Woche der Modellierung mit Mathematik zu veranstalten.

Nun ist leider für viele Menschen Mathematik ein Schulfach, mit dem sie eher unangenehme Erinnerungen verbinden. Umso erstaunlicher erscheint es, dass Schülerinnen und Schüler sich freiwillig melden, um eine ganze Woche lang mathematische Probleme zu wälzen - und dabei auch noch Spaß haben. Sie erleben hier offensichtlich die Mathematik auf eine Art und Weise, wie sie der Schulunterricht nicht vermitteln kann. Die jungen Leute arbeiten und forschen in kleinen Gruppen mit Wissenschaftler/innen an realen Problemen aus den verschiedensten Bereichen und versuchen, mit Hilfe mathematischer Modelle neue Erkenntnisse zu gewinnen. Sie arbeiten ohne Leistungsdruck, dafür mit Eifer und Enthusiasmus, rechnen, diskutieren, recherchieren, oft auch noch am späten Abend, in einer entspannten und kreativen Umgebung, die den Schüler/innen und betreuenden Wissenschaftler/innen gleichermaßen Spaß macht. Die Projektbetreuer konnten auch in diesem Jahr wieder erleben, wie eigenes Entdecken und Selbstmotivation das Verhalten der Schüler/innen während der ganzen Modellierungswoche bestimmen. Sie lernen eine Arbeitsmethode kennen, die in beinahe allen Details den Arbeitsmethoden einer Forschergruppe entspricht. Bei keiner anderen Gelegenheit erfahren Schüler/innen so viel über Forschung wie bei so einer Veranstaltung.

Modellierungswochen gab bzw. gibt es zum Beispiel auch in den USA, in Deutschland oder in Italien. Wir verdanken Herrn Prof. Dr. Stephen Keeling den Vorschlag, auch durch die Universität Graz so eine Woche zu veranstalten, und seiner unermüdlichen Organisationsarbeit das tatsächliche Zustandekommen. Er leitet nun bereits zum neunten Mal diese inzwischen zur Institution gewordene Veranstaltung. Ihm sei an dieser Stelle noch einmal ausdrücklich und herzlich gedankt. Besonders wichtig war in den vergangenen Jahren auch die Unterstützung durch den langjährigen Mentor der Modellierungswoche, Herrn o.Univ.-Prof. Dr. Franz Kappel, der oft auch eine eigene Gruppe mit interessanten Problemstellungen betreut hat.

Wir danken dem Landesschulrat für Steiermark, und hier insbesondere Frau Landesschulinspektorin Frau HR Mag. Marlies Liebscher, für die Hilfe bei der Organisation und ihre kontinuierliche Unterstützung der Idee einer Modellierungswoche. Ohne den idealistischen, unentgeltlichen und engagierten Einsatz der direkten Projektbetreuer Dr. Birgit Bednar-Friedl, Dipl.-Math. Carl Philip Trautmann, Daniel Kraft, BSc BSc MSc und Mag. Martin Holler – Institut für Mathematik und Wissenschaftliches Rechnen bzw. Institut für Volkswirtschaftslehre – hätte diese Modellierungswoche nicht stattfinden können.

Besonderer Dank gebührt ferner Herrn Mag. Dr. Christoph Gruber, M.A., der die ganze Veranstaltung betreut und auch die Gestaltung dieses Berichtes übernommen hat, Frau Michaela Seiwald für die tatkräftige Hilfe bei der organisatorischen Vorbereitung, und Herrn Alexander Sekkas für die Hilfe bei der Betreuung der Hard- und Software.

Finanzielle Unterstützung erhielten wir vom Land Steiermark durch Landesrätin Mag. Kristina Edlinger-Ploder, von der Karl-Franzens-Universität Graz durch Vizerektor Prof. Dr. Martin Polaschek und Dekan Prof. Dr. Karl Crailsheim, von Magna Powertrain und Comfortplan.

Pöllau, am 16. Februar 2013

Bernd Thaller
Institut für Mathematik und Wissenschaftliches Rechnen
Karl-Franzens-Universität Graz



Modellierungswoche der Mathematik

Internationaler Klimaschutz

Bernd Prach, Martina Svibic

Johann Grillitsch

Tobias Dünser, Danijel Obadic, Thomas Schnedl

15.2.2013

Inhaltsverzeichnis

1	Modell 1 - Kostenminimierung im Klimaschutz	3
1.1	Definition des Problems	3
1.2	Darstellung des Modells	3
1.2.1	Kostenfunktionen	3
1.2.2	Daten/Parametrisierung	7
1.2.3	Nash-Gleichgewicht	7
1.2.4	Kooperative Lösung	9
1.2.5	Abweichung einer Ländergruppe von der kooperativen Lösung	10
1.3	Resultate	11
1.4	Schlussfolgerungen	13
2	Modell 2 - Wohlfahrtsmaximierung für unterschiedlich entwickelte Weltregionen	14
2.1	Definition des Problems	14
2.2	Darstellung des Modells	14
2.3	Daten/Parametrisierung	16
2.4	Resultate	16
2.4.1	Nicht-kooperative Lösung: jedes Land maximiert seine eigene Wohlfahrt . .	16
2.4.2	Globale kooperative Lösung: Maximierung der gemeinsamen Wohlfahrt . . .	18
2.5	Schlussfolgerungen	19
3	Modell 3 - Maximierung des Bruttoinlandsprodukts über die Zeit	20
3.1	Definition des Problems	20
3.2	Darstellung des Modells	21
3.2.1	1.Fall: BIP-Maximierung für einen Zeitpunkt in der Zukunft (einmaliges Spiel)	21
3.2.2	2.Fall: BIP Maximierung in 5-Jahresschritten (wiederholtes Spiel)	22
3.3	Daten/Parametrisierung	22
3.4	Resultate	23
3.4.1	1.Fall: BIP Maximierung für einen Zeitpunkt in der Zukunft (einmaliges Spiel)	23
3.4.2	2.Fall: BIP Maximierung in 5-Jahresschritten (wiederholtes Spiel)	23
3.5	Schlussfolgerungen	27
A	Tabellen	28

1 Modell 1 - Kostenminimierung im Klimaschutz

Bernd Prach, Martina Svibic

1.1 Definition des Problems

Ziel des Modells „Kostenoptimierung“ ist es ein Klimaschutzabkommen aufgrund von Kostenoptimierungsproblemen zu simulieren. Die Akteure werden von folgenden drei Gruppen gebildet: Entwicklungsländer (später werden ihre Emissionsmengen als X_1 bezeichnet), Schwellenländer (X_2) und Industrieländer (X_3). Als Kosten werden dabei jene der Emissionsvermeidung und der Schäden aus dem Klimawandel berücksichtigt, wobei sie für die unterschiedlichen Ländergruppen unterscheiden. Unterschieden werden hierbei zwei Betrachtungsweisen, einerseits die Optimierung der Ausgaben aus Sicht des eigenen Landes und andererseits die Kostenoptimierung aus globaler Sicht. Ebenfalls wird der Fall, dass ein Akteur vom globalen Klimaabkommen abweicht, berücksichtigt.

1.2 Darstellung des Modells

Die Lösung der Fragestellung basiert auf zwei variablen Kostenfunktionen. Anschließend werden spieltheoretische Szenarien modelliert und visualisiert.

1.2.1 Kostenfunktionen

Die Gesamt-Kostenfunktion setzt sich aus zwei Funktionen zusammen. Eine dieser Funktionen beschreibt die Kosten, die zum Abbau von Emissionen, also in diesem Modell von CO_2 , notwendig sind, während die andere angibt, wie hoch die Kosten durch Umweltschäden sind in Abhängigkeit zu den vorhandenen globalen Emissionen.

Kostenfunktion für den Abbau von Emissionen Die Funktion P_4 stellt die Kosten für den Abbau von Emissionen in Abhängigkeit zu den aktuellen Landes-Emissionen EmX und dem Kostenproportionalitätsfaktor K dar. Da mit zunehmender Verschmutzung des Planeten die Kosten zuerst sinken, dann aber steigen müssen, und die Ausgaben für CO_2 -Einsparungen zudem rasant wachsen sollen, wurde ein Polynom vierten Grades gewählt und über bekannte Werte der Kosten für den Abbau von Emissionen die folgende Funktion aufzustellen:

$$P_4(X, K, EmX) = \left(1 * 10^{-30} * \left(\frac{Em}{EmX} * X + K \right)^4 + \right. \\ \left. + 7 * \left(\frac{Em}{EmX} * X + K \right) - 1 * 10^{-30} * K^4 - 7 * K \right) * \frac{EmX}{Em}$$

Diese Funktion lässt zwei grundlegende Betrachtungen zu wie in Abbildung 1 und 2 dargestellt. Für die zugrunde liegenden Parameterwerte, siehe Abschnitt 1.2.2. Diese Betrachtungen zeigen, dass es für die Schwellenländer aufgrund ihres sehr hohen CO_2 -Ausstoßes und der hohen Bevölkerung billig ist einzusparen, da geringe Maßnahmen wie Filter sehr effektiv sind. Wohingegen Industrie- wie Entwicklungsländer einen niedrigen CO_2 -Ausstoß haben und daher die geringen Einsparungen, die sie treffen können, sehr teuer sind. Allerdings wird es für die Entwicklungsländer später etwas

Abbildung 1: Kosten der CO₂-Einsparungen pro Einwohner

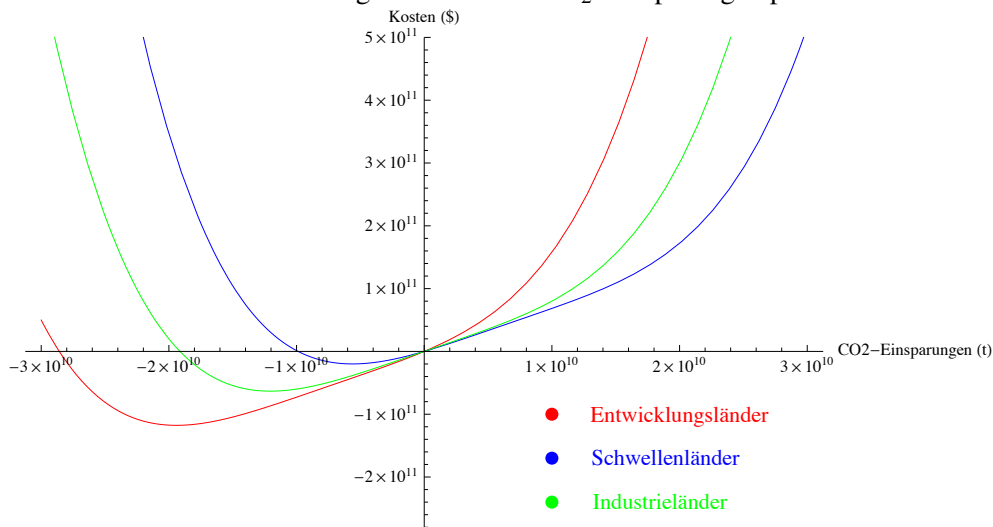
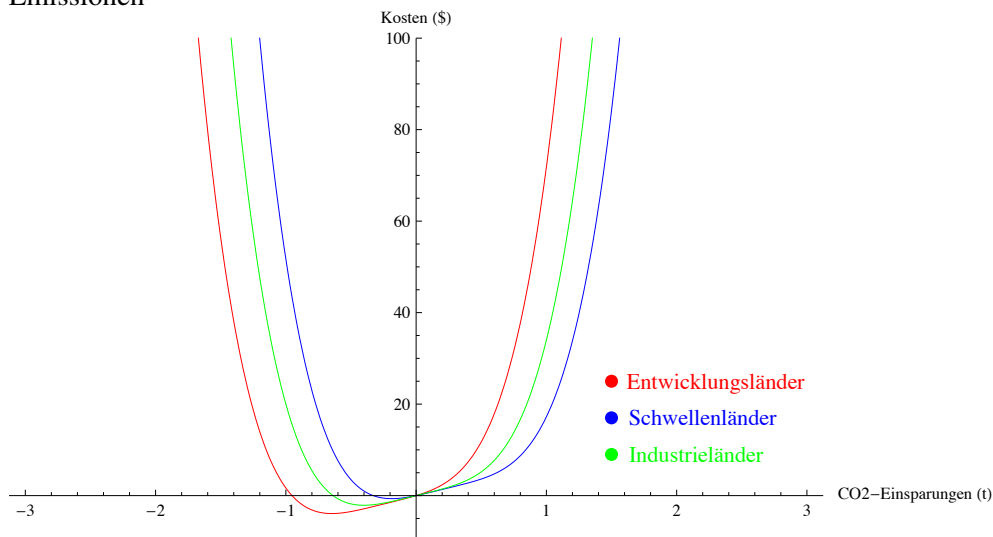


Abbildung 2: Kosten der CO₂-Einsparungen pro Einwohner und in Abhängigkeit von den aktuellen Emissionen



günstiger zu sparen als für die Industrieländer, da die ohnehin schon hochentwickelten Staaten noch effektivere Technologien hervorbringen müssen.

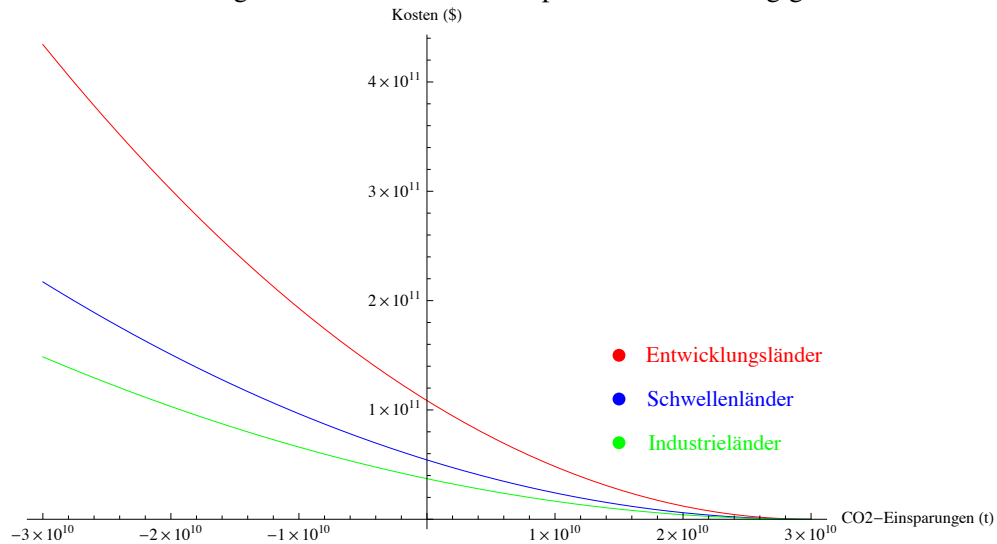
Kostenfunktion durch Umweltschäden Die zweite Kostenfunktion, siehe Abbildung 3, entspricht jenen Kosten, die durch Umweltschäden bei einem bestimmten Ausstoß an CO₂ anfallen in Abhängigkeit zu den Einwohnern EW und c , den aktuellen Schäden durch die aktuellen Emissionen, wobei WEW die Welteinwohner bezeichnet:

$$F(X, EW) = \frac{EW}{WEW} * c * X^2$$

Abzulesen aus dieser Funktion ist, dass die Entwicklungsländer am stärksten von zunehmendem CO₂-Ausstoß betroffen sind, vor den Schwellen- und dann den Industrieländern. D.h. die Funktion

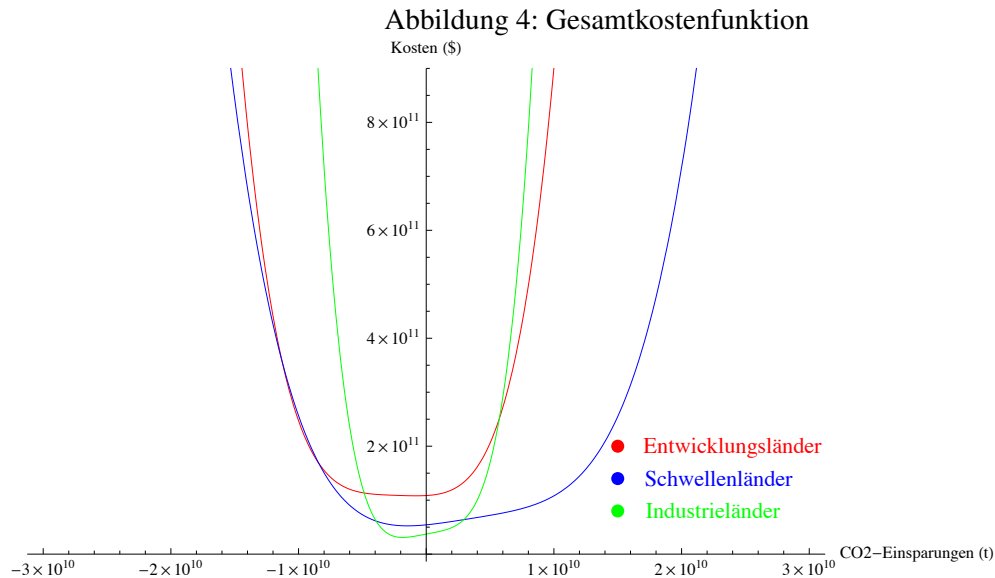
gibt auch ohne Betroffenheitskonstante die Gefahren durch Taifune, Tsunamis und andere Katastrophen wieder.

Abbildung 3: Kosten durch Umweltprobleme in Abhängigkeit zum CO₂-Ausstoß



Gesamtkostenfunktion Durch Addition der oben genannten Kostenfunktionen erhält man die Gesamtkostenfunktion abhängig von der Summe S , den Einwohnern EW , dem Kostenproportionalitätsfaktor K und den jeweiligen Emissionen EmX , wobei S die Summe der Emissionseinsparungen bezeichnet:

$$GK(S, X, EW, K, EmX) = F(Em - S, EW) + P_4(X, K, EmX)$$



1.2.2 Daten/Parametrisierung

Die Daten für das Kostenoptimierungsmodell gliedern sich in zwei Gruppen, nämlich in die Parameter und die Konstanten. Sie sind im folgenden angegeben.

Tabelle 1: Parameterwerte

	Einwohner EW	Emissionen Em	Kostenproportionalitätsfaktor K
Entwicklungsländer	$3,8 * 10^9$	$2.42 * EW_1$	$0.8 * Em_1$
Schwellenländer	$1,9 * 10^9$	$8,12 * EW_2$	$-0,7 * Em_1$
Industrieländer	$WEW - EW_1 - EW_2$	$Em - Em_1 - Em_2$	0

Die Konstanten sind gegeben mit: $c = \frac{2 * 10^{11}}{900 * 10^{18}}$, $WEW = 7 * 10^9$, $Em = 3 * 10^{10}$.

1.2.3 Nash-Gleichgewicht

Das Nash-Gleichgewicht beschreibt jene Situation in nicht-kooperativen Spielen, in der sich keiner der Spieler durch Änderung seiner Strategie verbessern kann. Dieses Gleichgewicht kommt zu Stande indem jeder der Beteiligten bestmöglich auf das Verhalten der Gegenspieler reagiert.

Um zu diesem Gleichgewicht zu kommen, wird die Summe aus F und P_4 nach den jeweiligen Variablen X_1 bis X_3 differenziert, Null gesetzt und schließlich nach der Variable, nach der differenziert wurde, aufgelöst, wobei die Variablen EW_1 , X_1 , K_1 und Em_1 respektive durchlaufen.

Für X_1 sieht die Rechnung wie folgt aus:

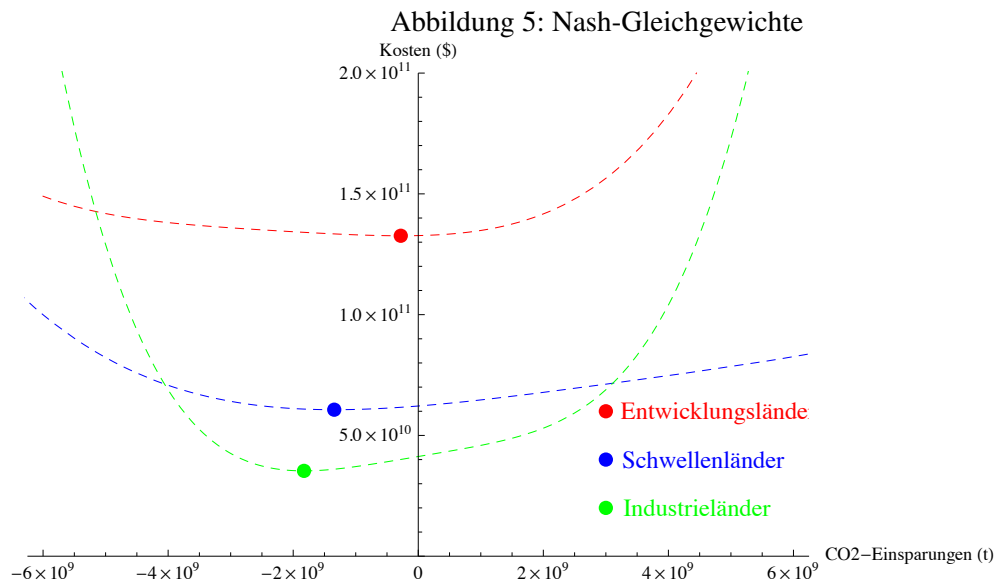
$$\frac{dF(Em - X_1 - X_2 - X_3, EW_1) + P_4(X_1, K_1, Em_1)}{dX_1} = 0$$

Nun werden die Funktionswerte berechnet, wobei NashX mit jeweiligem Index die Nash-Lösung an der angegebenen Stelle bezeichnet:

$$S_I = NashX_1 + NashX_2 + NashX_3$$

$$y_1 = GK(S_I, NashX_1, EW_1, K_1, Em_1)$$

Analoge Berechnungen führen zu allen Nash-Gleichgewichten.



1.2.4 Kooperative Lösung

Die kooperative Lösung entspricht in diesem Modell den minimalen Kosten auf globaler Ebene.

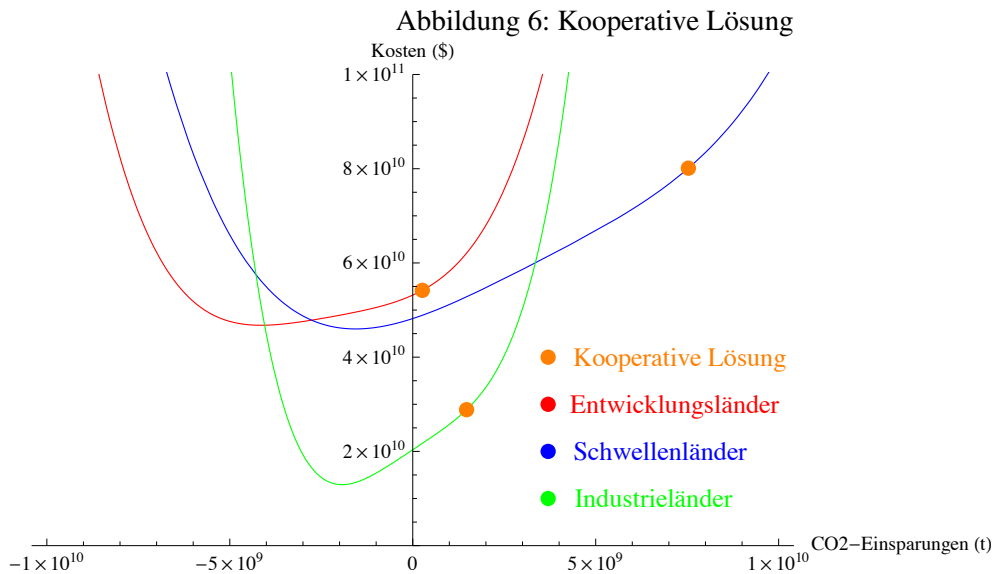
Es wird die Summe der addierten Funktionen F und P_4 aller Länder drei mal getrennt nach X_1 bis X_3 differenziert. Der erste Block sieht wie folgt aus:

$$\begin{aligned} & \frac{dF(Em - X_1 - X_2 - X_3, EW_1) + P_4(X_1, K_1, Em_1)}{dX_1} + \\ & \frac{dF(Em - X_1 - X_2 - X_3, EW_2) + P_4(X_2, K_2, Em_2)}{dX_1} + \\ & \frac{dF(Em - X_1 - X_2 - X_3, EW_3) + P_4(X_3, K_3, Em_3)}{dX_1} = 0 \end{aligned}$$

Dann werden das Gleichungssystem bestehend aus den genannten drei Differenzialen Null gesetzt, und nach den Variablen X_1 bis X_3 aufgelöst. Schließlich werden noch die Funktionswerte bestimmt, analog wie von X_1 , wobei $KoopX$ mit jeweiligem Index die kooperative Lösung an der angegebenen Stelle bezeichnet:

$$S_{II} = KoopX_1 + KoopX_2 + KoopX_3$$

$$y_1 = GK(S_{II}, KoopX_1, EW_1, K_1, Em_1)$$



1.2.5 Abweichung einer Ländergruppe von der kooperativen Lösung

Die minimalen Kosten einer Gruppe lassen sich ermitteln indem die Vereinbarung der kooperativen Strategie gebrochen wird, wobei vorausgesetzt wird, dass sich die anderen beiden Akteure an das Festgelegte halten.

Der Wert für das jeweilige Land setzt sich aus seinem Optimum und dem dazugehörigen Funktionswert zusammen, das hieße für X_1 :

$$\frac{dK(S_{II} - KoopX_1 + X_1, X_1, EW_1, K_1, Em_1)}{dX_1} = 0$$

Dann wird nach X_1 aufgelöst und als $OptiX_1$ bezeichnet. Der Funktionswert davon wird wie folgt berechnet:

$$y_1 = GK(S_{II} - KoopX_1 + OptiX_1, OptiX_1, EW_1, K_1, Em_1)$$

Wie Abbildung 7 verdeutlicht, können die Entwicklungsländer durch einseitiges Abweichen und damit einhergehenden höheren Emissionen geringere Kosten erzielen. Dies führt jedoch zu höheren Kosten für die beiden anderen Ländergruppen. Allerdings sind die Kostenoptima der Länder nicht anzustreben, da die Kosten stark ansteigen, sobald sich die anderen Akteure ebenfalls nicht an die Abmachung der kooperativen Lösung halten.

Analog sind die Berechnungen für X_2 und X_3 zu führen und man erhält jeweils folgende Graphen wie in Abbildung 8 und 9.

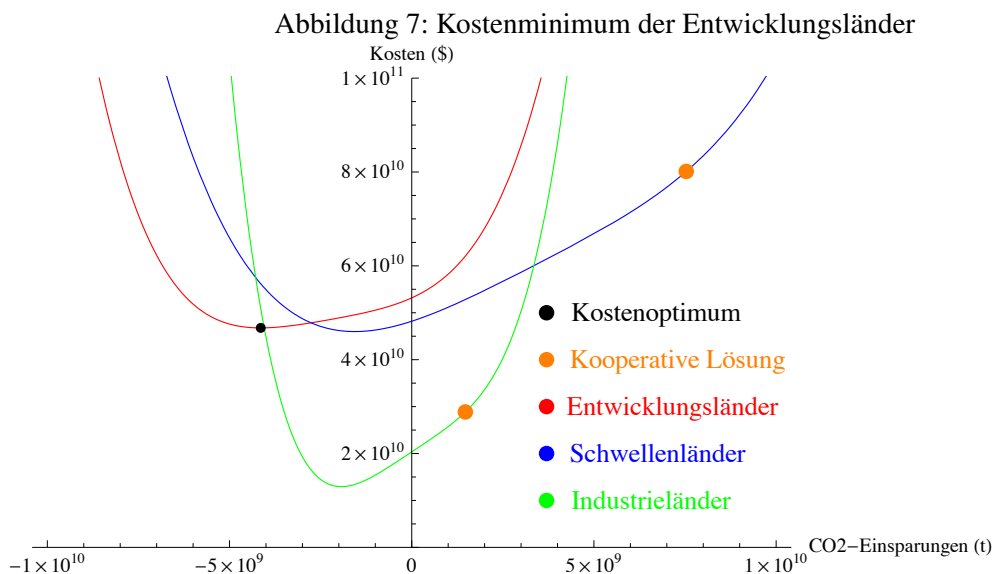


Abbildung 8: Kostenminimum der Schwellenländer

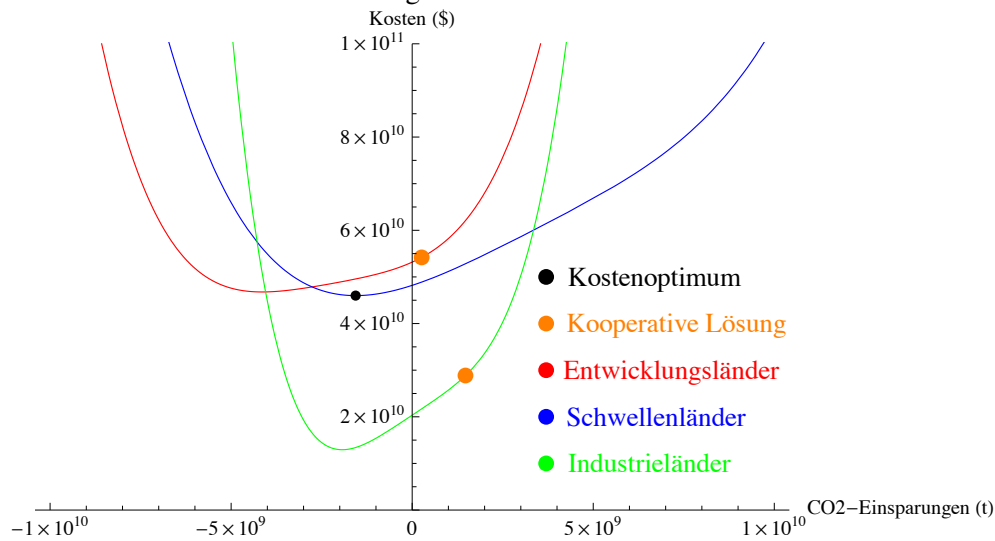
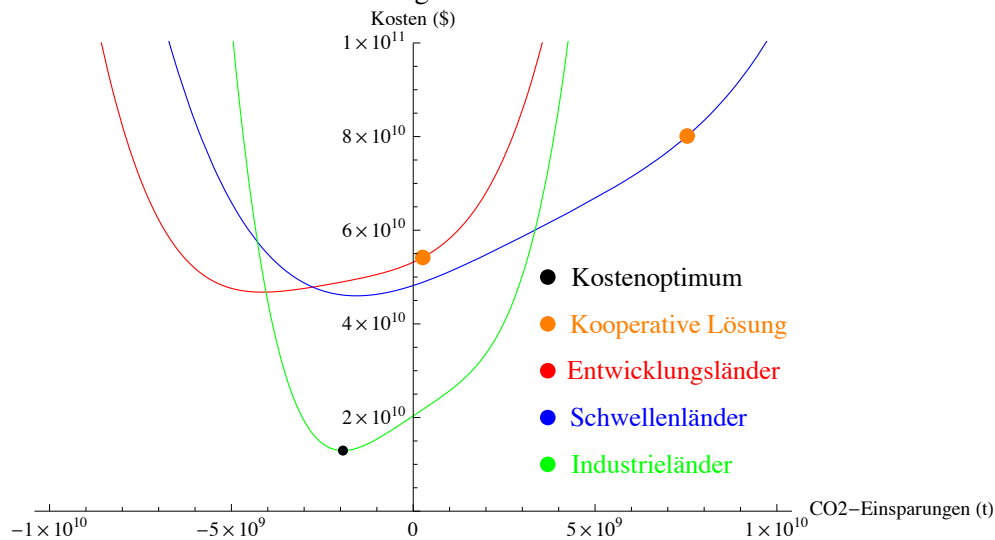


Abbildung 9: Kostenminimum der Industrieländer



1.3 Resultate

Betrachtet man die Werte von Abbildung 10, gelten für die Akteure folgende Schlüsse:

- Entwicklungsländer

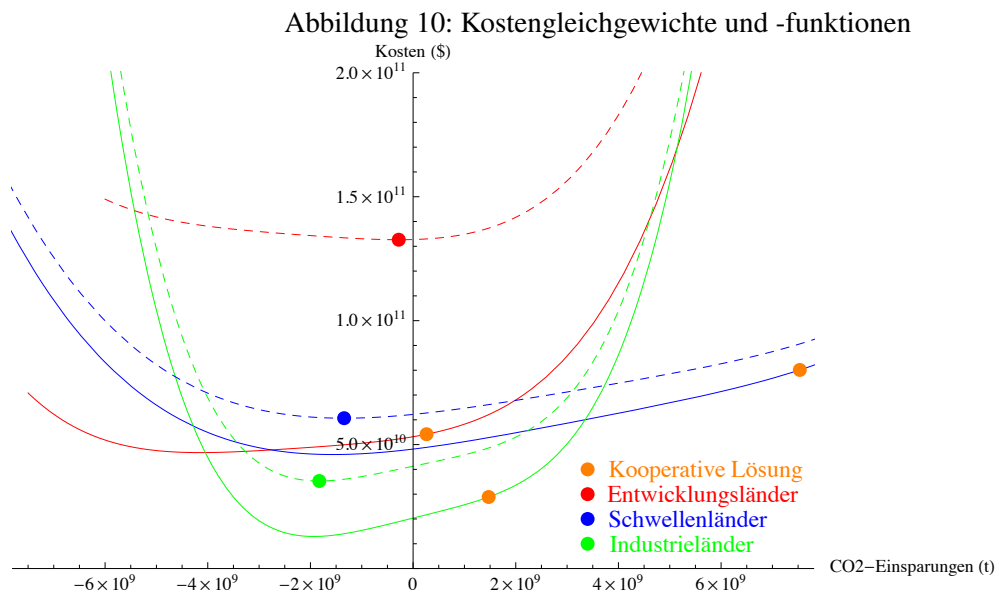
Da die derzeitigen Kosten aufgrund von Umweltschäden \$ 100 Mrd. betragen, und die Kosten eines Nash-Gleichgewichts (siehe strichlierte Linie) deutlich höher sind als bei der kooperativen Lösung (siehe durchgezogene Linie), sollten Entwicklungsländer auf ein globales Klimaschutzabkommen abzielen, in welchem sie 2,8% CO₂ einsparen müssen.

- Schwellenländer

Deren derzeitigen Kosten belaufen sich auf rund \$ 54 Mrd. Somit sind die kooperative Lösung und die damit verbundenen 48,4% an Abbau von Emissionen, wie auch das Nash-Gleichgewicht teurer. Ein Klimaschutzabkommen ist also nicht zu empfehlen.

- Industrieländer

Die jetzigen Kosten durch die Klimaverschmutzung betragen \$ 37 Mrd. Daher ist die kooperative Lösung anzustreben, und ein Nash-Gleichgewicht zu vermeiden. Um ein Klimaschutzabkommen zu erreichen, ist es für die Industriestaaten ratsam in die Schwellenländer zu investieren aufgrund des großen Einsparungspotenzials zu niedrigen Preisen und deren Unwillen zu kooperieren. Um zur kooperativen Lösung zu gelangen, müssen die Industrieländer 27,3% der jetzigen Emissionen einsparen.



Zur Veranschaulichung sind die wichtigsten Kosten in den Tabellen 2 und 3 zusammengefasst. Die jetzigen Kosten aller Staaten aufgrund des Umweltproblems betragen \$ 191 Mrd. Die geringsten Kosten fallen bei der kooperativen Lösung an, nämlich \$ 163 Mrd., also um \$ 28 Mrd. weniger als sie derzeit betragen. Somit ist ein Klimaschutzabkommen zu favorisieren.

Tabelle 2: Nash-Gleichgewicht

	CO ₂ Einsparungen	Kosten der Einsparungen	Kosten der Schäden
Global	$-34,5 \cdot 10^8 \text{ t}$	-20 Mrd. \$	249 Mrd. \$
Entwicklungsländer	$-2,8 \cdot 10^8 \text{ t}$	-2,3 Mrd. \$	135 Mrd. \$
Schwellenländer	$-13,4 \cdot 10^8 \text{ t}$	-6,8 Mrd. \$	68 Mrd. \$
Industrieländer	$-18,3 \cdot 10^8 \text{ t}$	-10,9 Mrd. \$	46 Mrd. \$

Tabelle 3: Kooperative Lösung

	CO ₂ Einsparungen	Kosten der Einsparungen	Kosten der Schäden
Global	92,6 * 10 ⁸ t	67,6 Mrd. \$	95,4 Mrd. \$
Entwicklungsländer	2,6 * 10 ⁸ t	2,3 Mrd. \$	51,8 Mrd. \$
Schwellenländer	75,3 * 10 ⁸ t	54,2 Mrd. \$	25,9 Mrd. \$
Industrieländer	14,7 * 10 ⁸ t	11,1 Mrd. \$	17,7 Mrd. \$

1.4 Schlussfolgerungen

Zusammenfassend ist es für Industrieländer optimal ca. 30% der derzeitigen Emissionen einzusparen. Anzuraten ist auch ein Vertrag in welchem Ausgleichszahlungen an Schwellenländer geleistet werden um diese zu CO₂-Einsparungen zu verpflichten. Zudem profitieren auch Entwicklungsländer sehr stark von einem solchen Abkommen.

Obwohl das vorliegende Modell die wesentlichen Charakteristika von Industrie-, Schwellen-, und Entwicklungsländern widerspiegelt, könnte durch eine stärkere Differenzierung der Regionen, exaktere Daten, Einbeziehung von anderen klimaschädlichen Stoffen und bessere Möglichkeit zur Variierung der Parameter eine Erhöhung der Aussagekraft erzielt werden.

2 Modell 2 - Wohlfahrtsmaximierung für unterschiedlich entwickelte Weltregionen

Johann Grillitsch

2.1 Definition des Problems

Mein Anteil an unserem gemeinsamen Projekt war es, die Kosten und den Nutzen des Klimaschutzes unter Einbezug der gesellschaftlichen Unterschiede der einzelnen Kontinente aufzuarbeiten und darzustellen. Somit lässt sich erkennen, für welchen Kontinent es sich für seine Bevölkerung lohnt, Klimaschutz bis zu welchem Grad zu betreiben.

Ebenfalls konnte dargestellt werden, bis zu welchem Grad es sich global für die gesamte Weltbevölkerung auszahlt, auf die Umwelt zu achten und inwieweit es für die einzelnen Ländergruppen schlechter oder besser ist, sich unkooperativ im Zuge des globalen Klimawandels zu verhalten.

2.2 Darstellung des Modells

Die Kosten von Land i für eine zusätzliche Einheit Klimaschutz werden definiert als:

$$G_i(x_i) = a * x_i^3 - b * x_i^2 + c * x_i - d, \quad x_i = \{x_1, x_2, x_3, x_4, x_5, x_6\}$$

Diese Kosten werden auch als Grenz-Vermeidungs-Kosten bezeichnet und unterscheiden sich zwischen den Ländern einerseits bezüglich der bereits erfolgten Klimaschutzmaßnahmen und andererseits aufgrund von länderspezifischen Einsparungspotenzialen.

- x_1 = Emissionsreduktion für Europa
- x_2 = Emissionsreduktion für Asien
- x_3 = Emissionsreduktion für Afrika
- x_4 = Emissionsreduktion für Australien
- x_5 = Emissionsreduktion für Nordamerika
- x_6 = Emissionsreduktion für Südamerika

Der Gewinn aus einer zusätzlichen Einheit Klimaschutz von Land i für Land i sei gegeben mit:

$$H1_i(x_i) = \frac{-g_i * x_i}{(f)} + g_i$$

Dieser Gewinn wird als Grenz-Gewinn bezeichnet und ist ident mit den vermiedenen Kosten des Klimawandels für Land i .

Der Gewinn aus einer zusätzlichen Einheit Klimaschutz von Land i für alle Länder sei gegeben mit:

$$H2_i(x_i) = \frac{-g * x_i}{(f)} + g$$

Dieser Gewinn wird als Grenz-Gewinn bezeichnet und ist ident mit den vermiedenen Kosten des Klimawandels für alle Länder.

Wohlfahrt für Land i ist die Differenz zwischen seinen Kosten für eine Einheit Klimaschutz und seinen Gewinnen aus dieser Einheit Klimaschutz. Die Wohlfahrt ist dann maximal, wenn die Kosten der letzten Einheit Klimaschutz genau dem Gewinn daraus entsprechen:

$$W1_i(x_i) = H1_i(x_i) - G(x_i) = 0 \implies x_i^{solo}$$

Wobei „solo“ die Strategie eines Landes bezeichnet, wenn es die Effekte des Klimaschutzes nur für sich selbst berücksichtigt.

Wenn Land i nicht nur die eigenen Gewinne aus dem Klimaschutz berücksichtigt, so wählt es sein Vermeidungsniveau x_i so, dass

$$W2_i[x_i] = H2_i[x_i] - G[x_i] = 0 \implies x_i^{koop}$$

Wobei *koop* die Strategie eines Landes bezeichnet, wenn es die Effekte des Klimaschutzes für alle Länder berücksichtigt.

Abhängig davon wie viele Länder vom Klimaschutz profitieren, wird die kooperative Emissionsreduktionsmenge x_i^{koop} deutlich größer als die nicht-kooperative Menge x_i^{solo} sein.

2.3 Daten/Parametrisierung

Parameter	Minimaler Wert	Maximaler Wert	Beschreibung
a	0,006		Kostenzuwachsparemeter für Vermeidungskosten
b_i	0,360	0,385	Kostenzuwachsparemeter für Vermeidungskosten (abhängig von länderspezifischen Einsparungspotenzialen)
c_i	8,15	8,40	Innovationsfähigkeit eines Landes
d_i	51,5	56,5	Entwicklungsstand eines Landes; je weniger entwickelt, desto größer das Einsparungspotenzial
f_i	33,5		Momentane CO ₂ Emissionen weltweit (10 ⁹ pro Gt)
g_i	1,2	32	Anteil an globalen Emissionen je Land (Annahme: Vulnerabilität variiert mit Anteil der Emissionen je Land an Gesamtemissionen)

2.4 Resultate

Die Ergebnisse dieser Arbeit bestätigen durchaus die bisherigen Anstrengungen der einzelnen Ländergruppen. Es zeigt, dass das Kyoto-Abkommen der EU-Staaten ein richtiger Schritt zum Ziel einer intakten Umwelt ist, aber eben nicht die optimale Lösung für alle ist. Ohne die Zusammenarbeit aller Länder wird wohl kaum eines freiwillig mehr einsparen, als es unbedingt sein

muss, da ansonsten die anderen Staaten von dem eigenen Einsatz profitieren. Doch werden wir uns überlegen müssen, was jeder einzelne zur Lösung der bisher größten Herausforderung der Menschheit beitragen muss, damit wir auch in Zukunft in einer intakten Umwelt auf diesem wundervollen Planeten verweilen dürfen.

2.4.1 Nicht-kooperative Lösung: jedes Land maximiert seine eigene Wohlfahrt

Abbildung 12 vergleicht die kooperativen Emissionsmengen und Preise für die verschiedenen Ländergruppen.

Europa Aufgrund der bereits jetzt schon relativ gut aufgestellten Klimaschutzgesetze im Kyoto-Abkommen liegt Europa derzeit relativ genau dort im meinem Modell, wo es im Idealfall liegen könnte. Die Auswertung der Daten hat ergeben, dass eine Investition in den Klimaschutz von 8,65 USD pro emittierter Tonne CO₂ den bestmöglichen Effekt erzielen würde, was einer Reduktion von 19% der Weltemissionen gleichkommen würde (siehe Tabellen 6 und 7). Damit würden die Einwohner der besagten Gebiete zwar ein wenig mehr für den Umweltschutz aufwenden müssen, würden dafür aber auch ein wenig von der verbesserten Umwelt und dem dadurch erhabenerem Lebensgefühl profitieren und somit 20% der dadurch weltweit entstehenden Wohlfahrt erhalten.

Asien Da der asiatische Weltteil noch ein gewaltiges Potenzial bei der Reduktion von CO₂ birgt, ist es auch dementsprechend rentabel größere Mengen einzusparen, da zusätzlich auch noch enorm viele Menschen davon profitieren würden. Natürlich würde dies mit Kosten von in etwa 7,90 USD pro emittierter Tonne CO₂ einhergehen, was bei den geschätzten zukünftigen Wirtschaftsleistungen aber nicht sonderlich ins Gewicht fallen dürfte. Somit erhalten sie aufgrund der vielen Nutznießer 28% des dadurch weltweit lukrierten Nutzens mit Einsparungen von nur einem Viertel der Weltemissionen (siehe Tabellen 6 und 7).

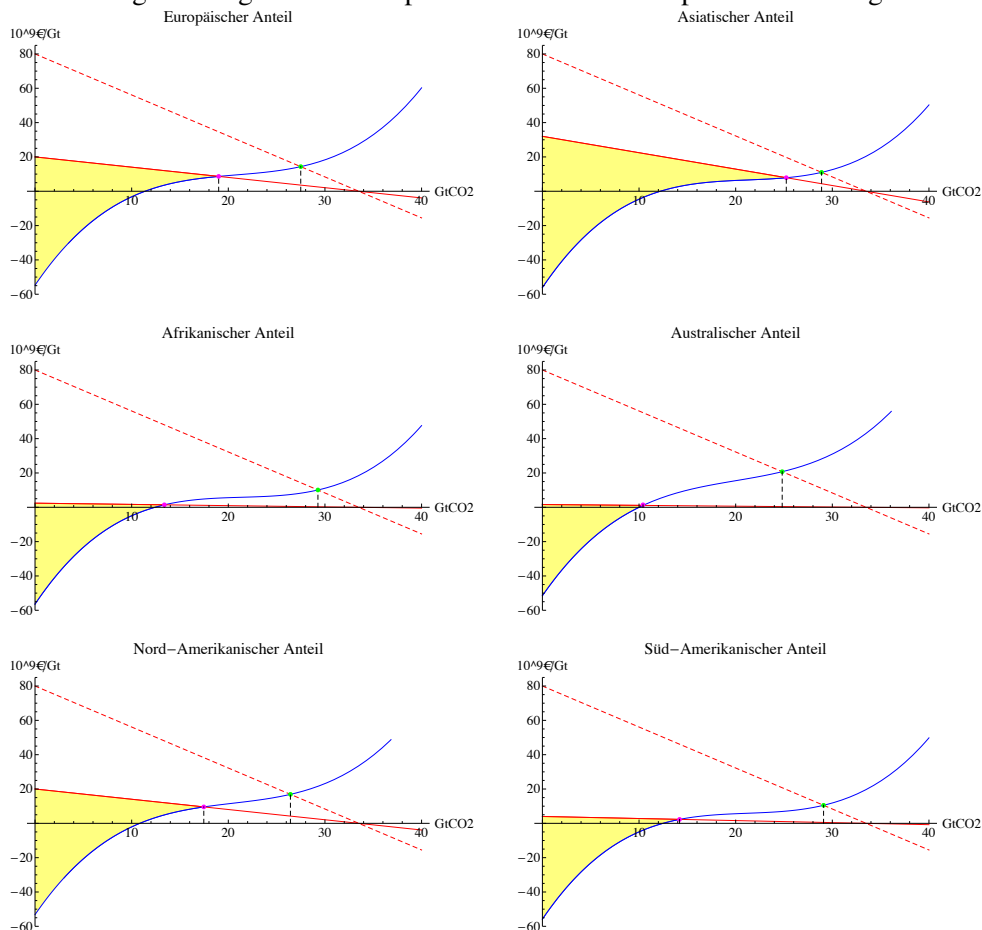
Afrika Bisher ist es noch so, dass Afrika keinen besonders großen Anteil zum Klimawandel beiträgt und daher auch kaum Emissionen reduzieren muss. Schon der Betrag von knapp 1,5 USD pro Tonne CO₂ den wohlfahrtsmaximierenden Effekt erzielen würde. Dadurch würden sie effektiv sogar weniger Wohlfahrt erhalten, als sie durch ihre Bemühungen in den Umweltschutz investieren. Jedoch muss man bedenken, dass Afrika in den nächsten Jahrzehnten wohl einen unglaublichen Aufschwung erleben wird und somit auch einen gewichtigen Faktor in der Weltklimapolitik.

Australien Australien steckt ein wenig in der Zwickmühle, da es solch einen geringen Anteil an den Weltemissionen beiträgt und gleichzeitig aber stark unter den Umweltschäden zu leiden hat. Wenn es sich alleine für einen stärkeren Umweltschutz einsetzen würde, wäre die nicht-kooperative Lösung gerade einmal bei knapp unter 1 USD pro Tonne CO₂, da die anderen Kontinente den Klimawandel trotzdem weitaus stärker beeinflussen würden. Deswegen würden sie genau wie die afrikanischen Staaten weniger Nutzen aus dem Klimaschutz ziehen, als sie in ihn investieren würden. Nord-Amerika: Durch den großen Einfluss von Nord-Amerika auf den Klimawandel würde es sich für diese Länder schon bald rentieren die Umwelt zu schützen, da auch genug Menschen dann von den verbesserten Lebensbedingungen profitieren würden. Jedoch ist es dort auch schon recht teuer,

in weitere Maßnahmen zu investieren und würde sich im Idealfall mit 9,5 USD je emittierter Tonne CO₂ zu Buche schlagen. Somit erhalten sie 19% der dadurch weltweit entstehenden Wohlfahrt während sie dafür nur 18% der weltweiten Emissionen einsparen (siehe Tabellen 6 und 7).

Süd-Amerika Da Südamerika derzeit noch keine besonders große Rolle als CO₂ Emittent im globalen Klimawandel spielt, würde es schon bei nur 2,3 USD pro Tonne CO₂ seine individuelle Wohlfahrt erreichen. Süd-Amerika wäre aber auch der größte Verlierer bei der Kosten/Nutzen Relation, da sie zwar insgesamt 14% der Weltemissionen einsparen müssten, aber nur 12% Nutzen daraus ziehen könnten (siehe Tabellen 6 und 7). Den größten Teil dieser Einsparungen würde wohl schon der Schutz des Regenwaldes ausmachen, was jedoch durch das enorme Wachstum von Süd-Amerika in den nächsten Jahrzehnten nicht mehr ausreichen würde.

Abbildung 11: Vergleich der kooperativen und nicht-kooperativen Lösungen nach Ländergruppen

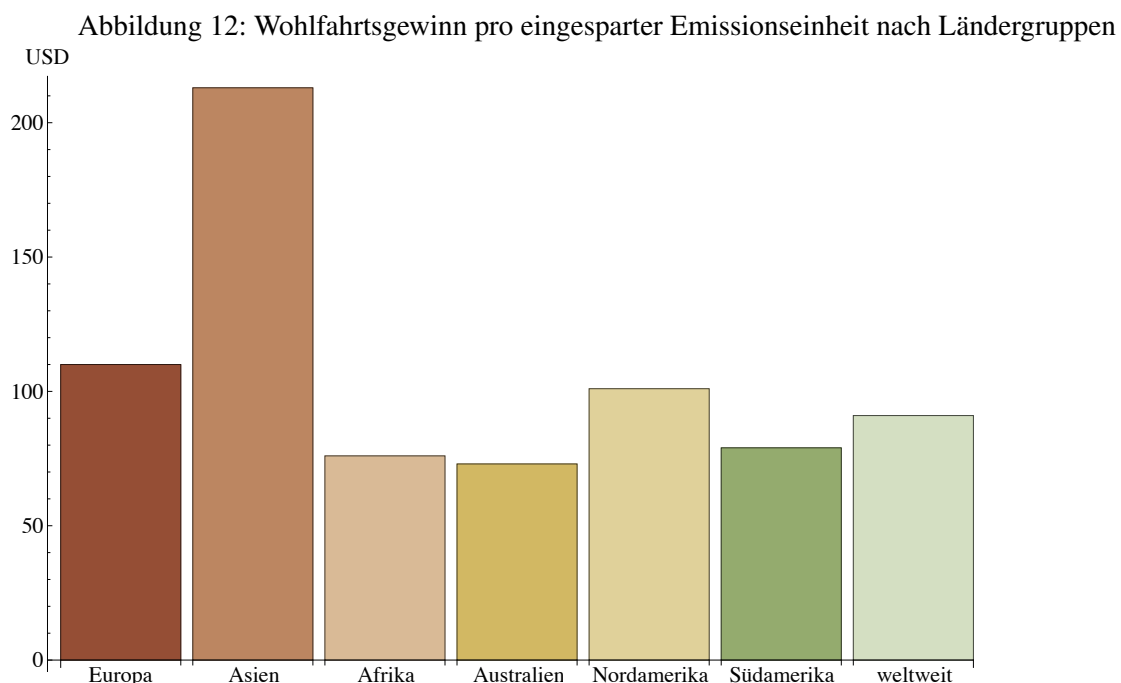


All diese Bemühungen müssten zuerst einmal von jedem Land konsequent erfüllt werden, was derzeit aber nur unzureichend umgesetzt wird. Doch selbst wenn alle Kontinente ihren Part erfüllen, ist das bei weitem noch nicht das Potenzial, welches durch die Kooperation aller Länder erreicht werden könnte, wie der folgende Abschnitt verdeutlicht. Dadurch würden sich in allen Regionen sowohl die Umweltschäden reduzieren und gleichzeitig das Leben durch die ganzen innovativen Verbesserungen unheimlich verbessern.

2.4.2 Globale kooperative Lösung: Maximierung der gemeinsamen Wohlfahrt

Jedes Land tut sein Möglichstes zum Schutz des Klimas, d.h. es wählt jenes Niveau von Emissionsvermeidung, dass die globale Wohlfahrt maximiert. Dadurch erhöhen sich die Kosten im Durchschnitt auf knapp 14 USD pro ausgestoßener Tonne CO₂. Außerdem zeigt sich, dass die Verteilung der Emissionsreduktionen auf die verschiedenen Kontinente deutlich homogener als in der unkooperativen Lösung ausfällt, da nun auch kleine Kontinente einen deutlich höheren Beitrag zum Klimaschutz im Interesse der anderen Kontinente bringen müssen.

Abbildung 12 zeigt die durchschnittliche Wohlfahrtsveränderung je eingesparter Einheit, wenn alle Länder anstelle ihrer unkooperativen Emissionsmenge die kooperative Menge wählen. Es zeigt sich dass diese Nutzen-Kosten-Relation für Asien besonders hoch ist, gefolgt von Europa und Nordamerika.



Dabei würden die ärmeren Länder mit den geringeren pro-Kopf-Emissionen auch einen ein wenig geringeren Beitrag leisten müssen, die höher entwickelten Länder mit den höheren pro-Kopf-Emissionen im Gegenzug etwas mehr. Aufgrund der höheren Beträge für den Klimaschutz würde sich die Lebensqualität durch das Ausbleiben von Naturkatastrophen und durch die Erleichterungen eines innovativeren Lebensstils weitaus stärker verbessern, als es durch die unkooperativen Lösungen jemals möglich gewesen wäre.

2.5 Schlussfolgerungen

Doch bis es zu einer kooperativen Zusammenarbeit unter den einzelnen Ländergruppen kommen kann, muss sich noch einiges tun: Die Forschung und Entwicklung von neuen Möglichkeiten zum effektiveren Umweltschutz muss gefördert werden; die ärmeren Regionen der Welt müssen auf den

Stand gebracht werden, sich die neuen Erfindungen überhaupt leisten zu können; die Menschen in den privilegierteren Ländern müssen umdenken und erkennen, dass desto länger sie warten sich um unseren Planeten zu kümmern, desto teurer wird es darauf zu (über)leben.

3 Modell 3 - Maximierung des Bruttoinlandsprodukts über die Zeit

Tobias Dünser, Danijel Obadic, Thomas Schnedl

3.1 Definition des Problems

Nachdem der Klimawandel zu einem gravierenden globalen Umweltproblem geworden ist, versuchen wir die Schwierigkeiten im Klimaschutz mit einfachen Modellen zwischen verschiedenen Ländern darzustellen, zu erklären und Lösungen zu erarbeiten. Fragen, die sich stellen, sind beispielsweise, wieso niemand mit einer Reduktion von Emissionen beginnt oder ob es sich wirtschaftlich überhaupt lohnt, jetzt oder in Zukunft den CO₂-Ausstoß seines eigenen Landes zu reduzieren. Außerdem wollen wir die Unterschiede zwischen statischen (Fall 1) und dynamischen (Fall 2) Lösungen untersuchen, wobei wir annehmen, dass dynamische und kurzfristige Lösungen unter Entscheidungsträgern der jeweiligen Staaten verbreiteter sind als langfristige, und herausfinden, welche der beiden Methoden die wirtschaftlich bessere ist.

Um dieses Problem zu untersuchen, entwickeln wir zwei verschiedene Modellvarianten. Im 1. Fall ist es das Ziel unseres Modells, das jeweilige Bruttoinlandsprodukt (BIP) der einzelnen Akteure Nordamerika, Südamerika, Afrika, Europa und Asien für einen bestimmten Zeitpunkt (in 50 Jahren) zu maximieren. Das BIP wächst einerseits mit einer gegebenen Wachstumsrate, andererseits wird es durch zeitlich verzögert auftretende Schäden aus dem Klimawandel sowie durch Ausgaben für Klimaschutz reduziert.

Jedes Land wählt nun jene Menge Klimaschutz (in unserem Fall CO₂ Reduzierung) um sein eigenes BIP zu maximieren. Wenn diese Maximierung die Effekte auf andere Länder nicht berücksichtigt, wird diese Lösung als nicht kooperative Lösung bezeichnet. Das bedeutet, dass sie sich entweder dafür entscheiden, einen Teil ihres BIPs für die Reduzierung von CO₂ einzusetzen, oder stattdessen noch mehr CO₂ auszustoßen und dafür aber verstärkte Kosten des Klimawandels (Katastrophen etc.) in Kauf zu nehmen. Das BIP besteht also aus den Einnahmen des Akteurs, den Ausgaben für den Klimaschutz und den Schäden durch den Klimawandel.

Die zweite Variante von Fall 1 hat das Ziel, das Bruttoinlandsprodukt der ganzen Welt nach einer bestimmten Zeit zu maximieren indem wiederum Emissionsreduktionen für die unterschiedlichen Länder gewählt werden. Es müssen also alle Akteure zusammenarbeiten und eine gemeinsame kooperative Lösung finden, damit das globale BIP, welches aus der Summe der BIPs aller Akteure besteht, den höchstmöglichen Wert erreicht.

Im 2. Fall ist es das Ziel unseres Modells, die Wahl von Emissionsreduktionen realistischer und aufgrund von unkooperativen Lösungen, d.h. Lösungen bei denen jeder Staat nur auf die Maximierung seines eigenen BIPs achtet, zu simulieren. Dazu wird das optimale Maß an Emissionsreduktion für alle fünf Jahre neu berechnet und das Modell wird somit dynamisch und realitätsnäher.

3.2 Darstellung des Modells

3.2.1 1.Fall: BIP-Maximierung für einen Zeitpunkt in der Zukunft (einmaliges Spiel)

Die Funktion für das eines Landes i nach 50 Jahren sieht folgendermaßen aus:

$$BIP_{i,t=50}(a_1, a_2, a_3, a_4, a_5) = BIP_{i,t=0} * (WWF^t) - P(a_1, a_2, a_3, a_4, a_5) * JS * BIP_{i,t=0} * (VDS^{t-VZ})(K(1 - a_i)) * BIP_{i,t=0} * (SKK^t), \quad i = 1, \dots, 5 \quad (1)$$

Wie aus dieser Gleichung ersichtlich ist, hängt das BIP des Landes i nicht nur von den eigenen Emissionen ab sondern auch von den Emissionsniveaus der anderen Länder. Der erste Ausdruck auf der rechten Seite steht für das Wachstum des BIP des jeweiligen Akteurs, der zweite Ausdruck steht für den Schaden, der durch den Klimawandel entsteht und der dritte Ausdruck steht für die Kosten des Klimaschutzes.

In Gleichung (1) bedeutet $BIP_{i,t=t}(a_1, a_2, a_3, a_4, a_5)$ das BIP des Akteurs nach einer gewissen Zeit bei gewissen Werten von a_1 - a_5 . $BIP_{i,t=0}$ steht für das anfängliche Bruttoinlandsprodukt des jeweiligen Akteurs in Prozent (also 100 Prozent). WWF repräsentiert den Wirtschaftswachstumsfaktor, der natürlich für jeden Akteur verschieden ist, weil nicht jede Region der Welt ein gleiches Wirtschaftswachstum hat t bedeutet die Zeit in Jahren. VZ gibt die Verzögerungszeit in Jahren, nach der sich der Klimawandel auswirken wird an; somit treten Klimafolgen erst mit einigen Jahren Verspätung auf, was einen Anreiz zu später oder niedriger Emissionsreduktion gibt. JS ist die jährlichen Schäden in Prozent vom Bruttoinlandsprodukt. VDS steht für die Verschlimmerung des Schadens des Klimawandels in Prozent pro Jahr. $K(1 - a_i)$ repräsentiert die Kosten von $1-a_i$ CO₂-Reduktion in Prozent vom BIP. SKK gibt die Senkung der Kosten des Klimaschutzes in Prozent pro Jahr an.

Der Gesamtausstoß des CO₂ in Prozent vom Anfangswert wird durch folgende Gleichung beschrieben:

$$P(a_1, a_2, a_3, a_4, a_5) = (a_1 * m + a_2 * n + a_3 * o + a_4 * p + a_5 * q)$$

Wobei a_1, a_2, a_3, a_4, a_5 für den CO₂-Ausstoß des jeweiligen Akteurs in Prozent vom Anfangswert steht.

m, n, o, p, q sind Faktoren, die für den Anteil an ausgestoßenem CO₂ in Prozent vom Gesamtausstoß stehen, die natürlich von Akteur zu Akteur aufgrund der unterschiedlichen Größen der Akteure verschieden sind.

Um das BIP maximierende Emissionsniveau (a_i -Wert) berechnen zu können, muss Gleichung 1 für jeden Akteur nach a_i abgeleitet werden und dann im nächsten Schritt Null gesetzt werden, um die Maximums-Stelle zu finden. So erhält man den optimalen a -Wert, und der Akteur weiß nun, wie er sich in den nächsten 50 Jahren bezüglich der CO₂-Reduzierung zu verhalten hat. Diese Lösung wird als nicht-kooperative Lösung bezeichnet.

In der nicht-kooperativen Lösung wird die Summe der BIPs aller Länder maximiert:

$$\sum_{i=1}^{i=5} BIP_{i,t=50}(a_1, a_2, a_3, a_4, a_5)$$

Leitet man diese Funktion nach a_1 bis a_5 ab, so ergibt sich die kooperative Lösung.

3.2.2 2.Fall: BIP Maximierung in 5-Jahresschritten (wiederholtes Spiel)

In der Realität ist es häufig so, dass Akteure nicht eine Strategie wählen, die ihren Nutzen oder Gewinn in 50 Jahren maximieren sondern dass die Entscheidungszeiträume deutlich kürzer sind. Wir unterstellen nun, dass die Akteure jeweils nur das BIP der nächsten fünf Jahre in ihrer Entscheidung berücksichtigen und somit alle fünf Jahre neu bestimmen, wieviel Emissionen ihr Land ausstoßen soll.

Auch in diesem Fall sieht die Zielfunktion wie Gleichung 1 aus mit dem einzigem Unterschied, dass t nun in 5 Jahresschritten gewählt wird.

$$BIP_{i,t=50}(a_1, a_2, a_3, a_4, a_5) = BIP_{i,t=0} * (WWF^t) - P(a_1, a_2, a_3, a_4, a_5) * JS * BIP_{i,t=0} * (VDS^{(t-VZ)}) \\ (K(1 - a_i)) * BIP_{i,t=0} * (SKK^t), \quad i = 1, \dots, 5$$

Um das Modell realitätsnäher zu gestalten, fügten wir durch eine Schleife Zeitschritte von fünf Jahren ein, an deren Enden die einzelnen Akteure die Möglichkeit hatten, die Menge an ausgestoßenem CO₂ so zu bestimmen, um ihr BIP für diese Zeitperiode von fünf Jahren zu maximieren, in diesem Fall nicht kooperativ, da dieser Fall für die Wirklichkeit viel realitätsnäher und wahrscheinlicher ist, da die einzelnen Akteure in Wirklichkeit immer zuerst auf ihr eigenes Wohl achten, als auf das der anderen. Durch diese Veränderung wurde das Modell dynamisch und damit auch auf die Wirklichkeit übertragbar. Die jeweiligen Maxima der BIPs wurden wieder durch die erste Ableitung der BIP-Funktion nach a_i und anschließendem Nullsetzen der Ableitung berechnet. Dieser Schritt wurde durch einen Loop, zu Deutsch Schleife, vom Computer automatisch alle fünf Jahre für einen Zeitraum von fünfzig Jahren wiederholt und es wurde jeweils das im vorigen Schritt neu berechnete BIP in die BIP-Funktion eingesetzt.

Für dieses Modell wird keine kooperative Lösung betrachtet, weil die Nachjustierung alle fünf Jahre einen kurzsichtigen Entscheidungsträger beschreibt, der weder die Effekte auf andere Länder zum gleichen Zeitpunkt noch zukünftige Effekte für das eigene Land berücksichtigt.

3.3 Daten/Parametrisierung

Tabelle 4 fasst die gewählten Parameterwerte zusammen.

Tabelle 4: Parameterwerte nach Ländergruppen

	WWF	$m/n/o/p/q$	JS	$BIP_{i,t=0}$	VDS	VZ	SKK	$K(1 - a)$
Nordamerika	1,02	0,2689	0,05	1	1,01	15	0,98	$1.3954 * \exp^{-6.868*a}$
Südamerika	1,032	0,0504	0,05	1	1,015	15	0,98	$1.3954 * \exp^{-6.868*a}$
Afrika	1,035	0,0420	0,05	1	1,015	15	0,98	$1.3954 * \exp^{-6.868*a}$
Europa	1,015	0,2185	0,05	1	1,01	15	0,98	$1.3954 * \exp^{-6.868*a}$
Asien	1,035	0,4202	0,05	1	1,015	15	0,98	$1.3954 * \exp^{-6.868*a}$

Die Unterteilung der Welt in diese Regionen haben wir so vorgenommen, dass Bereiche mit ähnlichen klimatischen(Meer, Wüsten, Wälder) und wirtschaftlichen(Industrie- oder Entwicklungsland) Voraussetzungen zusammengefasst wurden.

Die Parameter, wie jährliches Wirtschaftswachstum oder Kosten für Klimaschutz, wurden möglichst realitätsnah mit Daten aus dem Internet (worldbank.org) bestimmt.

3.4 Resultate

3.4.1 1.Fall: BIP Maximierung für einen Zeitpunkt in der Zukunft (einmaliges Spiel)

Die Industrieländer und Schwellenländer werden ihren CO₂-Ausstoß sicher zurückschrauben müssen, um ihr BIP zu maximieren, weil sie derzeit einen Großteil zum Gesamtausstoß beitragen. Tabelle 5 stellt unsere Ergebnisse dar, unter der Voraussetzung, dass jeder Akteur 50 Jahre lang das gleiche Konzept verfolgt, einmal unter dem Ziel das eigene BIP zu maximieren (nicht kooperativ), und einmal unter dem Ziel, das globale BIP zu maximieren (kooperativ).

Tabelle 5: Vergleich der kooperativen und nicht-kooperativen Lösung bei Optimierung für den Zeitpunkt in 50 Jahren

	CO ₂ Ausstoß rel. zu $t = 0$ in %		BIP rel. zu $t = 0$ in %		Differenz BIP in %
	koop.	nicht koop.	koop.	nicht koop.	
Nordamerika	76	50	264	0,15	
Südamerika	98	75	477	478	0,37
Afrika	100	78	552	554	0,33
Europa	79	54	205	206	0,32
Asien	67	44	552	552	0,02
Weltweit	75	51	2049	2054	0,23

Um das eigene BIP nach 50 Jahren zu maximieren, ist es für Nordamerika, Europa und Asien wichtig, ihren CO₂-Ausstoß stark zu reduzieren, denn andernfalls würden sie auf lange Sicht von den Kosten der Schäden des Klimawandels, wie zum Beispiel überschwemmungen, erdrückt werden. Afrika und Südamerika hingegen müssen kaum etwas verändern, um ihr eigenes BIP nach 50 Jahren maximiert zu haben, weil sie prozentuell sehr wenig zum Gesamtausstoß beitragen.

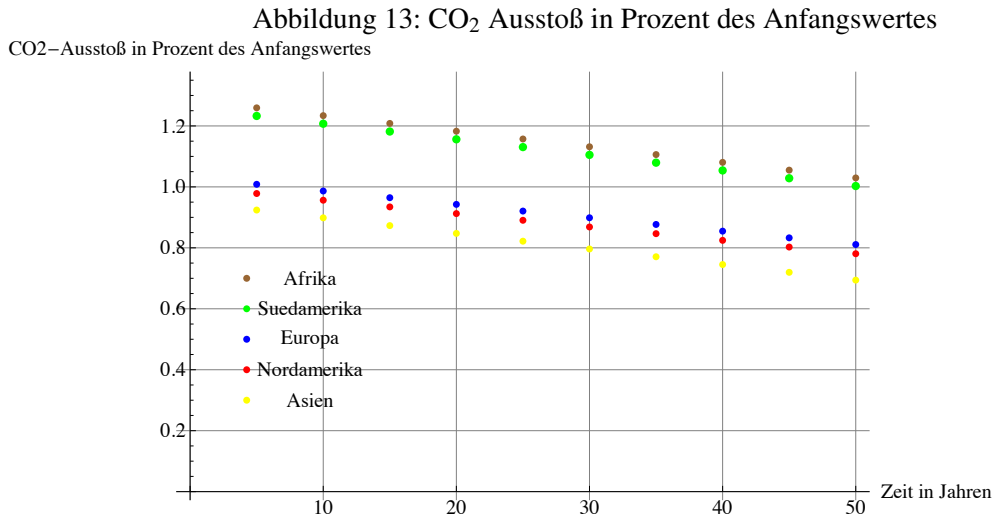
Um das globale BIP nach 50 Jahren zu maximieren, ist es für alle Akteure wichtig, ihren CO₂-Ausstoß zu reduzieren. Für Nordamerika, Europa und Asien jeweils um ca. die Hälfte und für Afrika und Südamerika jeweils ca. um ein Viertel.

Außerdem ist zu erkennen, dass das globale, kooperative BIP prozentuell größer ist, als die Summe aller unkooperativen BIPs der einzelnen Akteure. Dies ist insofern erstaunlich, als die globale Emissionsreduktion in diesem Fall im Schnitt um beinahe ein Drittel höher ausfällt und somit die Vermeidungskosten auch deutlich höher sind.

3.4.2 2.Fall: BIP Maximierung in 5-Jahresschritten (wiederholtes Spiel)

Die Industrieländer, vor allem aber auch die Schwellenländer wie China und Indien, werden ihre CO₂-Ausstöße stark reduzieren müssen, da sie derzeit einen für einen Großteil der globalen Gesamtausstöße verantwortlich sind.

Grafik 13 stellt den prozentuellen CO₂-Ausstoß relativ zum Basisjahr (t=0) für jeden Akteur dar, vorausgesetzt dass jeder Akteur nicht kooperativ handelt und auf die Maximierung seines eigenen BIPS für jede der 5-Jahres-Perioden bedacht ist.



In dieser ersten Grafik ist zu erkennen, dass die Entwicklungsländer ihren CO₂-Ausstoß anfangs sogar erhöhen müssen um ihr BIP zu maximieren und damit ihre Entwicklung zu fördern. Es würde für sie nichts bringen, ihren CO₂-Ausstoß zu verringern, weil die Kosten in keinem Verhältnis zum Nutzen stehen, da diese Länder nur für einen kleinen Teil des Gesamtausstoßes verantwortlich sind. Für große Länder wie z.B. China und Indien, die insgesamt sehr viel zum Gesamtausstoß beitragen, ist es von Vorteil, den CO₂-Ausstoß sofort zu senken um zukünftige Kosten durch Umweltschäden zu vermeiden.

In Grafik 14 zeigen wir das zugehörige und maximierte BIP zur nicht kooperativen Lösung jedes einzelnen Akteurs, basierend auf den prozentuellen Ausstoßmengen, die in der ersten Grafik dargestellt sind.

Aus der zweiten Grafik kann man entnehmen, dass durch den schrittweisen, stetigen Abbau von CO₂-Ausstößen über einen längeren Zeitraum, das BIP trotzdem relativ hoch gehalten werden kann. Vor allem die BIPs der Entwicklungs- und Schwellenländer steigen aufgrund des hohen Wirtschaftswachstums und des weltweiten, stetigen Abbaus an CO₂ stärker an, als die BIPs der Industrieländer.

Nun zeigen wir in Grafik 15, wie sich die Anteile am gesamten CO₂-Ausstoß der einzelnen Länder während des gesamten Zeitraums verändern.

Da es für Asien gewinnbringend ist, seine CO₂-Ausstöße stärker zu senken als die anderen Akteure, weil es einen Großteil zum Gesamtausstoß an CO₂ beiträgt, sinkt sein Anteil am gesamt

Abbildung 14: BIP dynamisch aller Länder

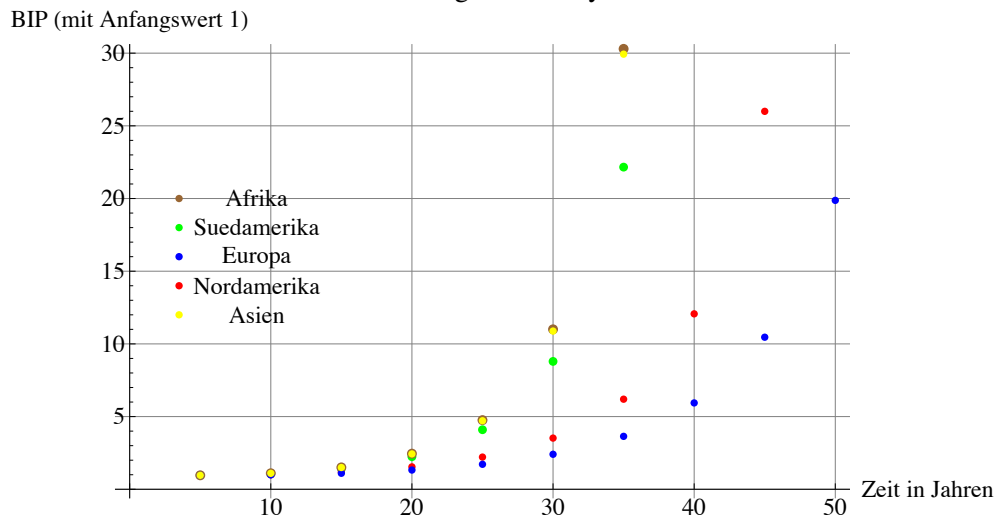
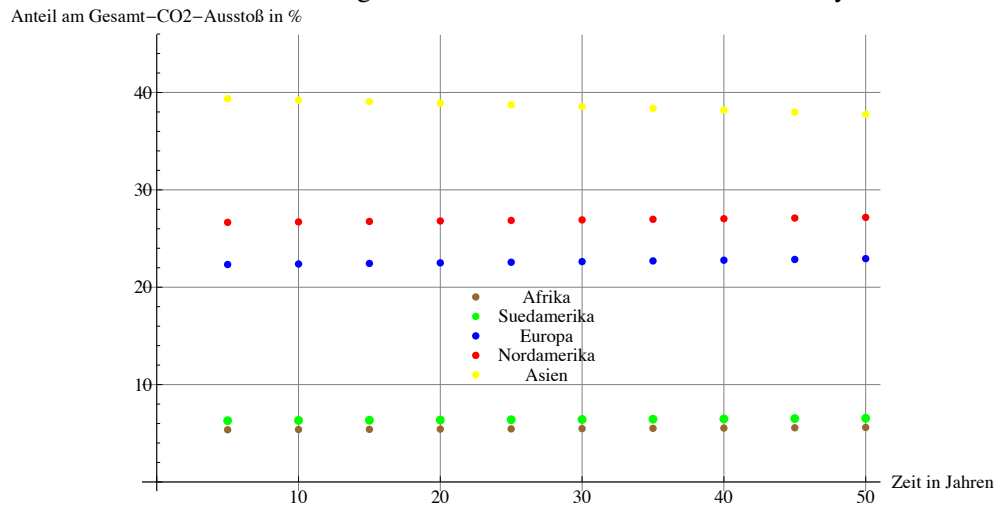


Abbildung 15: Anteil Länder am Gesamtausstoß dynamisch

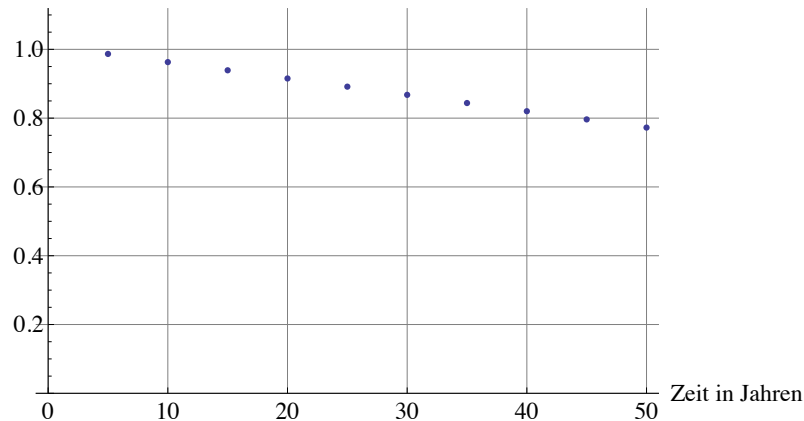


ausgestoßenen CO₂, während der der anderen Länder steigt, wie aus der dritten Grafik ersichtlich wird.

Zuletzt stellen wir noch die Veränderung des globalen CO₂-Ausstoßes in Prozent vom Anfangswert dar, die nötig ist, um das BIP aller Länder zu maximieren.

Aus Grafik 16 sieht man, dass der globale CO₂-Ausstoß mit zunehmender Zeit drastisch gesenkt werden muss, um die einzelnen BIPs der Akteure zu maximieren.

Abbildung 16: Globaler CO₂ Ausstoß fünfzig Jahre dynamisch
Globaler CO₂-Ausstoß in % vom Anfangswert



3.5 Schlussfolgerungen

Aufgrund unserer Ergebnisse können wir sagen, dass kein Weg daran vorbei führen wird, in den nächsten 15 Jahren damit zu beginnen, den CO₂-Ausstoß merklich zu verringern, um so längerfristige wirtschaftliche und klimatische Schäden in Industrie- wie auch in Entwicklungsgebieten zu vermeiden. Dies ist sogar in der nicht kooperativen Lösung und bei kurzfristiger Betrachtungsweise erforderlich. Wenn die Maximierung für einen Zeitpunkt in 50 Jahren erfolgt und zusätzlich die Wirkungen des Klimaschutzes auf alle Länder der Erde berücksichtigt werden, wäre eine deutlich höhere Reduktion erforderlich, die global auch zu deutlich höherem BIP führen würde. Dieses Ergebnis beschreibt das bekannte Problem des Gefangenendilemmas, bei dem jedes Land einen Anreiz hat, sich nicht kooperativ zu verhalten und ein zu schwaches Klimaziel zu verfolgen. Erst wenn es gelänge, ein globales Abkommen zu beschließen und somit alle Länder zu deutlichen Emissionsreduktionen zu verpflichten, könnte das globale Maximum erreicht werden.

Vor allem Industrieländer, und Länder mit großem Anteil am Gesamtausstoß des CO₂ wie beispielsweise die Region Südostasien, sollten sehr bald damit anfangen, einen Teil des BIPs für den Klimaschutz aufzuwenden. Andernfalls würden auf lange Sicht die durch den Klimawandel entstehenden Kosten überhand nehmen und vor allem (w)ärmere Länder in den Ruin treiben.

A Tabellen

Tabelle 6: Vergleich zwischen kooperativer und nicht-kooperativer Lösung bei Wohlfahrtsmaximierung

	Menge		Preis		Nettonutzen	
	kooperativ	nicht kooperativ	kooperativ	nicht kooperativ	nicht kooperativ	kooperativ
Europa	19	27	9	14	476	1407
Asien	25	29	8	11	694	1472
Afrika	13	29	1	10	694	1472
Australien	10	25	1	21	226	1283
Nordamerika	17	26	10	17	452	1358
Südamerika	14	29	2	11	303	1484
weltweit	100	166	5	14	2439	8493

Tabelle 7: Verteilung auf Ländergruppen zwischen kooperativer und nicht-kooperativer Lösung

	Menge in % von weltweit		Nettonutzen in % von weltweit	
	kooperativ	nicht kooperativ	kooperativ	nicht kooperativ
Europa	19	17	20	17
Asien	25	17	28	17
Afrika	13	18	12	18
Australien	10	15	9	15
Nordamerika	18	16	19	16
Südamerika	14	18	12	17
weltweit	100	100	100	100

Financial Contagion

Dipl.-Math. Carl Philip Trautmann
Moritz Hiebler, Valentin Haberl, Stefan Papst
Florian Lackner, Benjamin von Berg, Goda Mark

Zusammenfassung

Unsere Gruppe behandelte die Modellierung der Ausbreitung von Finanzkrisen in Wirtschaftssystemen. Dazu wurden zwei verschiedene Modelle erstellt, um einerseits die gegenseitige Ansteckung der Banken wie bei einem Virus und andererseits den Einfluss der Vernetzung untereinander auf die Ausdehnung einer Krise zu berechnen und darzustellen.

1 Modell der Ausbreitung

Der erste Teil der Gruppe modellierte die Finanzkrise in Form von einer Seuche, die sich wie eine Welle in Berücksichtigung von verschiedenen Parametern unterschiedlich ausbreitet. Ziel war es dieses Modell durch eine Differentialgleichung zu beschreiben und zu visualisieren. Der letzte Schritt der Arbeit war die Variation der Parameter, um das Verhalten des Modells kennen zu lernen.

1.1 Aufstellen einer chemischen Reaktion und Differentialgleichung

Die Seuche wird zunächst als chemische Reaktion modelliert.



X wird als die Anzahl an „gesunden“ Firmen und Y als die Anzahl an „kranken“ Firmen definiert. α entspricht hierbei der durchschnittlichen Anzahl an Y , mit denen ein X in Kontakt kommen muss, um selbst „krank“ zu werden. Da die Veränderung von X in einem gewissen Zeitabschnitt indirekt proportional zu den Grundkonzentrationen von X und Y ist, kann die Veränderungsrate von X (die immer < 0 ist) entsprechend der chemischen Reaktionskinetik folgendermaßen aufgeschrieben werden:

$$\frac{dX}{dt} = -\beta X \underbrace{Y \cdot Y \cdot Y \cdots Y}_{\alpha \text{ Mal}} = -\beta X Y^\alpha \quad (2)$$

β stellt die Reaktionsgeschwindigkeit dar.

Wir nehmen im Folgenden an, dass die Firmen nach einer gewissen Zeit alle Verbindungen zu anderen Firmen abbrechen und dadurch nicht mehr ansteckend wirken. Sei diese durchschnittliche Zeit, in der die Firmen überleben können, σ . Dann ergeben sich die folgenden zwei Differentialgleichungen, die beide abhängig von der Zeit t sind:

$$\begin{aligned} \dot{X} &= -\beta X Y^\alpha \\ \dot{Y} &= \beta X Y^\alpha - \frac{1}{\sigma} Y \end{aligned} \quad (3)$$

Die erste Differentialgleichung beschreibt die Abnahme der gesunden Population, wohingegen die zweite Gleichung die Veränderung der kranken Bevölkerung angibt. $\dot{X}$ ist immer negativ, da in der ersten Differentialgleichung X und $Y > 0$ sind. $\dot{Y}$ kann dagegen

positiv und negativ werden, je nach Startwert und den von der Wirtschaft ausgeschlossenen Firmen. Diese werden mit dem Subtrahenden $\frac{Y}{\sigma}$ ausgedrückt.

1.2 Funktion der Parameter

Unser Modell ist nun von 5 Parametern, nämlich α , β , σ , N und p , abhängig.

Als Gesamtanzahl der am Finanzgeschehen teilnehmenden Firmen wird N definiert, die sich zu Beginn nur aus „gesunden“ und „erkrankten“ zusammensetzt.

Als p wird der Anteil der erkrankten Firmen an N , also $\frac{Y}{N}$ definiert. Durch einen Preisschock kann p und somit auch Y erhöht werden.

α spiegelt auch die durchschnittliche Anzahl an Verbindungen wider, die sogenannte Diversifikation. Je höher nämlich α gewählt wird, desto mehr Verbindungen werden benötigt, um ein X erkranken zu lassen. Da α als Exponent von Y auch $\dot{Y}$ stark erhöht (siehe Graphik 5), sollte der Wert so groß wie möglich gehalten werden. Ist $\alpha < 1$, so bedarf es weniger als eines Y , um ein X zu infizieren. Damit erkranken unweigerlich alle Firmen an der Epidemie. In der Realität liegt α zwischen 1 und 10, in nicht extremen Wirtschaften ungefähr zwischen 1 und 2. Dieser Parameterwert ist kaum beeinflussbar.

β wird als Reaktionsgeschwindigkeit definiert und stellt in unserem Modell die durchschnittliche Umlaufgeschwindigkeit der Wirtschaft dar. Damit hängt β von der Größe der Ökonomie ab und ist von Land zu Land unterschiedlich. Es kann konkret kaum angegeben und wie α kaum beeinflusst werden.

σ , die durchschnittliche Zeitspanne, in der Y ansteckend ist, kann als einziger Parameter vom Staat gesteuert werden, da dieser entscheiden kann, ob er einer „erkrankten“ Firma schnell hilft oder nicht. Durch ein kleines σ nimmt Y schnell ab, da es weniger Zeit hat, andere anzustecken. Andererseits verursacht ein kleines σ auch hohe anfallende Kosten für den Staat. In dem erstellten Modell liegt der Wert von σ bei einer bis zu mehreren Wochen.

1.3 Erstellen einer Grenzliefenfunktion

Die Grenzliefenlinie, Frontierkurve genannt, wird durch $\dot{Y} = 0$ definiert und trennt den Bereich A, in dem $\dot{Y} > 0$, vom Bereich B, der durch $\dot{Y} < 0$ charakterisiert wird.

$$\begin{aligned} \beta XY^\alpha - \frac{1}{\sigma}Y &= 0 \Leftrightarrow Y \left(\beta XY^{\alpha-1} - \frac{1}{\sigma} \right) = 0 \\ Y &= \sqrt[\alpha-1]{\frac{1}{\beta\sigma X}} \quad \text{für } \alpha \neq 1 \end{aligned} \tag{4}$$

Ist $\alpha < 1$, ergibt sich eine Potenzfunktion, für $\alpha > 1$ ein positiver Hyperbelast. Im Spezialfall $\alpha = 1$ wird die Grenzlinie als Parallele zur Y -Achse dargestellt. Der Graph kann mithilfe des folgenden Quellcodes beschrieben werden:

```

1 function frontier(alpha,beta,sigma,n,t)
3 tau=t/n;
4 X = 0:tau:t;
6 if alpha==1
7     X=1/(beta*sigma)
8     Y=-10:0.001:10;
10 plot(X,Y)
11 axis([0 (4/beta/sigma/3) 0 1])
12 else
13     for i=1:(n+1)
14         Y(i) = (beta*sigma * X(i))^(1/(1-alpha));
15     end
16     plot(X,Y)
17 axis([0 1 0 1])
18 end

```

Man erhält somit einen „guten“ Bereich A , wo $\dot{Y} < 0$ und daher Y sinkt, und einen „schlechten“ B -Bereich, in dem $\dot{Y} > 0$. Hier steigt Y an.

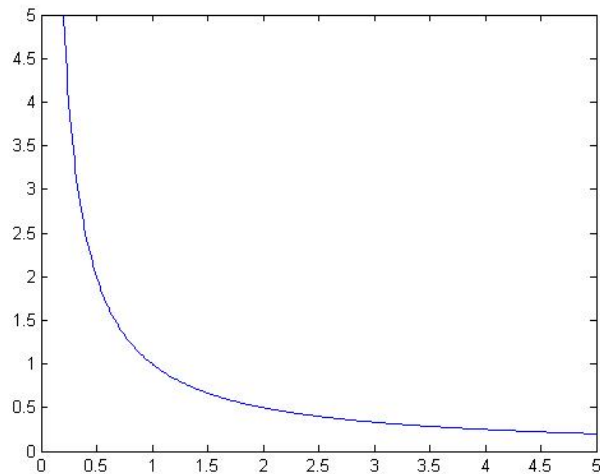


Abbildung 1: Grenzlinie - Frontier für $\alpha = 2$, $\beta = 0.0001$, $\sigma = 5$

1.4 Lösen der Differentialgleichungen

Zum Lösen des Differentialgleichungssystems wurde das explizite Euler-Verfahren verwendet. Dabei wird der Differentialquotient durch einen Differenzenquotienten angenähert und von dem letzten approximierten Wert auf den nächsten geschlossen.

In Formeln:

$$\begin{aligned}
 \frac{Y_k - Y_{k-1}}{\tau} &= \beta X_{k-1} Y_{k-1}^\alpha - \frac{1}{\sigma} Y_{k-1} \\
 \Leftrightarrow Y_k &= Y_{k-1} \left(\tau \beta X_{k-1} Y_{k-1}^{\alpha-1} + 1 - \frac{\tau}{\sigma} \right) \\
 \frac{X_k - X_{k-1}}{\tau} &= -\beta X_{k-1} Y_{k-1}^\alpha \\
 \Leftrightarrow X_k &= X_{k-1} (1 - \tau \beta Y_{k-1}^\alpha)
 \end{aligned} \tag{5}$$

Aus dieser Berechnung von X und Y folgen sofort die Funktionen für X und Y in Abhängigkeit von den Zeitschritten τ . Außerdem ergibt sich eine Kurve $X \mapsto Y$, die die Entwicklung der kranken Firmen Y in Abhängigkeit von den gesunden Firmen X widerspiegelt. Diese fällt insofern ins Auge, als dass sie von rechts nach links abgelesen werden muss, weil $\dot{X} < 0 \forall X, Y > 0$. Ferner kann als dritte Funktion der nicht mehr ansteckenden Firmen $g : t \mapsto (N - X - Y)$ aufgestellt werden.

```

1  function euler(n,t,N,p,alpha,beta,sigma)
3  tau=t/n;
5  X(1)=N*(1-p);
6  Y(1)=N*p;
8  for k=2:(n+1)
9      X(k)=X(k-1)*(1-beta*tau*Y(k-1)^alpha);
10     Y(k)=Y(k-1)*(1+beta*tau*X(k-1)*Y(k-1)^(alpha-1)-tau/sigma);
11 end
13 figure (1)
14 plot (T,X)
16 figure (2)
17 plot (T,Y)
19 figure (3)
20 plot (T,1-X-Y)
22 figure (4)

```


Nach einem rapiden Abstieg fängt sich die „gesunde“ Population (blau) wieder. Auf Grund von dem immer größer werdenden Teil $\frac{Y}{\sigma}$ besitzt die „kranke“ Bevölkerung (rot) einen Hochpunkt, fällt danach und nähert sich 0 an.

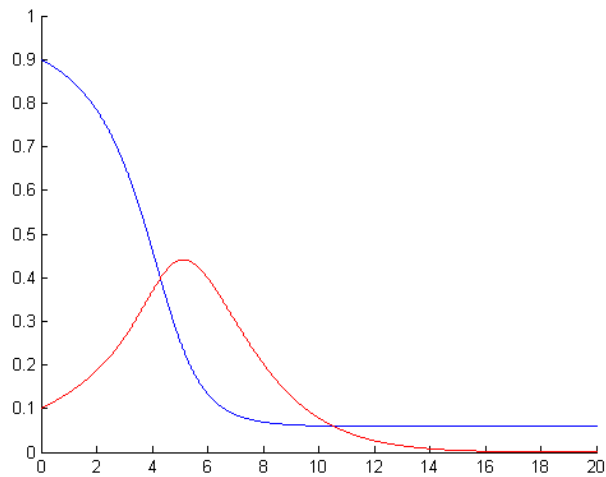


Abbildung 2: Populationsentwicklung nach der Zeit für $\alpha = 2, \beta = 0.0001, \sigma = 5$

1.5 Variation der Parameter

Durch Veränderung der p -Werte wird auch die Abbildung $X \mapsto Y$ variiert. Liegt Y (durch einen Preisschock) so hoch, dass es im B -Bereich startet, steigt Y streng monoton bis zum Schnittpunkt mit der Frontier-Kurve an, hat hier einen Hochpunkt und fällt dann streng monoton bis zur X -Achse ab.

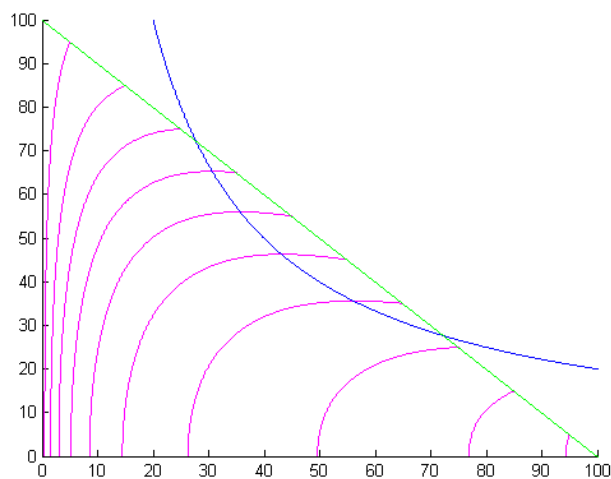


Abbildung 3: $p=0.05-0.95$ für $\alpha = 2, \beta = 0.0001, \sigma = 5$

Je größer N gewählt wird, umso anfälliger verhält sich auch das System, da sich die Frontierkurve mit größerem N umso mehr der Abszisse annähert. Damit verkleinert sich der Bereich A und der Startwert liegt leichter im „schlechten“ Bereich B .

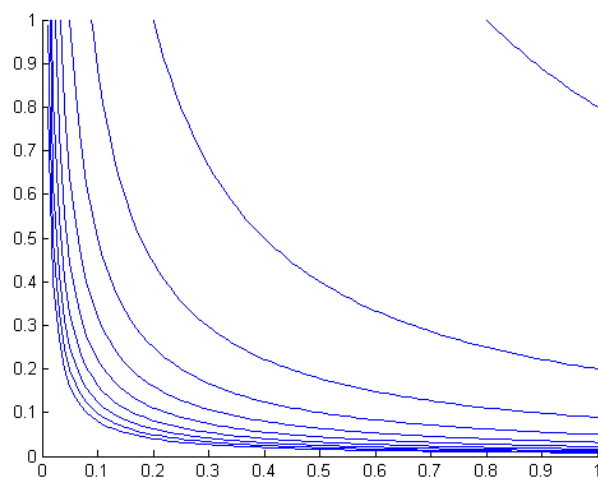


Abbildung 4: $N = 50 - 500$ für $\alpha = 2, \beta = 0.0001, \sigma = 5$

Mit der Variation von α nach oben verschiebt sich auch die Grenzkurve nach oben und damit erhöht sich auch die Wahrscheinlichkeit, dass der Startwert in dem „guten“ Bereich A liegt.

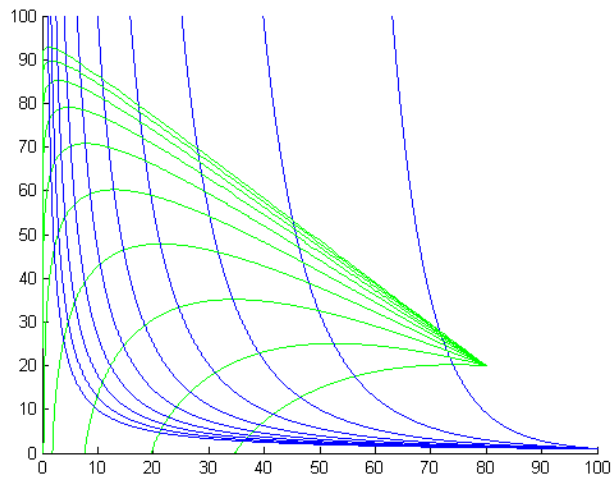


Abbildung 5: Grenzkurvenentwicklung mit Erhöhung von α , ($\beta = 0.0001, \sigma = 5$)

Im Gegensatz dazu flacht der Graph mit der Erhöhung von β mehr und mehr ab, das heißt die Wahrscheinlichkeit, den Startwert in dem „schlechten“ Bereich B anzutreffen, steigt.

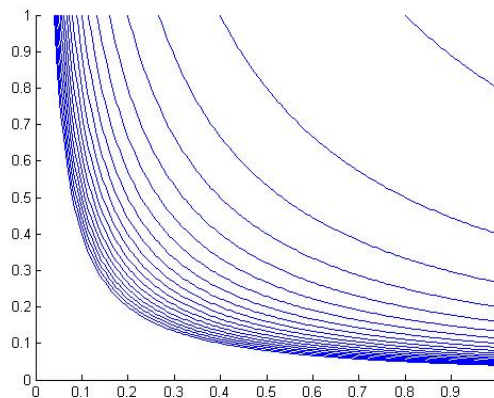


Abbildung 6: Abflachung der Grenzkurve durch Erhöhung von β

Mit der Erhöhung von σ nähert sich die Grenzlinie ähnlich zu β schnell der X -Achse an und vergrößert somit den „schlechten“ Bereich B .

1.6 Resümee

Die Financial Contagion kann durch ein Differentialgleichungssystem modelliert werden. Durch Veränderung der Parameter kann der Verlauf der Ausbreitung variiert werden.

Der Staat kann dabei auf den Parameter σ zurückgreifen, der den Verlauf der Krise stark beeinflusst. Auf α und β hat der Staat wenig Einflussmöglichkeiten. Da die Grundgesamtheit N meistens einen hohen Wert aufweist, ist das System anfällig für finanzielle Probleme. Aufgrund dessen sollte der Staat möglichst früh in das Geschehen eingreifen, indem σ ausreichend klein gewählt wird. Dies verursacht zwar hohe Kosten, jedoch kann dadurch die Krise eingedämmt werden.

2 Modell der Vernetzung

Der zweite Teil der Gruppe befasste sich mit der Ausdehnung der Finanzkrise auf Grund der weltweiten Vernetzung. Ziel dieses Modells ist es, die Ausbreitung in Abhängigkeit der Vernetzung der Banken untereinander zu betrachten und eine mögliche optimale Anzahl an Anteilen einer Bank an anderen Banken zu finden. Welches Netzwerk ist am stabilsten und welche Parameter beeinflussen die Stabilität des Netzwerkes?

Die Idee ist, eine Finanzkrise zu simulieren und Möglichkeiten zu finden sie zu verhindern. Immer wieder erscheinen Stresstests, die bei einem möglichen Preisschock Banken und Firmen auf ihre Stabilität prüfen. Von diesem Grundgedanken geleitet, wollen wir Firmennetzwerke aufstellen und einen Börsenkrach simulieren, um optimale Verbindungsstrukturen zu ermitteln. Von einem einfachen Modell ausgehend versuchen wir uns zu immer komplexeren Modellen heranzuarbeiten, damit die Simulation möglichst realitätsgetreu wird.

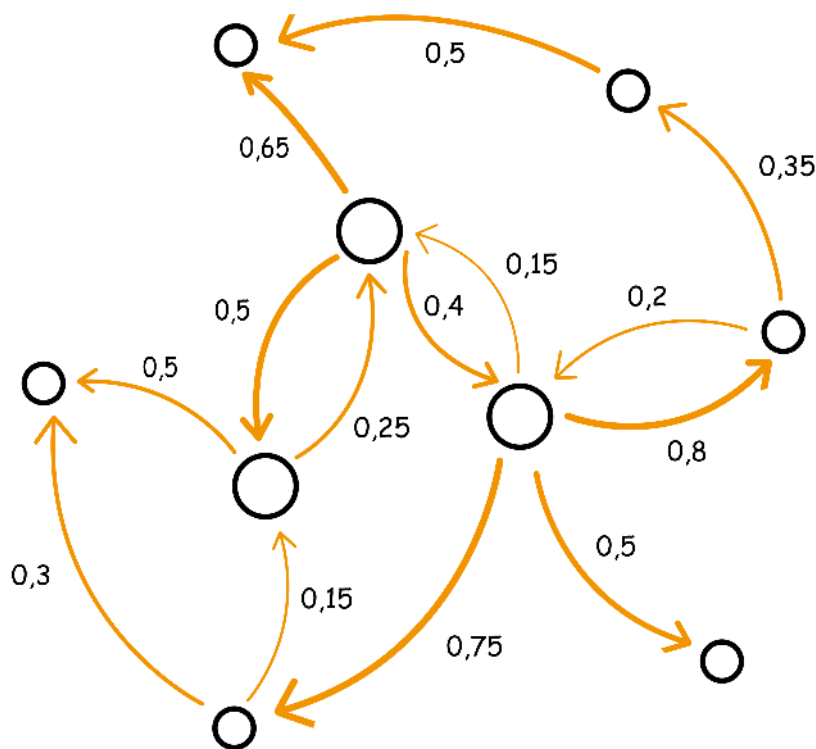


Abbildung 7: Crossholdings bei mehreren Firmen

2.1 Berechnung des Marktwertes der Firmen

Wir bestimmen den Wert einer Firma oder einer Bank einerseits anhand ihrer Assets, also der Investments in Immobilien, Autos, Flugzeuge, etc. und andererseits anhand der Anteile bei anderen Organisationen, den so genannten Crossholdings. Dabei nehmen wir an, dass eine Firma nicht sich selbst gehört, sondern anderen Firmen im Netzwerk und privaten Investoren gehört. Der Asset-Anteil wird von den Asset-Preisen bestimmt während die Crossholdings von den unterschiedlichen und schwankenden Marktwerten der anderen Firmen abhängen.

In unserem Modell sind n Firmen vorhanden, die über m Assets verfügen und des Weiteren untereinander gegenseitige Anteile haben. Der Rest der Anteile wird von Privatpersonen, den so genannten „Outsideholders“ getragen. Mit der Matrix C stellen wir die Crossholdings dar. Der Eintrag C_{ij} beschreibt den Anteil von Firma i an Firma j in Prozent. Die Matrix D beinhaltet die Mengen, die eine Firma an den einzelnen Assets besitzt.

$$V_i = \sum_{k=1}^m D_{ik} p_k + \sum_{j=1}^n C_{ij} V_j \quad (6)$$

Der Wert jeder Firma i ist die Summe der eigenen Assets gewichtet mit den Asset-Preisen p und der Summe der Anteile am Wert der anderen Firmen V . Daher ergibt sich folgende Gleichung:

$$V = Dp + CV \quad (7)$$

Da diese Gleichung jedoch ein implizites System darstellt, müssen wir zur Berechnung des Firmenwertes V ein lineares Gleichungssystem aus Matrizen lösen:

$$\begin{aligned} V &= Dp + CV \\ V - CV &= Dp \\ (I - C)V &= Dp \\ V &= (I - C)^{-1} Dp \end{aligned} \quad (8)$$

Aus dieser Berechnung ergibt sich jedoch ein weiteres Problem: In den Wert der anderen Firmen sind möglicherweise auch die eigenen Assets mit einberechnet, wenn die andere Firma einen gewissen Anteil an der eigenen Firma besitzt. Dies würde zu einem fiktiven Wert jeder Firma führen, der überhöht ist.

Um dieses Problem zu lösen, präzisieren wir unsere Formel und fügen der Berechnung einen weiteren Wert hinzu, der die Anteile an Firma i , die von fremden Firmen gehalten werden, abzieht.

$$V_i = \sum_{k=1}^m D_{ik} p_k + \sum_{j=1}^n C_{ij} V_j - \sum_{j=1}^m C_{ji} V_i \quad (9)$$

Des Weiteren führen wir die Matrizen $\tilde{C}$ und $\hat{C}$ ein, wobei die Diagonalelemente von $\tilde{C}$ als die Summe der Crossholdings von fremden Firmen definiert wird.

$$\begin{aligned} \tilde{C}_{ii} &= \sum_{j=1}^n C_{ji} \\ \hat{C} &= I - \tilde{C} \end{aligned} \quad (10)$$

Die Matrix $\hat{C}$ enthält auf der Hauptdiagonale jetzt nur mehr die Anteile der privaten Shareholder. Da die Gleichung (8) nur fiktive Werte enthält, benennen wir den hypothetischen Firmenwert V auf $\tilde{V}$ um. Nun können wir mit Hilfe der aufgestellten Gleichung (9) die realen Marktwerte der einzelnen Firmen ermitteln.

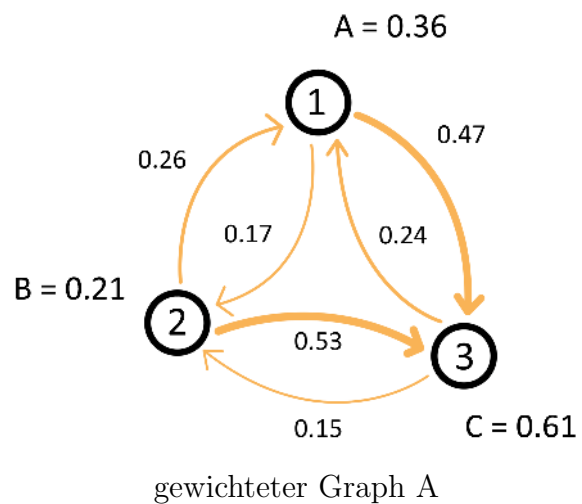
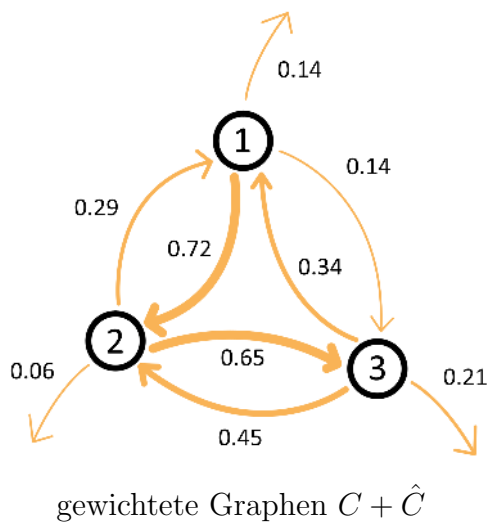
$$\begin{aligned} \tilde{V} &= (I - C)^{-1} Dp \\ V &= C\tilde{V} - \tilde{C}\tilde{V} + Dp \\ V &= (C - \tilde{C})\tilde{V} + Dp \\ V &= (C - \tilde{C})(I - C)^{-1} Dp + Dp \\ V &= ((C - \tilde{C})(I - C)^{-1} + I) Dp \\ V &= ((C - \tilde{C}) + (I - C))(I - C)^{-1} Dp \\ V &= (I - \tilde{C} + C - C)(I - C)^{-1} Dp \\ V &= \hat{C}(I - C)^{-1} Dp \\ V &= ADp \end{aligned} \quad (11)$$

Durch die Substitution von $\hat{C}(I - C)^{-1}$ durch A erhalten wir eine Matrix, welche die Abhängigkeit der Marktwerte einer Firma i von den Assets einer Firma j verkörpert.

Diese Abhängigkeit kann durch direkte oder indirekte Crossholdings, in Form von Kettenverbindungen entstehen. Wenn nun ein Asset vollkommen abstürzt und seinen Wert verliert, können dadurch nicht nur Firmen, die dieses Asset besitzen, sondern auch alle Organisationen, die indirekt an diesem Asset über Crossholdings beteiligt sind, bankrott gehen. Diese Abhängigkeit soll die Matrix abbilden.

$$C = \begin{pmatrix} 0.00 & 0.29 & 0.34 \\ 0.72 & 0.00 & 0.45 \\ 0.14 & 0.65 & 0.00 \end{pmatrix}$$

$$A = (I - C)^{-1} = \begin{pmatrix} 0.36 & 0.26 & 0.24 \\ 0.17 & 0.21 & 0.15 \\ 0.47 & 0.53 & 0.61 \end{pmatrix}$$



Die oben abgebildeten Grafiken verdeutlichen den Unterschied zwischen den zwei gewichteten Graphen, die in C und A beschrieben sind. Die Matrix $C + \hat{C}$ setzt sich aus den einzelnen Crossholdings zwischen den Firmen sowie den Outsideholdings der privaten Investoren zusammen. Bei der Matrix A hingegen ist die gegenseitige Vernetzung miteinbezogen bzw. ausgerechnet; so hält die Firma 1 durch die Anteile an der Firma 2 auch gewisse Crossholdings an der Firma 3. Dementsprechend wirkt sich eine Insolvenz einer der Firmen auf alle anderen Firmen aus, auch wenn sie keine direkten Kontakte hätten. In den veränderten Zahlenwerten wird sichtbar, wie sich die Relationen zwischen den einzelnen Organisationen durch die komplexen Verknüpfungen ändert und wie sich die Weitervernetzung auswirkt. Umgelegt auf das weltweite Wirtschaftssystem würde die Matrix A die Abhängigkeit einer einzelnen Bank vom globalen Banken- und Firmennetzwerk demonstrieren. Eine einzige Verbindung einer kleinen Bank zu einer internationalen Großbanken führt folglich auch zu einer Vernetzung mit dem gesamten System.

Die Umsetzung der Theorie mit einem MATLAB-Programm war nach der theoretischen Behandlung die zweite große Hürde. Der Einfachheit halber nahmen wir für die Berechnungen an, dass jede Firma nur ein Asset mit dem Wert 1 besitzt. Demnach ist der Vektor p , der die Preise beschreibt, der Einsvektor. In einer Matrix G generieren wir zufällige Verbindungen zwischen den Firmen, ein randomisiertes Netzwerk, wobei wir die durchschnittliche Anzahl eines Knotens mit der entsprechenden Konfigurationsvariable „average_connections“ vorgeben. Dies wird mit Hilfe der BERNOULLI-VERTEILUNG realisiert, indem wir jedes Matrix-Element als BERNOULLI-verteilte Zufallszahl setzen.

```

1 % Asset-Preise :
2 p = ones(firmen_anz,1);

4 %Definition von G (Verbindungen zwischen Organisationen, ungewichteter Graph)
5 G = zeros(firmen_anz);

7 % weise 0 fuer keine Verbindung und 1 fuer Verbindung zu einer anderen Firma
8 G = binornd(1,average_connections/(firmen_anz-1), firmen_anz, firmen_anz);
9 % In der Diagonale der Matrix darf nur 0 stehen
10 G(eye(firmen_anz)~=0) = 0;

```

Der nächste Schritt ist es, aus der ungewichteten Matrix G , in der die Verbindungen der Firmen untereinander durch die Werte 0 und 1 gekennzeichnet sind, in eine gewichtete Matrix C zu verwandeln. Durch eine weitere Variable „max_crossholds“ kann festgelegt werden, welcher Anteil der Firma auf Privatinvestoren, den so genannten Outsideholdings, und welcher Anteil auf andere Firmen aufgeteilt wird. Für die Simulation der Finanzkrise sind eigentlich nur die Anteile von anderen Firmen entscheidend, da diese durch den Konkurs der Firma auch Pleite gehen können. Daher wird der Anteil, der im Netzwerk verbleibt, gleichmäßig auf die Verbindung zu anderen Organisationen verteilt. Diese Werte werden in die Matrix C eingeschrieben und demnach ergibt die Summe jeder Spalte immer genau die Variable „max_crossholds“.

```

1 % connetions_anz zeigt wie viele Firmen in die Firma i investiert haben
2 % Durchschnitt der Anteile anderer Firmen an einer Firma berechnen
3 C=zeros(firmen_anz);

5 for i=1:firmen_anz
6     for j=1:firmen_anz
7         if connections_anz(i) == 0
8             C(j,i) =0;
9         else
10            C(j,i) = C(j,i)*(max_crossholds/connections_anz(i));
11        end
12    end
13 end

```

Der gewichtete Graph enthält nun in der Zeile i die Anteile der Firma i an den anderen Firmen in Prozent, wobei der Anteil an der eigenen Firma auf 0 gesetzt wurde, wodurch die Diagonale der Matrix C nur aus 0 besteht. Die Verbindungen zeigen die Abhängigkeit der anderen Organisation an der Firma durch ihre Anteile. Geht die Firma bankrott, so müssen alle Organisation, die Anteil daran haben, für die entstandenen Kosten haften, wodurch sie ebenfalls in Konkurs geraten können.

Die eigentliche Berechnung erfolgte im nächsten Schritt. Dazu setzten wir die berechneten Matrizen in die umgeformte Formel ein, um tatsächliche Werte zu erhalten.

```

1 % Berechnung von C_
2 summe_jeder_spalte = ones(1,firmen_anz)*C;
3 komplementaer_summen = ones(1,firmen_anz)-summe_jeder_spalte;
4 C_ = diag(komplementaer_summen);

6 % Berechnung des Wertes der Firmen
7 Z = C_ \ (eye(firmen_anz)-C);
8 V = Z \ p;

```

2.2 Preisschock

Nach der Berechnung der Marktwerte der unterschiedlichen Firmen, gilt es als nächstes einen Asset-Crash zu simulieren und daraus folgende Pleiten zu errechnen. Wir nehmen dazu an, dass zunächst ein Asset-Preis 0 wird und versuchen dadurch die Auswirkungen auf die Firmenwerte zu ermitteln. Wann eine Organisation Konkurs geht, wird dabei von einem Grenzwert entschieden, ab dem die Firma sich selbst nicht mehr erhalten kann. Bestimmt wird dieser Konkurswert, indem der ursprüngliche Firmenwert noch vor dem Preisschock mit einer bestimmten variablen Größe „threshold“ multipliziert wird. Wenn nun der Firmenwert nach einem Preissturz eines Assets unter den Grenzwert gelangt, geht diese pleite.

Da wir in unserem vereinfachten Modell jeder Firma nur ein einziges Asset mit dem Wert 1 zugewiesen haben, geht durch den Preissturz eines einzigen Assets bereits eine Firma pleite und verursacht einen Schaden. In der Realität entsprechen diese Kosten den finanziellen Maßnahmen wie Fabriksauflösungen, Kündigungen, etc. Bei einer Fluggesellschaft muss zum Beispiel ein gewisses Budget vorhanden sein, damit Flugzeuge gekauft werden können und das Personal bezahlt werden kann. Füllt die Firma unter einen gewissen Wert können diese Kosten nicht mehr abgesichert werden und die Firma kann nicht mehr wirtschaften, sie geht pleite. Die entstandenen Kosten sollen dabei von

den Investoren, also den privaten Unterstützern und den fremden Firmen mit Crossholdings am Pleitegänger, entsprechend der Abhängigkeitsmatrix A aufgeteilt und bezahlt werden.

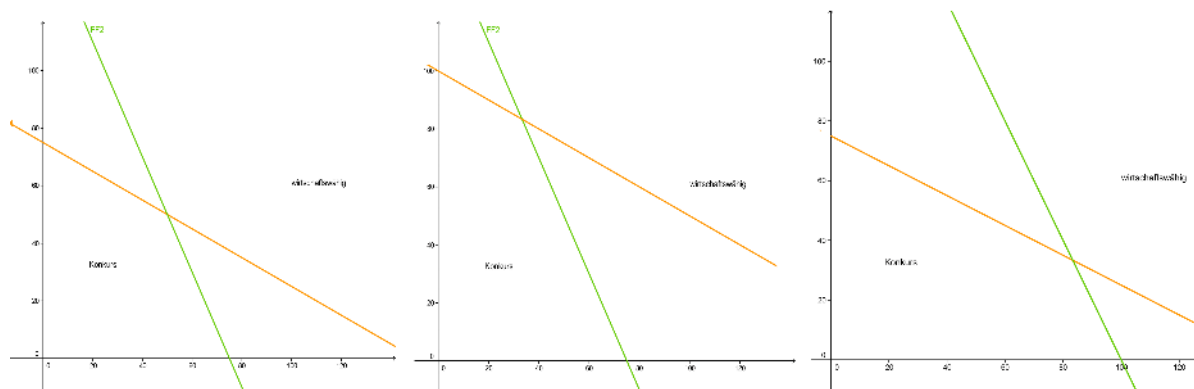


Abbildung 8: Konkursgrenzen bei zwei Firmen

Die Grafik verdeutlicht die Auswirkungen der Schulden nach einem Bankrott einer Firma auf die andere anteilhabende Firma. Die zwei durchgezogenen Geraden beschreiben die eigentliche Konkursgrenze, wenn sonst keine zusätzlichen Kosten entstehen. Links von der Linie ist jeweils der Konkursbereich und rechts davon der Bereich, in dem die Firma wirtschaften kann. Durch die Insolvenz der einen Organisation A müssen die entstandenen Kosten von der anderen Firma B getragen werden. Dadurch sinkt des Weiteren auch der Marktwert von B . Durch den kleineren Firmenwert ist die Firma auch näher zur orangenen Konkursgrenze. Sie ist durch die Kosten anfälliger für den Bankrott geworden. Analog gilt diese Folgerung auch für die Firma A , wenn die Organisation B zugrunde geht. Auch ihr Marktwert wird durch die Anteile an der verschuldeten Firma geringer und sie kommt der Pleite näher, was durch die Verschiebung der grünen Linie deutlich wird. Der Konkurswert selber ändert sich nicht, die Grafik soll nur durch die Verschiebung der Grenze die größere Anfälligkeit zur Pleite demonstrieren.

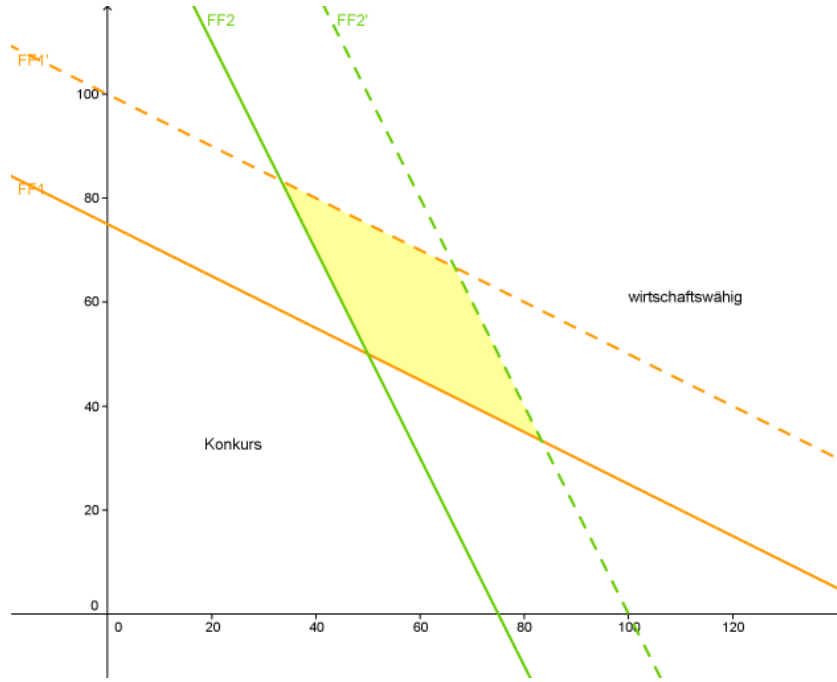


Abbildung 9: Konkursbereich

Der gelbe Bereich in der Darstellung zeigt einen sensiblen Bereich. Wenn sich nämlich die Marktwerte beider Firmen A und B in diesem Segment befinden und eine Firma plötzlich durch einen Preisschock pleite geht, wird automatisch auch die zweite Firma in den finanziellen Ruin mitgerissen, weil die Konkursgrenze auf jeden Fall über den Marktwert springt.

Um den entstandenen Schaden zu berechnen, wird ein neuer Vektor β eingeführt, welcher die Differenz zwischen dem ursprünglichen Firmenwert und der Konkursgrenze darstellt, was den Schulden entspricht. Beahlt werden diese von den Firmen mit Anteilen an der pleite gegangenen Firma, indem man β in die Wertberechnung miteinbezieht. Um Schulden nur bei zahlungsunfähigen Firmen zu erfassen, nehmen wir einen Indikator Ind an, der den Werten 0 annimmt, wenn der derzeitige Marktwert über der Grenze und 1 wird, wenn sie darunter liegt.

$$V_i = \sum_{j=1}^n C_{ji} V_j + \sum_{k=1}^n D_{ik} p_k - \beta_i Ind \quad (12)$$

Vereinfacht auf unsere Matrizenrechnung bedeutet dies:

$$\begin{aligned}
\tilde{V} &= (I - C)^{-1}(Dp - \beta(V)) \\
V &= \hat{C}(I - C)^{-1}(Dp - \beta(V)) \\
V &= A(Dp - \beta(V))
\end{aligned} \tag{13}$$

In Folge des plötzlichen Preisschocks werden auch alle Marktwerte der Firmen, die an der bankrotten Firma Anteil haben, ebenfalls kleiner was erneute Konkurse auslösen kann. Die Anteile an fremden Firmen, also die Crossholdings, stehen hier wiederum in der Matrix A - sie verdeutlicht die Abhängigkeit einer Organisation von den Assets anderer Organisationen. Je nach Vernetzung sind die Organisationen untereinander so gekoppelt, dass alle Firmenwerte von einem einzigen Preissturz eines einzigen Assets und den durch die Konkurse entstandenen Schulden fallen können. Die ganze Krise pendelt sich ein, wenn keine Firma mehr konkurs geht und dadurch keine Kosten mehr entstehen.

Im Programm wird zunächst der neue Firmenwert berechnet und mit dem Grenzwert, der durch das Multiplizieren des ursprünglichen Firmenwerts mit dem „threshold“ erzeugt wird, verglichen. Eine If-Abfrage ersetzt dabei die Indikatormatrix *Ind*. Ist eine Firma bankrott, wird der Schaden in den Vektor *beta* eingetragen. Im nächsten Schritt werden die Marktwerte der Firmen neu berechnet, wobei der β -Vektor mit den entstandenen Kosten von den Assetspreisen abgezogen wird. Ausgelöst wird der Preisschock durch das Nullsetzen von zufälligen Asset-Preisen; die Anzahl ist durch die Variable „percent_of_failures“ regelbar.

```

1 % zufälliger Aktiencrash
2 p(1:floor(percent_of_failures*firmen_anz)) = 0;

4 % Neu-Berechnung des Firmenwerts
5 V = Z\(p-beta)

7 % Finden der Firmen die Pleite gingen
8 % beta => wenn 0 = nicht pleite, sonst zu bezahlender Betrag
9 for i=1:firmen_anz
10     if (V(i) < V_alt(i)*threshold)
11         beta(i) = (V_alt(i)*threshold)-V(i);
12     end
13 end

```

Wenn eine Firma insolvent wird, verringern sich dadurch auch andere Firmenwerte, wodurch die einzelnen Marktwerte neu berechnet werden müssen. Daher läuft der oben beschriebene Vorgang mehrmals ab, bis keine Organisation mehr pleite geht. Die Berechnung des Preisschocks wird deswegen in eine große While-Schleife verpackt. Dabei muss in jeder Runde mitgezählt werden, wie viele Firmen unter den Grenzwert gekom-

men sind, damit man für die Abbruchbedingung fest stellen kann, ab welchem Zeitpunkt keine Firma mehr in Konkurs gegangen ist. Dies erfolgt in unserem Programm mit der Variable „pleite“, die nach jedem Durchlauf auf „pleite_old“ gespeichert wird und dann als Schleifenbedingung mit dem neuen Wert verglichen wird.

```

1 % zufälliger Aktiencrash
2 p(1:floor(percent_of_failures*firmen_anz)) = 0;

4 % Anzahl der Pleiten deklarieren
5 pleite = 0;
6 pleite_old = 1;

8 while pleite ~= pleite_old

10 % Berechnung des Wertes der Firmen
11 V=Z\(p-beta);

13 % Finden der Firmen die Pleite gingen
14 % beta => wenn 0 = nicht pleite, sonst Schulden
15 for i=1:firmen_anz
16     if (V(i) < V_alt(i)*threshold)
17         beta(i) = (V_alt(i)*threshold)-V(i);
18     end
19 end

21 % Berechnung der Firmen, die bankrott sind
22 pleite_old = pleite;
23 pleite = 0;

25 for i=1:firmen_anz
26     if beta(i) ~= 0
27         pleite = pleite +1;
28     end
29 end
30 end

```

Anhand dieser Schleife können wir nun eine Finanzkrise simulieren und durch das Verändern der Variablen „firmen_anz“, „average_connections“ und „percent_of_failures“ verschiedene Situationen darstellen. Durch das Untersuchen des Netzwerks unter gewissen Bedingungen und das Auswerten der Resultate, ist es so möglich, ein optimiertes Firmennetzwerk zu entwickeln.

2.3 Visualisierung

Für die Interpretation ist es sehr wichtig die Ergebnisse zu visualisieren. Zunächst entwickelten wir ein Modell, das den Verlauf einer Krise in Echtzeit darstellen kann. Dabei

war wunderbar ersichtlich in welcher Art und vor allem wie schnell sich die Schulden ausbreiteten und immer mehr Firmen in den Abgrund rissen. Eine weitreichende Analyse der Auswirkung der Änderung vieler Parameter in Bezug zueinander kann aber mit diesem Aufbau nicht getätigt werden.

Hierfür programmieren wir ein kleines Script, das einen oder mehrere Parameter änderte, das Simulationsprogramm jede Parameterkombination x -mal rechnen ließ und die Mittelwerte als Graph darstellte. So konnte die Simulation noch besser benutzt werden um unser Finanzsystem in Grundzügen zu testen, die Ergebnisse zu interpretieren und die gewonnen Erkenntnisse festzuhalten.

2.4 Ein Beispiel

Zur Veranschaulichung der Theorie und des Programms wird hier im Folgenden ein konkretes Beispiel mit gegebenen Werten angeführt. Der Einfachheit halber wird die Firmenanzahl n auf 5 beschränkt und der maximale Anteil von Firmen an anderen Firmen auf 0.7 festgelegt.

- Firmenanzahl n
firmen_anz = 5
- Durchschnittliche Verbindungen
average_connections = 2
- Maximaler Anteil der Firmen an anderen Firmen
max_crossholds = 0.7

$$C = \begin{pmatrix} 0.00 & 0.00 & 0.175 & 0.35 & 0.00 \\ 0.00 & 0.00 & 0.175 & 0.00 & 0.00 \\ 0.00 & 0.00 & 0.00 & 0.35 & 0.00 \\ 0.00 & 0.00 & 0.175 & 0.00 & 0.70 \\ 0.70 & 0.70 & 0.175 & 0.00 & 0.00 \end{pmatrix}$$

$$A = \begin{pmatrix} 0.3910 & 0.0910 & 0.1396 & 0.1857 & 0.1300 \\ 0.0136 & 0.3136 & 0.0655 & 0.0277 & 0.0194 \\ 0.0774 & 0.0774 & 0.3741 & 0.1580 & 0.1106 \\ 0.2213 & 0.2213 & 0.2118 & 0.4516 & 0.3161 \\ 0.2967 & 0.2967 & 0.2090 & 0.1770 & 0.4239 \end{pmatrix} \quad V = \begin{pmatrix} 0.9373 \\ 0.4396 \\ 0.7977 \\ 1.4220 \\ 1.4034 \end{pmatrix}$$

2.5 Vernetzung & Bankrotts

Nach der theoretischen Behandlung des Problems und dem Programmieren gelangen wir zu unserem eigentlichen Vorhaben: Die Simulation von unterschiedlichen Finanzkrisen. Unser Augenmerk liegt auf den Vernetzungen der Firmen untereinander. Wie stark sollten Banken vernetzt sein, damit sie eine Krise gut überstehen können? Wann ist das Bankensystem am instabilsten und am meisten anfällig? Welche Parameter beeinflussen die Ausbreitung der Krise?

Welche grundlegenden Bedingungen ermöglichen die Ausbreitung einer Krise?

- Preisschock (First Failure)
Eine gewisse Anzahl an Asset-Preisen muss als Auslöser der Krise plötzlich tief genug fallen.
- Ansteckung (Contagion)
Genügend Organisationen müssen für die Krise anfällig sein, damit sie durch den Preisschock überhaupt bankrott gehen können.
- Vernetzung (Interconnections)
Die einzelnen Organisationen müssen durch Crossholdings untereinander vernetzt sein, damit sich die Krise großflächig ausbreiten kann.

Auf diesen Grundlagen arbeiteten wir uns von einfachen Grundmodellen bis zu komplexen Simulationen. Durch das Verändern der Grundvariablen und das Hinzufügen von neuen Funktionen, erzielten wir durchaus realistische Ergebnisse und folgerten auf brauchbare Erkenntnisse. Unsere MATLAB- und C#-Programme entwickelten sich stetig weiter, vor allem starke Performance-Optimierungen waren notwendig um oft genug rechnen zu können und somit die Genauigkeit der Graphen auf ein akzeptables Niveau zu heben.

2.5.1 P2P-Modell

Das Programm simuliert ein Netzwerk aus n Firmen, die dargestellt durch Punkte in einem großen Kreis angeordnet werden, wobei kein Unterschied zwischen den einzelnen Firmen im Kreis besteht. Die Verbindungen sind dabei zufällig, jedoch kann die Anzahl der durchschnittlichen Verbindungen mit der Variable „average_connections“ bewusst verändert werden.

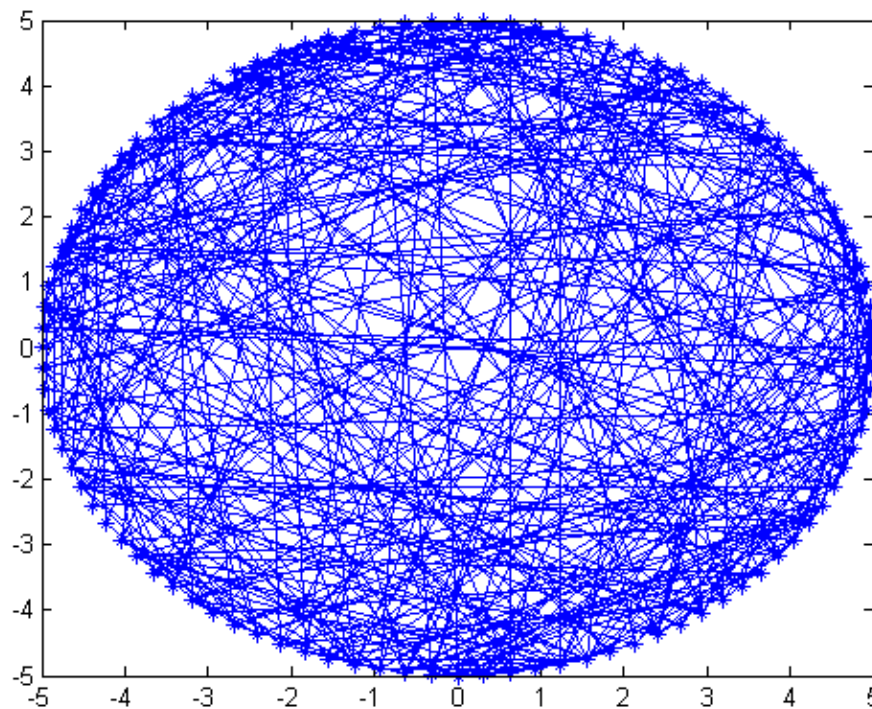


Abbildung 10: Ein gesundes Firmennetzwerk

Die vielen Connections der Firmen untereinander weisen auf die Crossholdings, den gegenseitigen Anteilen, hin. Bei dieser graphischen Darstellung wurde eine durchschnittliche Verbindungsanzahl von 6 gewählt, wodurch ein recht vernetztes Gebilde entsteht. Dieses komplexe Gefüge stellt ein gesundes System dar, welches in sich ohne Probleme funktioniert.

Im nächsten Schritt soll nun der Preisschock simuliert und die Auswirkungen auf das System betrachtet werden. Dazu wird der Wert der Assets von sieben Firmen auf 0 gesetzt, was diese Firmen in den Konkurs treibt und Kosten verursacht, welche weitere Firmen in den Ruin stürzt. Aus dem anfänglich relativ kleinen Verlust entwickelt sich eine Sturmflut der Pleitegänge. Obwohl sie keine direkten Verbindungen zu den bankrotten Firmen hatten, verlieren viele Firmen trotzdem an Wert oder gehen pleite, nur

weil sie Teil des Netzwerks sind. Die vielen Crossholdings reichen aus, um über mehrere Organisationen trotzdem mit irgendeiner zahlungsunfähigen Firma verbunden zu sein.

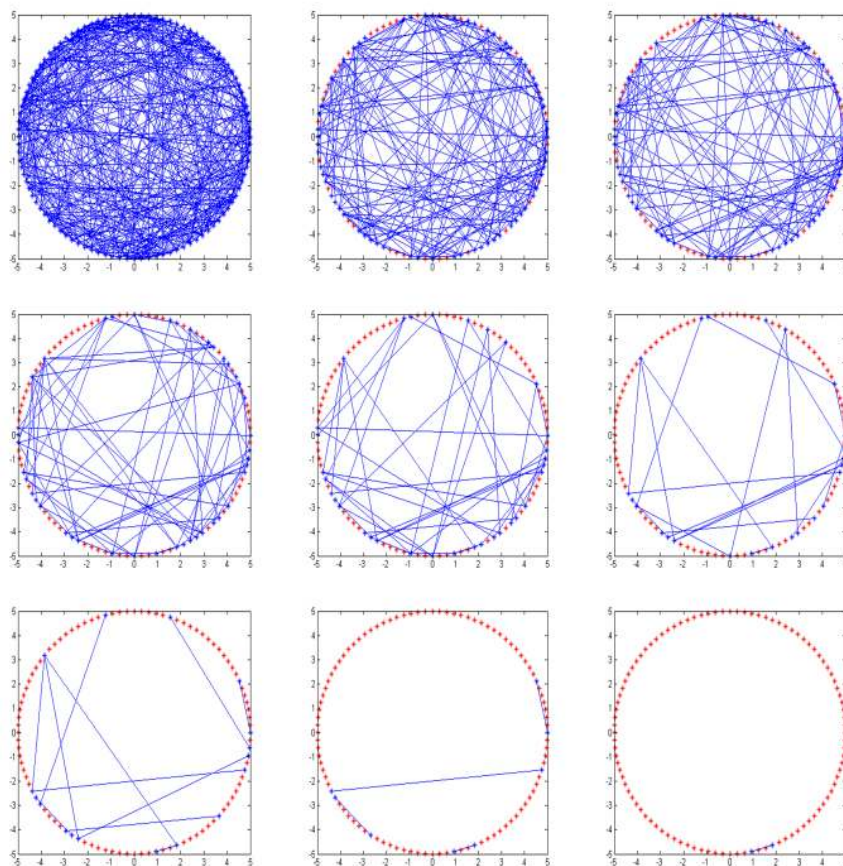


Abbildung 11: Verlauf einer simulierten Krise

Es überleben nur zwei Firmen, die miteinander verbunden sind. Obwohl andere Firmen an ihnen Aktien besaßen, sind sie übrig geblieben, weil sie keine Anteile an den anderen Firmen hatten. Durch Zufall bildeten sie dadurch ein autarkes System, das von der Krise nicht getroffen wurde.

Untersuchung der Pleiten und durchschnittlichen Verbindungsanzahl Mit Hilfe von Schleifen untersuchen wir die Anzahl der bankrotten Firmen in Abhängigkeit der durchschnittlichen Verbindungsanzahl. So kann eine erste wichtige Frage ,verdeutlicht in einem Graphen, beantwortet werden, nämlich bei wie vielen durchschnittlichen Verbindungen ein Wirtschaftssystem am anfälligsten bzw. stabilsten ist. Die ganze Simulation erfolgt in mehreren Runden, um auch bei willkürlichen Zufallszahlen durchschnittliche Werte

zu erhalten.

Benutzte Parameter:

- 100 Firmen
- Threshold = 93%
- Rundenanzahl pro Wert = 100
- Gecrashte Assets = 1

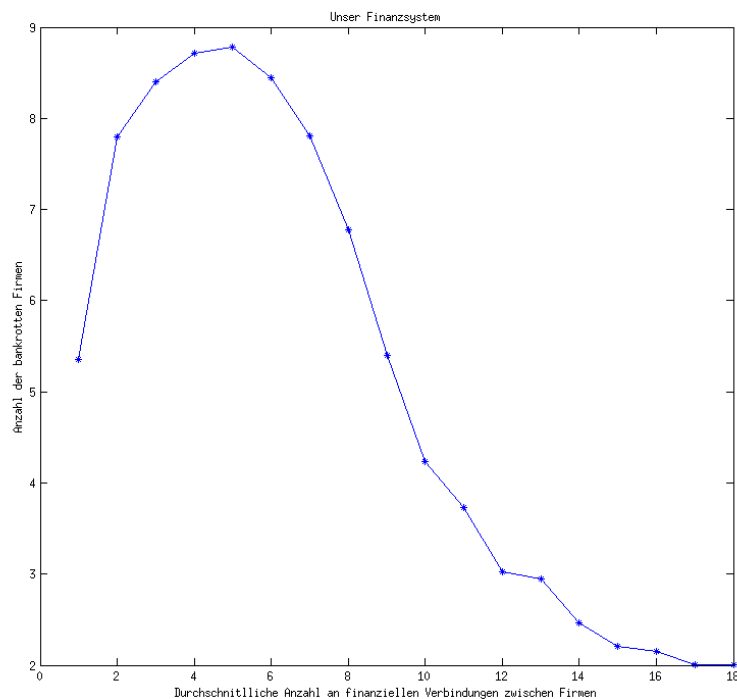


Abbildung 12: bankrotte Firmen in Abhängigkeit der Crossholdings

Am Anfang ist die Verbindungsdichte im Netzwerk so gering, dass nicht viele Firmen in den finanziellen Ruin mitgerissen werden. Die Kurve steigt bis zum X -Wert 5, etwa an diesem Punkt ist das Netzwerk am instabilsten. Danach sinkt die Kurve wieder und nähert sich der X -Achse an, weil das Netzwerk der Firmen so stark ist, bzw. jede Firma ihre Anteile auf so viele andere Firmen verteilt, dass sie ein Pleitegang allein nicht abstürzen lässt.

Variiert man nun die Crossholds, erkennt man, dass die umso höhere Vernetzung den Pleitegang einer Firma wunderbar abfangen kann (siehe schwarzer Graph). Jede Kurve

im nachfolgenden Graphen repräsentiert die durchschnittliche Verbindungsanzahl (X -Achse) und die Anteile an anderen Firmen (rot: 20%; schwarz: 80%).

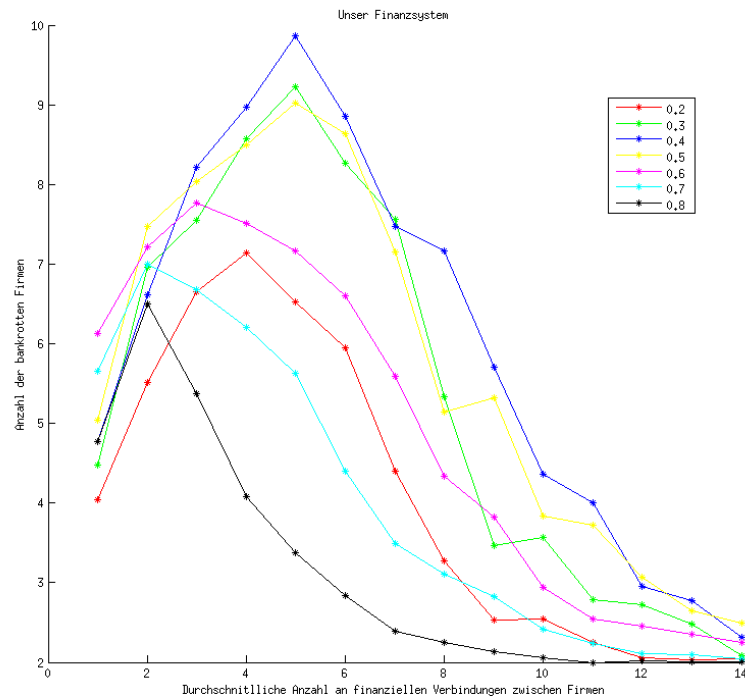


Abbildung 13: 1% der Firmen ging am Anfang pleite

Das Netzwerk scheint mit steigender Crosshold-Anzahl stabiler zu werden, vor allem der schwarze Graph, der die meisten Crossholds darstellt, zeigt für das Netzwerk ein sehr positives Verhalten. Je größer jedoch der anfängliche Crash ist, desto schlechter verhält sich ein Netzwerk mit hoher Crosshold-Anzahl unter Stress, anders ausgedrückt einem Preisschock von mehreren Assets.

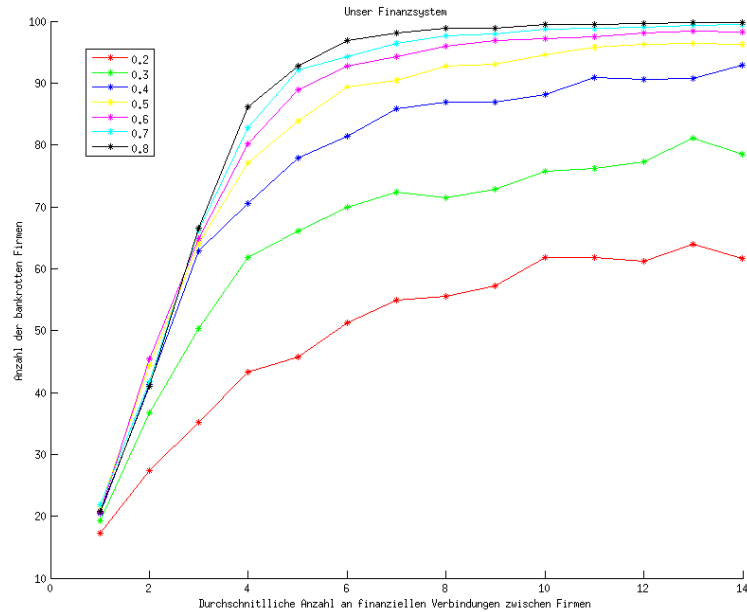


Abbildung 14: 7% der Firmen ging am Anfang pleite

Wie ersichtlich wird das Netzwerk mit steigenden Crossholds empfindlicher gegen große Schocks. Der schwarze Graph verhält sich nun am schlechtesten, die meisten Firmen gehen pleite. Im Gegensatz dazu zeigt der rote Graph, der eine eher niedrige Anzahl an Crossholds repräsentiert, dass, im Vergleich zu den anderen, nur sehr wenig Firmen Pleite gehen.

Mit stärkerer Vernetzung der Firmen erhöht sich die Wahrscheinlichkeit eines Kollaps zusätzlich, siehe x-Achse, allerdings nur bis zu einer gewissen Grenze. Diese ist im oberen Graphen leider nicht ersichtlich, da nicht genug Platz ist, lediglich am roten Graphen lässt sich eine leichte Abflachung erahnen. Dieser Trend ist in seinem weiteren Verlauf bei allen Graphen zu beobachten, je mehr Crossholds, desto später sinkt der Graph wieder. Ist die Vernetzung schließlich stark genug (im oberen Beispiel bei durchschnittlich 90 Verbindungen) steigt die Stabilität erneut stark an. Das Netzwerk kann den größeren Schock wieder so gut aufteilen, dass eine Krise verhindert wird.

Erweiterung des Modells und alternative Berechnung in C# Die Berechnung der Konkurskosten kann auf zwei Arten erfolgen:

A $\beta = \text{Threshold} * \text{ursprünglicher Wert der Firma} - \text{aktueller Wert}$

B $\beta = \text{Threshold} * \text{ursprünglicher Wert der Firma}$

Im Vergleich zu den *A*-Konkurskosten reicht bei den *B*-Konkurskosten ein kleinerer Initialschock um das System aus dem Gleichgewicht zu bringen. Die Berechnungen in MATLAB beruhen auf den *A*-Konkurskosten, während die *C#* Berechnungen auf die *B*-Konkurskosten baut. In den folgenden Graphiken werden immer zwei Werte variiert, dabei wird ein Wert auf der *X*-Achse aufgetragen und der andere durch die Schar dargestellt, wobei hellere Farben höhere Werte bedeuten.

Es werden meist verschiedene Werte der Ergebnisse geplottet:

- rot: Durchschnittliche Anzahl an pleite gegangenen Firmen (eigentlicher Messwert)
- grün: Standardabweichung
- blau: durchschnittliche Ausbreitungsdauer der Krise

In diesem Modell wird nicht mehr jede Firma durch ein für alle gleich gewichtetes Asset gekennzeichnet. Die Anzahl der Assets, deren maximaler Wert und die Anzahl wie viele verschiedene Assets und wie viele von diesen Assets eine Firma besitzt sind variabel. Da keine realistischen Werte bekannt sind, ist es schwierig mit diesem Modell zu arbeiten. Um trotzdem Vergleichbarkeit zu wahren, wurden meist gleiche oder ähnliche Werte verwendet. Meist werden nur hundert Testläufe pro Eingabewert verwendet, da höhere Zahlen einen enormen zusätzlichen Rechenaufwand bedeuten würden. Der Threshold liegt standardmäßig bei 93%, jede der 100 Firmen ist im Schnitt an 9 der 200 Assets beteiligt. Die Schockgröße liegt meist zwischen 1 und 4.

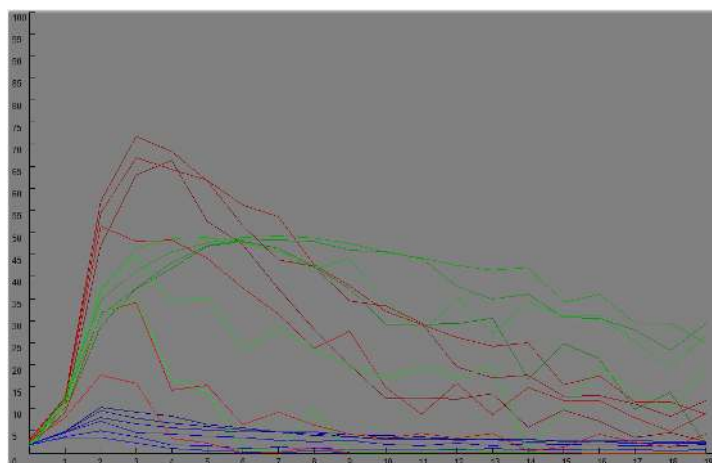


Abbildung 15: Simulation 1

Die durchschnittliche Anzahl der Verbindungen wird auf der *X*-Achse dargestellt. Die geringen Werte an Pleiten an den Stellen Null und Eins, also den roten Graphen, erklären sich dadurch, dass nicht genügend Verbindungen vorhanden sind um die Krise

weiter zu verbreiten. Das darauf folgende Maximum der Kurve entsteht dadurch, dass zwar genügend Verbindungen existieren um die Krise zu verbreiten, jedoch nicht genug um die Konkurskosten der Firmen, welche durch den Initialschock zu Grunde gingen, zu stützen. Die verbundenen Firmen müssen also Insolvenz anmelden, wodurch neue Kosten entstehen. Dies resultiert in einer Kettenreaktion, die in den meisten Fällen alle Firmen in den Konkurs zieht. Nach diesem kritischen Bereich sinkt die Anzahl der pleite gegangenen Firmen stetig, bis eventuell, bei einem leichten Schock, nicht einmal die Firmen, welche vom Initialschock betroffen waren, in Konkurs gehen.

Die Crossholds werden hier durch die Schar dargestellt und reichen von 30% bis 80% Prozent, jeweils in 10-er Schritten. Sie bestimmen wie stark die Firmen vernetzt sind, ohne Verbindungen haben sie jedoch keinen Einfluss. Sind die Crossholds gering wird die Kurve der durchschnittlichen Pleiten (rot) gestreckt. Dies hat seine Ursache darin, dass der Schaden stärker verteilt wird und so die Krise sofort eingedämmt wird. Vor allem im kritischen Bereich zwischen einer und fünf Verbindungen spielt dies eine wesentliche Rolle. Die Größe der Crossholds bestimmt also, wie anfällig das Netzwerk ist. Weniger Crossholds zögern die Ausbreitung der Krise hinaus, wie man dem blauen Graphen entnehmen kann.

Die grünen Graphen zeigen, dass das System den Schock entweder sehr schnell auffängt, oder aber vollständig an ihm zugrunde geht, sich also fast nie bei der Hälfte der Firmenanzahl wieder fängt, was einem die roten Graphen vielleicht suggerieren. Die roten Graphen zeigen also meistens nicht die Anzahl der Firmen die überleben, sondern meistens die Wahrscheinlichkeit mit der das System kollabiert.

Konklusion: Zumindest wenn der Schock gering ist (1-2 Assets) ist es für das Gesamtsystem also zuträglicher, wenn die Firmen nicht isoliert, sondern vernetzt sind, so dass der entstehende Schaden aufgeteilt wird, was durch die Verbindungsichte geregelt wird. Wie stark diese Dämpfung ist wird durch die Crossholds, also den Teilen die nicht an der Börse gehandelt werden ausgedrückt.

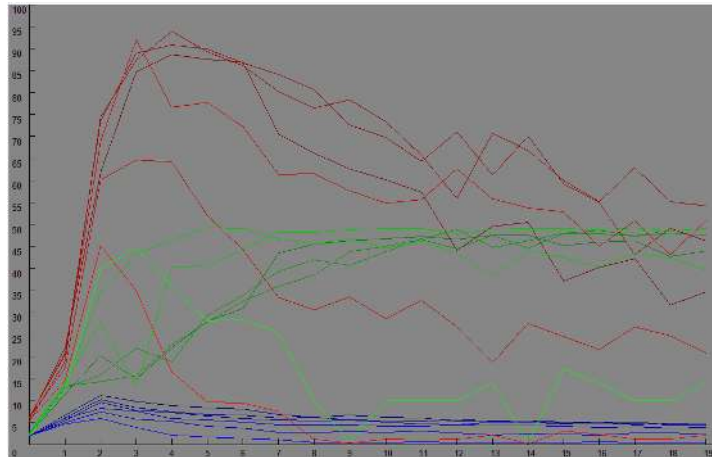


Abbildung 16: Simulation 2

Betrachtet man die Prozentzahl der Pleiten, so ähnelt der Graph dem vorherigen stark, die einzige Änderung ist, dass der Graph extremer ist, eben da der Initialschock stärker ist. Dadurch wird der Schock näherungsweise um den Faktor zwei gestreckt.

Konklusion: Erhöht man die Stärke des Schocks, so steigt die Initiallast, demnach benötigt man mehr Verbindungen um sie so zu verteilen, damit sich die Krise nicht ausbreitet. Ab einer gewissen Schockstärke kann das System die Last nicht mehr tragen, sogar bei maximaler Vernetzung. Also kann eine stärkere Vernetzung zwar kleine Krisen verhindern, macht jedoch anfälliger für große Schocks.

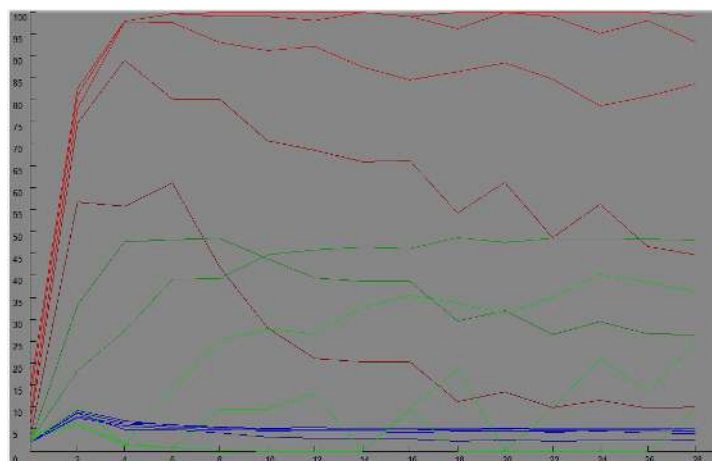


Abbildung 17: Simulation 3

Die Ausmaße der Krise steigen logischerweise mit der Stärke des Initialshocks. Die Anzahl der Verbindungen welche nötig sind um den Schock auszugleichen steigt dabei enorm, da anfangs mehrere Firmen bankrott gehen, woraus ein höherer Gesamtschaden

entsteht. Dieser wird nun ja auf andere Firmen übertragen, da diese jedoch nicht ebenfalls bankrott gehen, muss der Schaden auf viele Firmen übertragen werden.

Konklusion: Je höher der Initialschock ist, umso weniger Erfolg hat die Verhinderung der Krise durch Steigerung der Konnektivität. Diese Erkenntnisse gelten also nicht nur für die *A*-Konkurskosten sondern finden auch bei den *B*-Konkurskosten Anwendung.

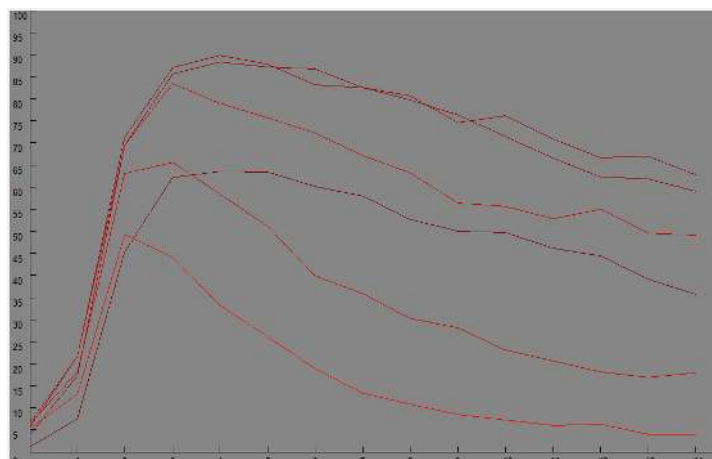


Abbildung 18: Simulation 4

Wie man vor allem an den beiden hellsten Graphen erkennen kann, nähert sich die Kurve einem gewissen Wert an, der von den Assets die die Firmen im Schnitt halten abhängig ist. Je größer der Schock ist, umso länger dauert die Annäherung. Interessant hierbei ist, dass in beiden Fällen der Schaden zunächst steigt und dann langsam wieder absinkt. Im Bereich zwischen 20 und 25 Assets pro Firma sinkt der Schaden für alle Varianten der Verbindungsichte zu Firmen unter den Schadenswert von einem Asset pro Firma, also der minimalen Anzahl an Assets an denen eine Firma beteiligt sein kann. Lediglich wenn alle Firmen für sich alleine stehen steigt die Chance für eine Firma zu überleben, wenn sie nur an wenigen Assets beteiligt ist, da dann die Wahrscheinlichkeit, dass jenes Asset, von welchem die Firma abhängt, von dem Schock betroffen ist, sinkt. In allen anderen Fällen ist es sicherer in mehrere Assets zu investieren.

Man könnte auch ein sehr stabiles Netzwerk erschaffen, indem man die Crossholds weglässt. Der folgende Graph visualisiert diesen Gedanken, die dunkle, rote Kurve am unteren Rand des Plots zeigt ein Netzwerk dieses Typs. Jede Kurve zeigt eine durchschnittliche Verbindungsanzahl von 0 bis 12, auf der X-Achse werden durchschnittliche Verbindungen zu Assets angezeigt. Wie man an den gezeigten Parametern erkennen kann, ist diese Simulation ein wenig abgeändert worden.

Wenn Firmen keine Crossholds an anderen Firmen besitzen, also nicht mehr vernetzt sind, kann sich keine Krise ausbreiten. Allerdings kann auch nicht mehr von einem Netzwerk gesprochen werden. Interessant ist, dass bei einem kleinen Schock und hohen Verbindungen weniger Firmen pleite gehen als bei keinen Verbindungen, weil das Netzwerk den Schock sofort anfängt.

Benutzte Parameter:

- 100 Firmen
- Asset-Anzahl einer Firma = 0 – 28
- Verbindungsanzahl zwischen den Firmen = 0 – 12 Verbindungen
- Gecrashte Assets = 2
- Crosshold der Firmen = 50%

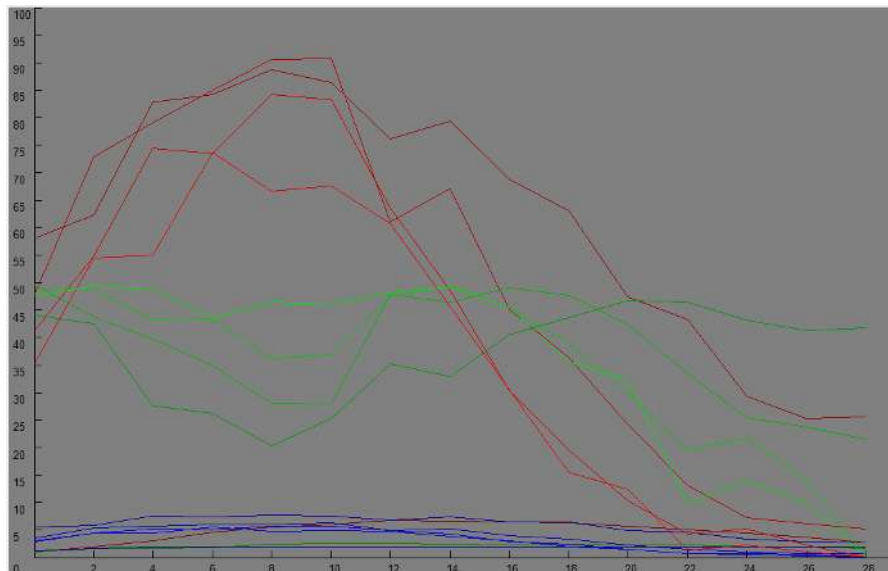


Abbildung 19: 7% der Firmen ging am Anfang pleite

Diese Graphik bestätigt, was die Vorherigen schon deutlich gemacht haben. Nur bei 0 Verbindungen schlägt der Graph nie aus. Diese Methode rentiert sich also bis die Konnektivität so stark ist, dass nicht einmal mehr die direkt vom Initialschock betroffenen Firmen bankrott gehen. Steigt der Schock so steigt auch die benötigte Anzahl an Verbindungen um dies auszugleichen, dies geht jedoch nur bis zu einer gewissen Schockstärke. Schwache aber vorhandene Konnektivität stellt sich immer wieder als letal dar.

Konklusion: Je mehr eine Firma an verschiedenen Assets beteiligt ist, umso weniger stützt sich die Firma auf ein einzelnes Asset. Hat eine Firma viele Assets ist sie mit höherer Wahrscheinlichkeit von dem Schock betroffen, ab einer bestimmten Grenze jedoch spielt das keine Rolle mehr, da ein einzelnes Asset (oder auch mehrere) den Wert der Firma nicht mehr unter den Threshold ziehen kann (siehe hellster und dunkelster roter Graph im Bereich 20-22). Auch hier gilt natürlich, dass die Größe des Schocks starken Einfluss hat. Wenn der Schock zu groß wird, hilft lediglich vollständige Isolation.

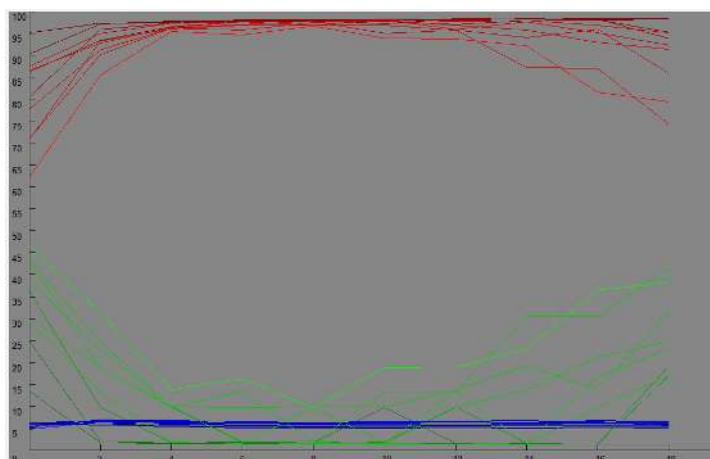


Abbildung 20: Simulation 6

An je mehr Assets man beteiligt ist, umso mehr steigt die Wahrscheinlichkeit, dass eines der eigenen Assets von dem Schock betroffen ist. Wenn man allerdings an genug Assets beteiligt ist, dann ist bei einem gewissen Schock im Schnitt ein gewisser, in diesem Fall unbedeutender, Prozentsatz davon betroffen. Dieser geringe Prozentsatz hat nun natürlich auch nur einen verhältnismäßig geringen Anteil an dem Gesamtwert der Firma. Je mehr Assets verfügbar sind, umso geringer ist die Wahrscheinlichkeit, dass ein Asset, von dem man abhängt, von dem Schock betroffen ist.

Maximalbelastung eines finanziellen Netzwerks Wenn jede Firma mit allen andern Firmen verbunden ist und alle Anteile an andere Firmen verteilt, anders ausgedrückt: Wenn das Netzwerk maximal verbunden ist, ergibt sich ein maximaler, überstehbarer Schock von 1-threshold, da ab diesem Wert der Gesamtwert der Assets, über den auch die Firmenwerte errechnet werden, genau unter den Threshold fällt.

Ist der Threshold z.B. 0.93% ergibt sich ein maximaler Schock von 7%. Im oberen Plot erkennt man dies sehr schön, je mehr sich die Crosshold-Anzahl 100% annähert, desto mehr Firmen gehen pleite. Ein guter Vergleich ist der schwarze Graph mit 80%

Crossholds im Gegensatz zum roten Graphen mit 20% Crossholds. Der rote Graph zeigt eindeutig weniger Pleitegänge.

2.5.2 Core-Periphery-Modell

Das *Core-Periphery-Modell* ist realistischer als das *P2P-Modell*. In der Realität gibt es nämlich einen Ring an großen Banken und Firmen die mit „peripheren“ Unternehmen finanziell zusammenhängen. In der Core-Periphery-Simulation werden die Core-Unternehmen als Kreis in der Mitte dargestellt und die peripheren Organisationen rundherum in einem größeren Kreis angeordnet.

Eine große Bank besitzt ein Asset von achtfachem Wert eines peripheren Unternehmens. Eine große Bank hat viele Anteile an peripheren Firmen, diese hingegen nur ein bis zwei Anteile an großen Banken. Je nach Simulation haben auch die Peripheren untereinander Firmenanteile.

Die ungewichtete Matrix G mit den Verbindungen zwischen den Firmen wird während der Generierung in vier Matrizen unterteilt, um die Verbindungsanzahl und die Crossholds der zwei verschiedenen Bankentypen einzeln verändern zu können.



Abbildung 21: Die Unterteilung der C-Matrix

- gelb: Anteile von großen Firmen an großen Firmen
- rot: Anteile von großen Firmen an kleinen Firmen
- blau: Anteile von kleinen Firmen an großen Firmen
- grün: Anteile von kleinen Firmen an kleinen Firmen

Den gerichteten und gewichteten Graphen C kann man anschließend mit jeweiligen Crosshold-Werten für jede der oben genannte Verbindungsarten erstellen. Die Berechnung des Marktwertes der Firmen funktioniert exakt wie im Core-Modell.

Eine große Bank besitzt ein Asset von achtfachem Wert eines peripheren Unternehmens. Eine große Bank hat um die 12 Anteile an peripheren Firmen, diese 2 – 3 Anteile an großen Banken. Je nach Simulation haben auch die Peripheren untereinander Firmenanteile. Der Preisschock wird dadurch ausgelöst, dass der Wert mehrerer kleiner Firmen gleichzeitig auf 0 gesetzt wird.

Bounce-Effekt Beim *Core-Periphery-Modell* kann man etwas interessantes beobachten. Wir tauften das Phänomen „*Bounce-Effekt*“. Der Schaden springt zwischen der Peripherie und dem Core hin und her. Im Folgenden wird eine Serie von Graphen gezeigt, auf der der *Bounce-Effekt* zu beobachten ist.

Benutzte Parameter:

- 100 Firmen
- Keine Verbindungen zwischen peripheren Firmen
- Durchschnittliche Verbindungen von Peripherie zu Core = 1.5
- Durchschnittliche Verbindungen von Core zu Peripherie = 12
- Anteile von Peripherie an Core = 65%
- Anteile von Core an Peripherie = 10%
- Anteile von Core an Core = 16%
- Threshold = 93%
- Jede große Bank ist mit den anderen großen verbunden
- Crosshold der Firmen = 50%
- 10 Firmen gingen am Anfang Pleite

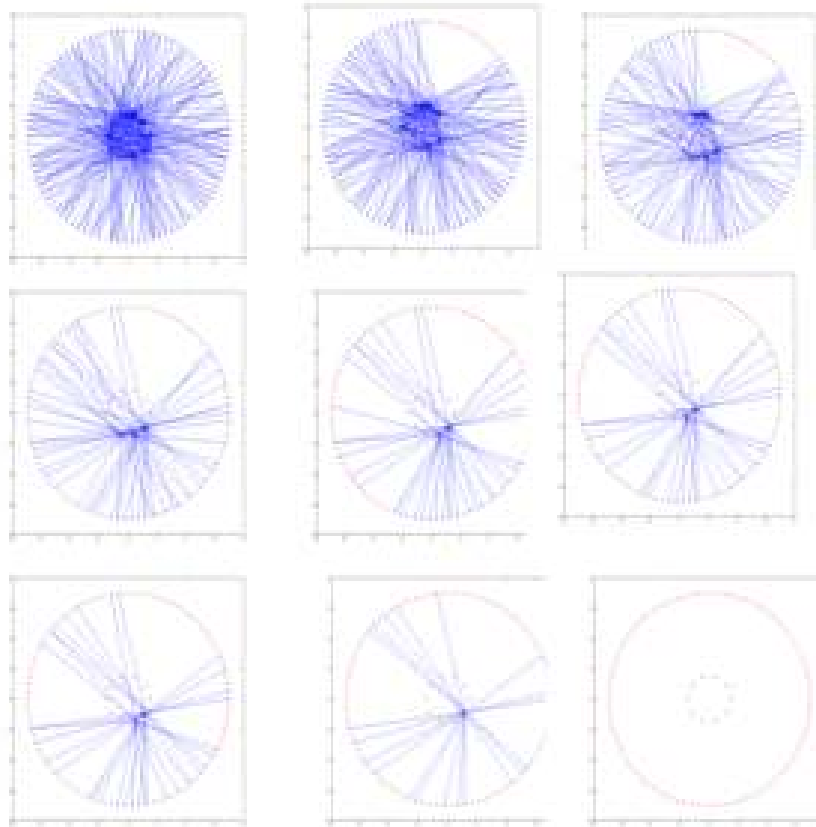


Abbildung 22: *Bounce-Effekt*

Am Anfang brechen zehn Firmen der Peripherie zusammen und diese lösen mehrere Pleiten im Core-Bereich aus. Doch dann gehen nur noch periphere Firmen Pleite, bis der Schaden so groß ist, dass eine große Bank fällt. Daraufhin kann sich auch die letzte überlebende Großbank nicht halten. Der *Bounce-Effekt* entsteht unter anderem durch die fehlende Vernetzung von kleinen zu kleinen Firmen. Der Schaden kann am Anfang nur einen Weg gehen: Er zerstört eine oder mehrere Großbanken. Diese sind sehr stark vernetzt und gehen auch pleite, womit sie wiederum kleine Firmen mitreißen. Der Schaden ist nach außen gewandert. Von dort aus kann er nur wieder in die Mitte und der Zyklus beginnt von vorne bis entweder der Schaden annulliert ist oder alle Firmen pleite sind.

Abhängigkeit zwischen Core- und Periphery-Bereich Die folgenden Plots zeigen sowohl die gegenseitige Abhängigkeit der Crossholds von kleinen Firmen an großen Firmen, als auch von großen Firmen an kleinen Firmen. Die Berechnung wurde mit 500 Runden pro Wertepaar gerechnet und am Anfang gingen 10% der kleinen Firmen pleite, wobei diese wie auch schon beim *Bounce-Effekt* keine Verbindungen untereinander haben. Die

Thresholds von kleinen an großen Firmen und umgekehrt wird variiert (Siehe Achsenbeschriftung).

Versuchsaufbau:

- 100 Firmen
- Rundenanzahl pro Wert = 100
- Gecrashte Assets = 4 (kleine Firmen)
- Keine Verbindungen zwischen peripheren Firmen
- Durchschnittliche Verbindungen von Peripherie zu Core = 2.5
- Durchschnittliche Verbindungen von Core zu Peripherie = 12
- Anteile von Peripherie an Core = 20%
- Anteile von Core an Peripherie = 20%
- Anteile von Core an Core = 16%
- Threshold = 93%
- Jede große Bank ist mit den anderen großen verbunden
- Crosshold der Firmen = 50%
- 10 Firmen gingen am Anfang Pleite

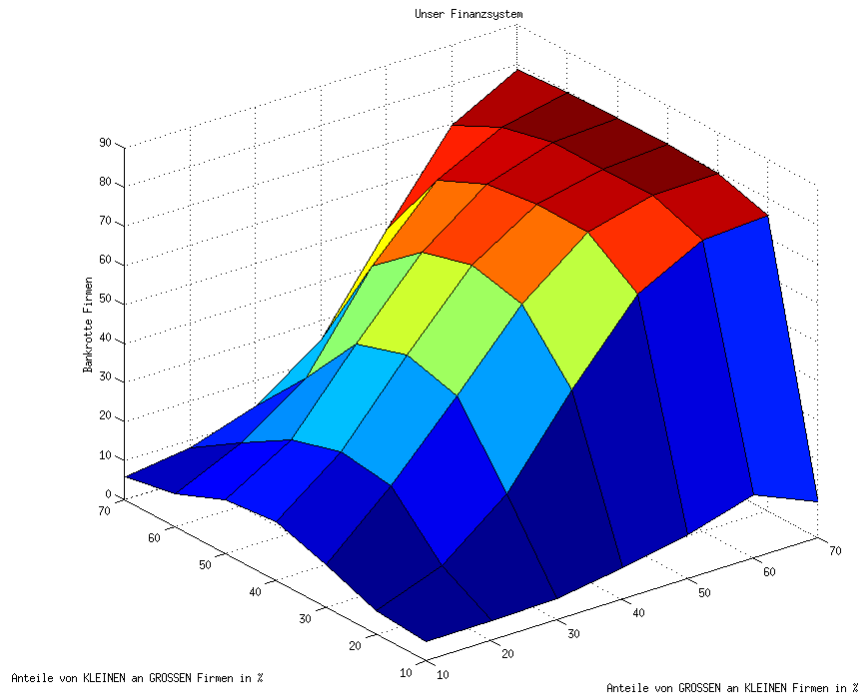


Abbildung 23: Abhängigkeit von Core und Peripherie

Am Koordinatenursprung ist der Schaden am geringsten. Je mehr Anteile große Firmen an kleinen besitzen, desto eher gehen viele Firmen pleite, da der Wert des Assets einer großen Firma das Achtfache des Wertes des Assets einer kleinen ist. Geht eine große Firma pleite, reißt sie viele Kleine mit.

Das Abflachen der Kurve auf der anderen Seite des Plots lässt sich, wie bereits bei anderen Plots, durch die breitere Verteilung des Schadens erklären. Geht eine große Firma pleite, teilen sich sehr viele kleine den Schaden auf und das Netzwerk kann diesen so kompensieren, bevor der Großteil der Firmen bankrott ist. Erhöht man allerdings den Preisschock verschiebt sich das Abflachen der Kurve aus dem Bereich des Plots, bzw. verschwindet ganz, da das Netzwerk den Schaden nicht mehr abfangen kann und sich die hohen Crossholds in diesem Fall als großer Nachteil erweisen.

Vergleich zwischen P2P- und Core-Periphery-Modell Wie man erkennen kann, sind die Graphen mit der Abhängigkeit zwischen den durchschnittlichen Verbindungen und den bankrotten Firmen beim *P2P-Modell* 24 und beim *Core-Periphery-Modell* 25 sehr ähnlich. Lediglich die Extremstelle hat sich verschoben. Dies lässt sich durch die nicht

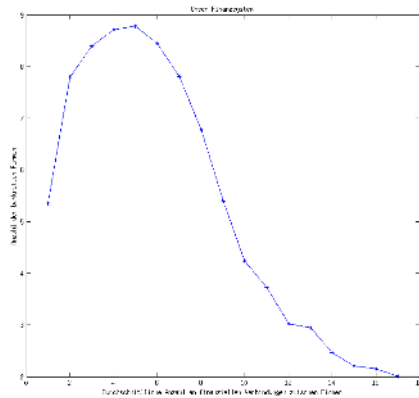


Abbildung 24: *P2P-Modell*

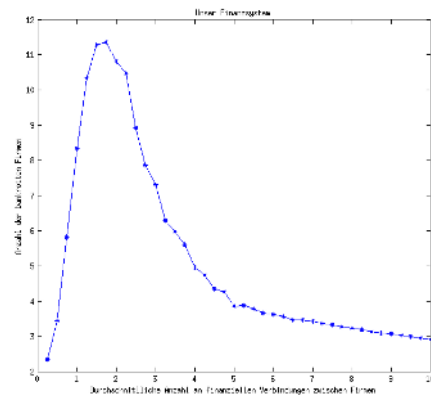


Abbildung 25: *Core-Periphery-Modell*

100%-ig aufeinander abgestimmten Parameter der beiden Modelle erklären. Nichtsdestotrotz ist das Ergebnis dasselbe, der Verlauf verhält sich identisch.

Daraus kann man schließen, dass sich ein *P2P-Netzwerk* beinahe identisch zu einem *Core-Periphery-Netzwerk* verhält. Vorteil des *Core-Periphery-Modells* ist jedoch, dass sich dieses mit den „Core“ Elementen sehr gut kontrollieren und steuern lässt, ein Staat kann zum Beispiel mithilfe seiner Nationalbank stark auf den Finanzmarkt einwirken. Das ist meist erwünscht, Staaten können so die Wirtschaft vor Spekulationen schützen oder ihre Währung stabilisieren. Es ist lediglich wichtig, die Stabilität der Core-Banken zu gewährleisten, da das System ohne sie zusammenstürzt.

Es geht aus den Berechnungen eindeutig hervor, dass das Modell und entsprechend auch unser Finanzsystem nicht gut mit Pleiten umgeht. Es passiert zu schnell und leicht, dass eine Pleite mehrere Firmen mitreißt und sie eine Kaskade an Konkursen auslöst. Das Netzwerk ist instabil, die Vernetzung wirkt sich fast ausschließlich negativ aus.

Weder das *P2P-* noch das *Core-Periphery-Modell* konnten sich unter Stress behaupten. Die einzige Chance ein angeschlagenes System zu retten ist den Schaden durch die Bankrotte möglichst früh zu bezahlen. Genau das haben die Staaten in der Wirtschaftskrise 2007 getan. Doch das Retten des Netzwerks kostet immens viel Geld. Nach unseren Berechnungen ist das Netzwerk am stabilsten wenn gar keine Crossholds existieren. Da wir keine Wirtschaftsexperten sind ist es unklar inwiefern ein finanzielles Netzwerk ohne Crossholds denkbar wäre. Dennoch kann eindeutig die Schwäche des jetzigen Systems bewiesen werden. Unser Vorschlag für ein Krisen-sicheres System ist daher, ohne Crossholds auszukommen.

2.6 Resümee

Einige interessante Erkenntnisse konnten gewonnen werden. Generell verhielten sich alle Graphen gleich, einige waren nur extrem gestreckt bzw. verzerrt. Jede Kurve kehrte früher oder später in die Nähe der X -Achse zurück, auf der die durchschnittliche Verbindungsanzahl aufgetragen wurde.

Folgendes lässt sich schließen: Hohe Diversität ist immer gut für ein Netzwerk. Mit einer Ausnahme: Es konnte stets ein kritischer Punkt festgestellt werden, an dem die Anzahl der Konkurse förmlich explodierte. Dies lässt sich durch Folgendes erklären: Zuerst gibt es zu wenig Verbindungen um die Krise zu verbreiten. Je mehr Verbindungen, desto besser verbreitet sich die Krise. Doch ab einem gewissen Punkt verteilt sich die Last immer stärker, bis zum Schluss kaum noch Firmen pleite gehen, weil der Schaden auf so viele Firmen aufgeteilt wird, dass keine von ihnen pleite geht.

Sozioökonomische Physik

Macht von Politikern und Fairness in einer Gesellschaft

Datum: 10. - 16. Februar 2013

Betreuer: Daniel Kraft

Gruppe:

Clemens Bischof

Martin Maritsch

Lukas Taus

Markus Tibet

Daniel Windisch



Inhaltsverzeichnis

1	Einleitung	2
2	Definitionen	3
2.1	Das Glück des Einzelnen	3
2.2	Die Gesellschaft als Ganzes	4
3	Konvexität	5
4	Einführung der Gerechtigkeit	6
4.1	Problem an der Formel $\mathcal{H}$ (1)	6
4.2	Definition von γ	6
5	Minimumsberechnung	7
5.1	Lagrange-Funktion	7
5.2	Ableitungen	8
5.2.1	Erste Ableitung(Gradient)	8
5.2.2	Zweite Ableitung(Hesse-Matrix)	8
5.3	Newton-Verfahren	8
6	Boltzmann-Verteilung	10
7	Berechnungen	12
7.1	Mittelwert der Unzufriedenheit berechnen	12
7.2	Zufriedenheit einer Gesellschaft	13
7.3	Bruttoinlandsprodukt	14
7.4	Gini-Koeffizient	15

1 Einleitung

Im Laufe der Menschheitsgeschichte kam es immer wieder zu tiefgreifenden politischen, gesellschaftlichen und wirtschaftlichen Veränderungen. Diese gingen von der Subsistenzwirtschaft vor dem Mittelalter, über den Merkantilismus der absolutistischen Herrschaftsformen in Mittelalter und Neuzeit, den Kapitalismus und Kommunismus im 19. und 20. Jahrhundert bis zur sozialen Marktwirtschaft, die heute in den meisten Staaten ihre Anwendung findet.

Wenn man sich die Frage nach der Gerechtigkeit in all diesen verschiedenen Systemen stellt, so wird man schnell zu einer Antwort kommen, denn es ist offensichtlich, dass die Mehrheit der Menschen in autoritären Machtverhältnissen vollständig unterdrückt war und für den Wohlstand des Herrschers oder des Vorgesetzten sorgte, während dieser sein Leben genießen konnte, ohne jemals viel geleistet zu haben.

In der Zeit des Merkantilismus war die Wirtschaft nur auf das Wohl der Obrigkeit also des Adels ausgelegt. Darunter mussten die niedrigeren Gesellschaftsschichten wie Bauern leiden, die oftmals Leibeigene waren und somit beinahe keinen Nutzen von ihrer harten Arbeit hatten, sondern alles an den Landesherren abliefern mussten. Auch im Kapitalismus war die Situation nicht anders: Viele Arbeiter leisteten unter meist unmenschlichen Verhältnissen ihren Dienst, während sich wenige Kapitalisten, die im Besitz der Fabriken waren, bereicherten. Die Idee Karl Marx sollte dieses ständige Ungleichgewicht in der Gesellschaft ausgleichen und allen die Erfüllung und das Glück bringen. Im Nachhinein kann man jedoch erkennen, dass dieser Gedanke von der besseren Welt auch eine Utopie ist. Denn gerade in den kommunistischen Staaten lebten und leben die Menschen heute noch in Unterdrückung und unter widrigen Verhältnissen. Das ist vor allem auf den Egoismus einiger mächtiger Personen, die nicht auf den Wohlstand der Bevölkerung achten, sondern darauf aus sind sich selbst zu bereichern, zurückzuführen. Und obwohl es uns heute in den westlichen Demokratien wirklich gut geht, ist die Tendenz da, dass die Vielzahl der Menschen völlig unterbezahlt ist und eine Minderheit an Mächtigen nicht viel für ihr Einkommen leisten. Es gibt also Systeme, in denen das Vermögen gerechter verteilt ist als in anderen, man wird aber keine Staatsform finden, die für jeden einzelnen das Optimum darstellt.

Vielmehr geht es darum die Glücklichkeit der Gesellschaft als ein Ganzes möglichst hoch zu halten. Es ist klar, dass das private Glück jedes einzelnen mehr oder weniger willkürlich ist und man im Prinzip keinen Einfluss darauf hat. Worüber man sich jedoch durchaus Gedanken machen kann, ist die wirtschaftliche Zufriedenheit jedes Menschen und wie stark sich dieser Parameter auf die gesamte Gesellschaft auswirkt.

Mit diesem Thema hat sich unsere Gruppe beschäftigt und versucht ein möglichst passendes mathematisches Modell zur Problemstellung zu entwickeln. Dazu mussten wir uns zuerst überlegen welche Größen überhaupt für das wirtschaftliche Glück notwendig sind und wie wir diese mathematisch definieren sollten. Danach konnten wir unsere Überlegungen umsetzen und mit Hilfe von technischen Mitteln ausweiten und verbessern, sodass wir am Ende zu einem aussagekräftigen Ergebnis kamen.

2 Definitionen

2.1 Das Glück des Einzelnen

Die erste Schwierigkeit lag darin entsprechende Parameter zu finden, die die Zufriedenheit der Gesellschaft im wirtschaftlichen Sinn bestimmen sollten. Eine solche Größe ist sicher der Wohlstand der Bevölkerung, der natürlich direkt vom Einkommen abhängt. Hat man einen hohen Lohn, so wird man sich mehr zu seinem Luxus leisten können und deswegen auch zufriedener sein. Das Geld allein ist jedoch relativ und kann nur dann das Glück eines Menschen widerspiegeln, wenn es zusammen mit seinem Arbeitsaufwand betrachtet wird. So wird eine Person, auch wenn sie im Vergleich zu allen anderen viel verdient, nicht bereit sein unverhältnismäßig viel für das Einkommen zu leisten. Im Gegensatz dazu wird jemand, der viel verdient und relativ dazu wenig arbeiten muss, umso glücklicher sein.

Um nun eine Funktion aufstellen zu können, die genau diese Eigenschaften besitzt, muss zuerst eine Überlegung zur Gesellschaft als ein Ganzes gemacht werden. Dazu definierten wir die einzelnen Personen als Teilchen in einem Stoff, wie zum Beispiel ein Gas. Ein solches System strebt in der Physik in den meisten Fällen nach einem möglichst energiearmen Zustand. Das heißt, dass das Glück in der Gesellschaft steigen würde, wenn unsere Funktion sinkt. Umso höher also der Funktionswert einer einzelnen Person ist, umso niedriger ist seine Zufriedenheit. Die Folge davon ist, dass die die Funktion jedes einzelnen zunimmt, wenn seine Arbeitsleistung steigt und abnimmt wenn sein Lohn steigt.

Aus dieser Definition kann man nun mehrere Darstellungen einer solchen Funktion ableiten. Eine der einfachsten Möglichkeiten ist Arbeit und Lohn in einen Bruch zu schreiben, sodass er die Bedingungen erfüllt. $f(a, l) = a/l$. Hier steht a für die Arbeit und l für den Lohn.

Warum diese Lösung für unser Modell jedoch nicht optimal ist, ist mehr oder weniger offensichtlich. Das hängt mit zwei Tatsachen zusammen, die die Arbeit und den Lohn betreffen. Eine Person ist nur imstande eine begrenzte Arbeitsleistung zu bringen. Diese Obergrenze nennen wir A . Auf der anderen Seite muss jeder ein gewisses Mindesteinkommen L haben, das über 0 liegt, damit ein Überleben überhaupt möglich ist. Für a und l gelten also folgende Bedingungen:

$$a \in [0; A], l \geq L$$

Nun bietet es sich an eine Funktion zu wählen, die diese Bedingungen bereits inkludiert. Eine solche wäre beispielsweise

$$f(a, l) = -\ln(A - a) - \ln(l - L)$$

Diese erreicht die Werte A und L wegen der enthaltenen Logarithmusfunktion nie, hat aber die zuvor beschriebenen Eigenschaften. Wenn man nun zur Vereinfachung L gleich 0 und A gleich 1 setzt und die Funktion in Abhängigkeit von a und l für eine Person darstellt, so kommt man zu folgendem grafischen Ergebnis.

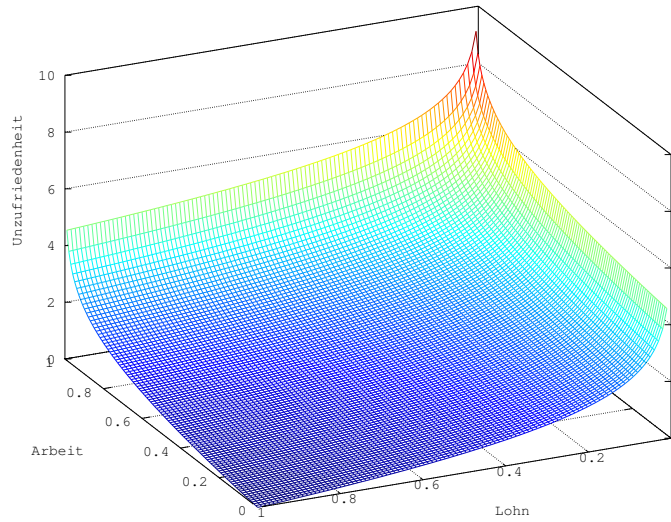


Abbildung 1: Unzufriedenheit des Einzelnen in Abhängigkeit von γ

Hier kann man feststellen, dass die Forderungen, die wir an die Funktion stellten, zutreffen.

2.2 Die Gesellschaft als Ganzes

Um nun die einzelnen Teilfunktionen zur gesamten Glücklichkeit der Gesellschaft zusammenfassen will, so wäre eine Möglichkeit dies einfach durch Addition durchzuführen. Man muss sich jedoch überlegen, ob tatsächlich alle Menschen gleich viel zählen. Die offensichtliche Antwort ist "nein". Deswegen muss noch eine weitere Variable m eingeführt werden, die die Macht jeder einzelnen Person, also ihren Einfluss auf die gesamte Gesellschaft, als einen Anteil von 1 darstellt. Das heißt die Summe der Machteinflüsse aller n Personen ist immer gleich eins. Da wir die Gesellschaft als ein abgeschlossenes System betrachten, können davon ausgehen, dass die Summe der Arbeit der Summe der Löhne entsprechen muss. Dazu ist es notwendig, dass Arbeit und Lohn die selbe Einheit besitzen. Man könnte sie beispielsweise an der entsprechenden Menge eines Gutes messen. In diesem Fall würde man zum Beispiel sagen, dass umgerechnet in einem Staat eine bestimmte Menge an Getreide erarbeitet wird. Man kann dann auch nicht mehr Lohn ausgeben, als diese Menge an Geld einbringt. Diese Summe der Arbeit beziehungsweise der Löhne entspricht genau dem Bruttoinlandsprodukt des betrachteten Staates. Es ergeben sich also zwei weitere Bedingungen:

$$\sum_{i=1}^n m_i = 1; \sum_{i=1}^n a_i = \sum_{i=1}^n l_i$$

Eine Person wird bestimmt durch $x_i = (m_i, a_i, l_i)$ und ist aus $U = \{(m_i, a_i, l_i) | m_i \in [0; 1], a_i \in [0; A], l_i \geq L\}$. Die entsprechende Definitionsmenge lautet also

$$\Omega := \{(x_1, \dots, x_n) | x_i \in U, \sum_{i=1}^n a_i = \sum_{i=1}^n l_i, \sum_{i=1}^n m_i = 1\}$$

Wenn wir nun alle Parameter in einer Formel verbinden, sodass m in einem Funktionswert eine bestimmte Wichtigkeit verleiht, so kommt man zu folgendem Bildungsgesetz für die Glücklichkeit der betrachteten Gesellschaft:

$$\mathcal{H} = \sum_{i=1}^n m_i \cdot f(a_i, l_i) \quad (1)$$

3 Konvexität

Die Konvexität einer Funktion bringt viele positiven Eigenschaften mit sich. Eine Funktion ist konvex, wenn sie keine Sprünge aufweist und wenn sie selbst unter der Verbindungsline zweier ihrer Punkte liegt. Das heißt, dass sie linksgekrümmt ist.

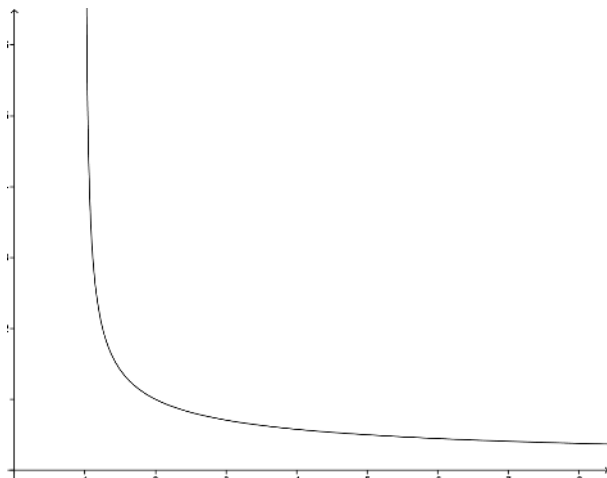


Abbildung 2: Beispiel für eine konvexe Kurve

Wenn eine Funktion f konvex ist gilt:

$$\frac{d^2 f(a, l)}{da^2} > 0; \frac{d^2 f(a, l)}{dl^2} > 0 \quad (2)$$

Eine konvexe Funktion ist in unserem Fall sehr gut anwendbar, weil sie durch ihre ab-, oder zunehmende Steigung die Weite eines mehr bekommenen Euro deutlich abbildert, was zu einem einfachen mathematischen Lösungsverfahren führt. Zweitens bewirkt die Funktion bei einer Zusammenarbeit einen Zuschuss an Einkommen, als bei einer geraden Funktion. Ohne diese Eigenschaft würde sich Zusammenarbeit gar nicht rentieren. Zum Beispiel würden Herr X und Herr Y, wenn sie zusammen arbeiten würden mehr verdienen als sie um Durchschnitt verdienen würden, wenn sie getrennt arbeiten. So wäre es von Vorteil wenn die von uns gesucht Funktion positiv wäre. Dann müsste gelten:

$$f\left(\sum_{i=1}^n m_i \cdot a_i, \sum_{i=1}^n m_i \cdot l_i\right) \leq \sum_{i=1}^n m_i \cdot f(a_i, l_i) \quad (3)$$

$$\text{wenn gilt } \sum_{i=1}^n m_i = 1 \quad (4)$$

Für unsere Werte wäre das dann:

$$f((m_1 \cdot a_1 + m_2 \cdot a_2), (m_1 \cdot l_1 + m_2 \cdot l_2)) \leq m_1 \cdot f(a_1, l_1) + m_2 \cdot f(a_2, l_2) \quad (5)$$

Die Konvexität kann man mit der sogenannten Hesse-Matrix beweisen, welche später noch genauer erklärt wird.

4 Einführung der Gerechtigkeit

4.1 Problem an der Formel $\mathcal{H}$ (1)

Falls eine Person keine Macht besitzt wird sie für die allgemeine Glücklichkeit vollkommen irrelevant, was nicht der Wirklichkeit entspricht, da das Glück, wenn nur eine Person die gesamte Macht besitzt, nur von dieser abhängt. Also führen wir eine neue Variable ein: γ

Formel mit neuer Variable:

$$\mathcal{H} = \sum_{i=1}^n (((1 - \gamma) \cdot m_i + (\frac{\gamma}{n})) \cdot f(a_i, l_i))$$

4.2 Definition von γ

$\gamma \in [0; 1]$... ist ein Gerechtigkeitsfaktor. Je größer γ wird, umso gerechter wird die Macht verteilt.

$n \in [0; \infty)$... steht für die Personenzahl.

$\gamma = 1$... Die Macht wird zu gleichen Teilen auf alle Personen ausgeteilt

$\gamma = 0$... Es herrscht keine Gerechtigkeit, es ist wieder die ursprüngliche Formel

$\gamma \in]1; 0[$... Es ist weder möglich, dass eine Person die gesamte Macht besitzt, noch dass die Macht vollkommen gerecht aufgeteilt ist

5 Minimumsberechnung

5.1 Lagrange-Funktion

$$L(x, \lambda) = f(x) + \lambda e(x) \quad (6)$$

In der Lagrange-Funktion (6) steht $f(x)$ für die eigentliche Funktion die minimiert werden soll. Dazu kommt der Lagrange-Multiplier (λ) und eine Funktion $e(x)$ die die Bedingung beschreibt, welche 0 sein muss.

Die Lagrange-Funktion (6) ersetzt die erste Ableitung und wird 0 gesetzt.

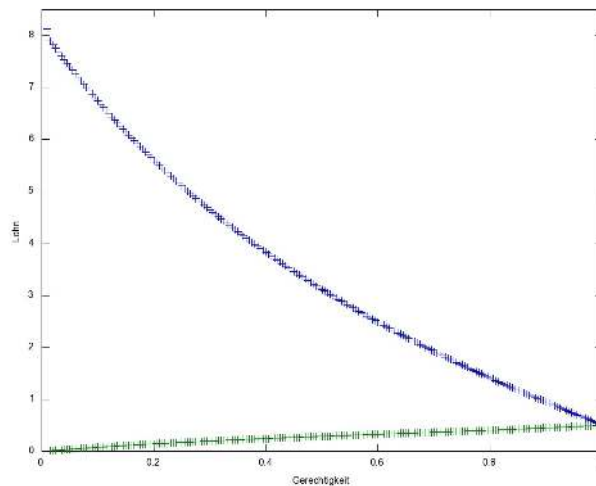


Abbildung 3: Gerechtigkeit und Lohn

Diese Abbildung zeigt den Zusammenhang zwischen der Gerechtigkeit und dem Lohn. Man kann ablesen dass der Lohn der mächtigen Person (blau), je mehr die Gerechtigkeit ansteigt, sich immer mehr dem Lohn der Person ohne Macht (grün) annähert.

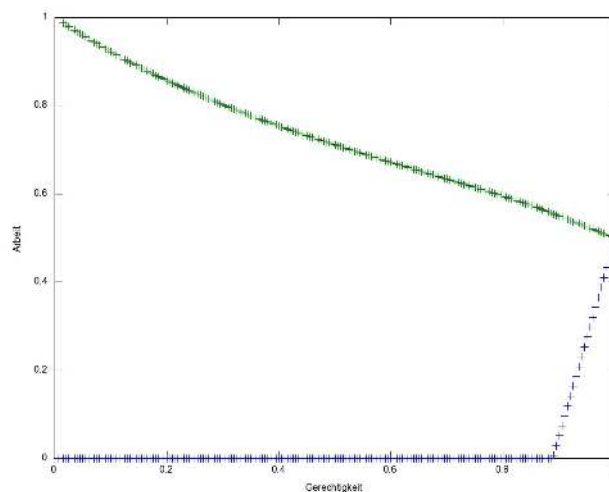


Abbildung 4: Gerechtigkeit und Arbeit

Aus dieser Grafik kann man das Zusammenspiel von Gerechtigkeit und Arbeit ablesen. Die Arbeit einer Person mit wenig Macht (grün) nimmt, je gerechter es wird, immer mehr ab, während die Arbeit einer

Person mit viel Macht (blau) ab einem gewissen Punkt stark ansteigt und sich der Arbeit der anderen Person immer weiter annähert.

5.2 Ableitungen

5.2.1 Erste Ableitung(Gradient)

Man muss in einer festgelegten Reihenfolge nach allen Variablen ableiten von der die Funktion abhängig ist. Diese Ableitungen werden in einem Vektor zusammengefasst, dem sogenannten Gradienten. Dieser Gradient ist die 1. Ableitung der Funktion. In unserem Fall sieht dieser dann so aus:

$$\begin{pmatrix} \frac{\partial f}{\partial a_0} \\ \frac{\partial f}{\partial l_0} \\ \frac{\partial f}{\partial a_1} \\ \frac{\partial f}{\partial l_1} \\ \frac{\partial f}{\partial \lambda} \end{pmatrix}$$

5.2.2 Zweite Ableitung(Hesse-Matrix)

Die Hesse-Matrix besteht aus den zweiten Ableitungen aller Variablen einer Funktion. Wir verwendeten zur Erstellung dieser Matrix das Programm Maxima, da dieses bereits eine Funktion dafür eingebaut hat.

In unserem Fall sieht die Struktur der Matrix so aus:

$$\begin{pmatrix} [3ex] \frac{\partial^2 f}{\partial^2 a_1} & \frac{\partial^2 f}{\partial a_1 \partial l_1} & \frac{\partial^2 f}{\partial a_1 \partial a_0} & \frac{\partial^2 f}{\partial a_1 \partial l_0} & \frac{\partial^2 f}{\partial a_1 \partial \lambda} \\ \frac{\partial^2 f}{\partial l_1 \partial a_1} & \frac{\partial^2 f}{\partial^2 l_1} & \frac{\partial^2 f}{\partial l_1 \partial a_0} & \frac{\partial^2 f}{\partial l_1 \partial l_0} & \frac{\partial^2 f}{\partial l_1 \partial \lambda} \\ \frac{\partial^2 f}{\partial a_0 \partial a_1} & \frac{\partial^2 f}{\partial a_0 \partial l_1} & \frac{\partial^2 f}{\partial^2 a_0} & \frac{\partial^2 f}{\partial a_0 \partial l_0} & \frac{\partial^2 f}{\partial a_0 \partial \lambda} \\ \frac{\partial^2 f}{\partial l_0 \partial a_1} & \frac{\partial^2 f}{\partial l_0 \partial l_1} & \frac{\partial^2 f}{\partial l_0 \partial a_0} & \frac{\partial^2 f}{\partial^2 l_0} & \frac{\partial^2 f}{\partial l_0 \partial \lambda} \\ \frac{\partial^2 f}{\partial \lambda \partial a_1} & \frac{\partial^2 f}{\partial \lambda \partial l_1} & \frac{\partial^2 f}{\partial \lambda \partial a_0} & \frac{\partial^2 f}{\partial \lambda \partial l_0} & \frac{\partial^2 f}{\partial^2 \lambda} \end{pmatrix}$$

5.3 Newton-Verfahren

Dieses Verfahren versucht näherungsweise Lösungen zu der Gleichung $f(x) = 0$ zu bestimmen, also die Nullstellen. Die Idee dahinter ist, in einem Ausgangspunkt der Funktion die Tangente zu bestimmen und die Nullstelle der Tangente als Näherung zur Nullstelle der Funktion zu verwenden. Der nächste Schritt ist dann die Iteration. Man führt den ersten Schritt solange aus bis eine festgelegte Schranke unterschritten ist. Die Formel für dieses Verfahren lautet:

$$x_{n+1} = x_n - \frac{f(x_n)}{f'(x_n)} \quad (7)$$

Wir wollten aber nicht die Nullstellen der Gleichung bestimmen sondern das Minimum. Deshalb versuchen wir mit dem Verfahren eine Lösung für die Gleichung $f'(x) = 0$ zu finden. Die Formel für unseren Fall lautet daher:

$$x_{n+1} = x_n - \frac{f'(x_n)}{f''(x_n)} \quad (8)$$

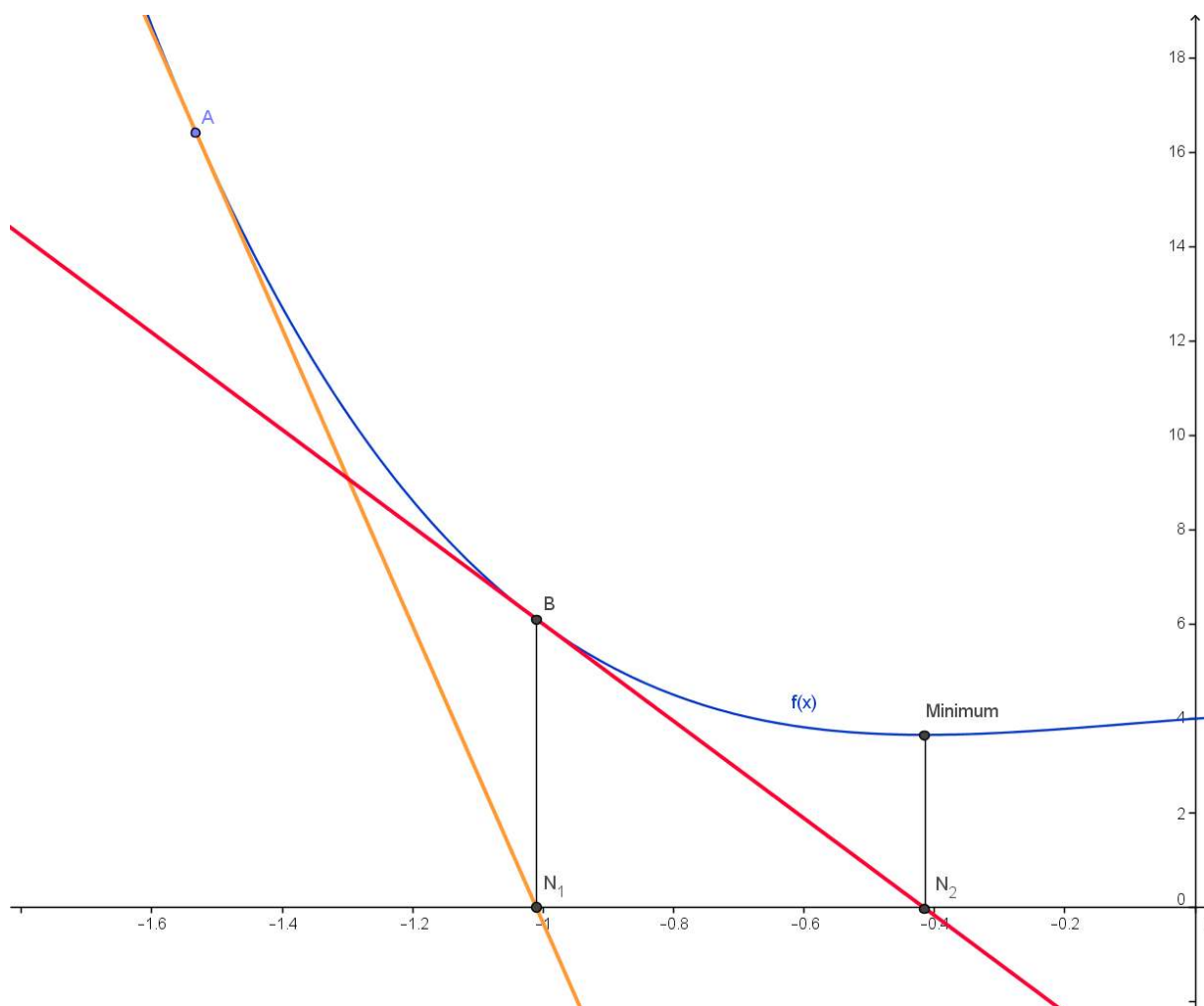


Abbildung 5: Grafische Darstellung des näherungsweisen Minimums

6 Boltzmann-Verteilung

Die Boltzmann-Verteilung wird normalerweise für Temperaturen verwendet. Aber da sich die Gesellschaft auch in ihren Strukturen ändert kann man diese Verteilung auch auf unser Modell anwenden. Diese Formel haben wir in unseren Metropolis-Hastings-Algorithmus implementiert. Dieser übernimmt mit einer gewissen Wahrscheinlichkeit zufällig berechnete Änderungen.

$$\rho = \frac{e^{-\beta \cdot H(a1,l1,m1)}}{e^{-\beta \cdot H(a,l,m)}} \quad (9)$$

β ist eine Konstante. Umso kleiner β ist, umso eher wird eine Änderung übernommen.

Falls eine Zufallszahl zwischen 0 und 1 kleiner als ρ ist dann wir die Änderung übernommen. Das heißt, dass eine Änderung eher übernommen wird, wenn das neue a,l,m ($a1,l1,m1$) $\mathcal{H}$ minimiert. Damit wir neue Werte für den Lohn, die Arbeit und die Macht bekommen haben wir die Macht untereinander verteilt, aber nicht erhöht da wir $\sum_{i=1}^n m_i = 1$ als Bedingung gesetzt haben. Den Lohn und die Arbeit haben wir auch untereinander vertauscht, aber auch erniedrigt und erhöht. Allerdings mussten wir sie bei derselben Person um denselben Wert erhöhen und erniedrigen da $\sum_{i=1}^n l_i = \sum_{i=1}^n a_i$ wir auch als Bedingung festgelegt haben. Wir haben die Ergebnisse in Abhängigkeit von γ abgebildet, da dies eine sinnvolle Darstellung ergibt. Wir haben alle Graphen in 2-D abgebildet.

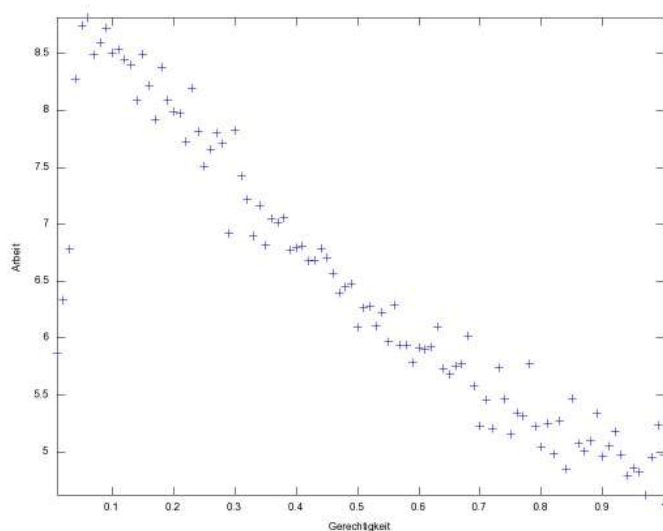


Abbildung 6: Gerechtigkeit und Arbeit

Bei dieser Abbildung (6) haben wir für $n = 10$ und für β einen kleinen Wert angenommen. Daher ergibt sich diese Streuung. Umso größer γ ist umso weniger Einfluss übt die Macht aus auf den Lohn. Bei einem kleinem γ müssen die Leute mit wenig Macht mehr arbeiten, dafür, dass sie weniger bekommen. Bei $\gamma = 1$ bekommt jeder seinen Anteil von dem was er gearbeitet hat.

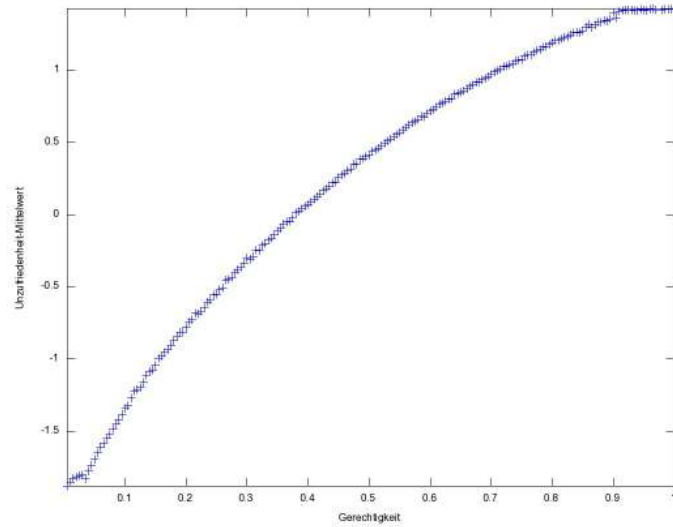


Abbildung 7: Unzufriedenheit und Gerechtigkeit

Bei dieser Abbildung (7) nahmen wir an das $n=10$ ist und $\beta = 300$ daher gibt es fast keine Streuung. Wir nahmen für $f(a, l) = -\ln(1 - a) - \ln(1)$ an, es gibt fast keine Sprünge in dieser Abbildung im Gegensatz zur folgenden Abbildung:

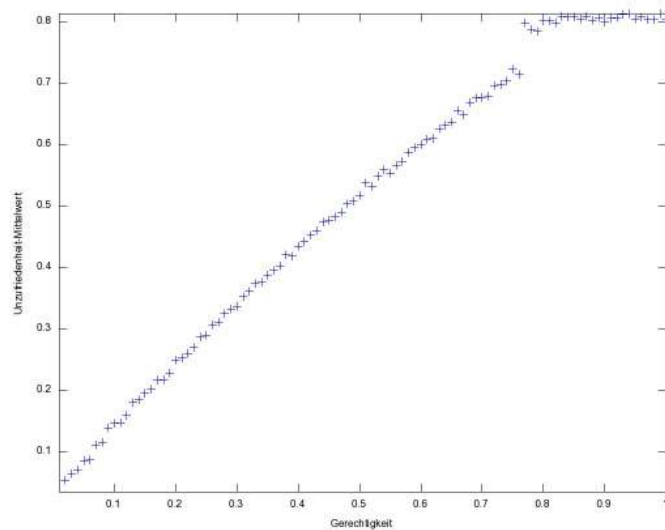


Abbildung 8: BIP in Abhängigkeit von β und γ

Hier (8) nahmen wir an für $n=10$ und $\beta = 300$ die Streuung ist sehr gering. Aber der Unterschied zur vorherigen Abbildung ist, dass wir $f(a, l) = e^{a^2-l}$ als Formel genommen haben. Dies verursacht das die Werte näher zusammen liegen, dafür gibt es aber Sprünge bei $\gamma = 0.75 : 0.8$

7 Berechnungen

Die Berechnungen wurden einerseits mit Matlab und andererseits mit GNU Octave durchgeführt. Mit Octave wurden die Berechnungen nach C++ ausgelagert um einen zusätzlichen Geschwindigkeitsvorteil zu erzielen. Zum Vergleich: Die Erzeugung der BIP-Kurve nahm in Matlab über 40 Minuten in Anspruch, mit Octave und zusätzlichen Optimierungen in C++ war die Berechnung nach 114 Sekunden erledigt.

7.1 Mittelwert der Unzufriedenheit berechnen

Den Mittelwert der Unzufriedenheit berechneten wir mit folgenden Code.

Das Zeichnen der Funktion wurde mit Octave durchgeführt:

```
for i=1:length(beta)
    for j=1:length(gamma)
        [m, a, l, mittelH(i,j)] = mittelwert(m, a, l, K, gamma(j), beta(i), L, 0);
    end
end
```

Die neuen Werte werden in C++ berechnet:

```
calculate(m,a,l,gamma,beta,1000);
for (int i = 1; i<=K; i++)
{
    calculate(m,a,l,gamma,beta,L);
    h += H1(m,a,l,gamma);
}
h /= K;
results(0) = octave_value(m);
results(1) = octave_value(a);
results(2) = octave_value(l);
results(3) = h;
return results;
```

7.2 Zufriedenheit einer Gesellschaft

Es wurde versucht die Zufriedenheit in der Gesellschaft, abhängig von β und γ , zu veranschaulichen, wobei $\gamma \in [0; 1]$ für die Gerechtigkeit in der Gesellschaft (bei $\gamma = 1$ herrscht vollständige Gerechtigkeit) und $\beta \in [0; \infty)$ für die Bereitschaft der Gesellschaft, Änderungen einzugehen, steht. Sinnvollerweise wurde das Intervall für β auf $[0; 150]$ beschränkt. Für größere β 's ändert sich der Wert kaum mehr. Je kleiner β ist, desto mehr ist die Gesellschaft bereit, Veränderungen zu machen. Es mag verblüffend wirken, dass bei der Gerechtigkeit $\gamma = 0$ die Gesellschaft am zufriedensten ist, wobei zu beachten ist, dass nur 1 Person mit der vollständigen Macht $m = 1$ die Zufriedenheit bestimmt und alle anderen irrelevant sind, da sie keine Macht haben. Auch die komplette Gerechtigkeit $\gamma = 1$ wird nie zu Zufriedenheit führen. Bei der Grafik ist zu beachten, dass der Mittelwert der Unzufriedenheit angegeben wird (blau = zufrieden, rot = unzufrieden).

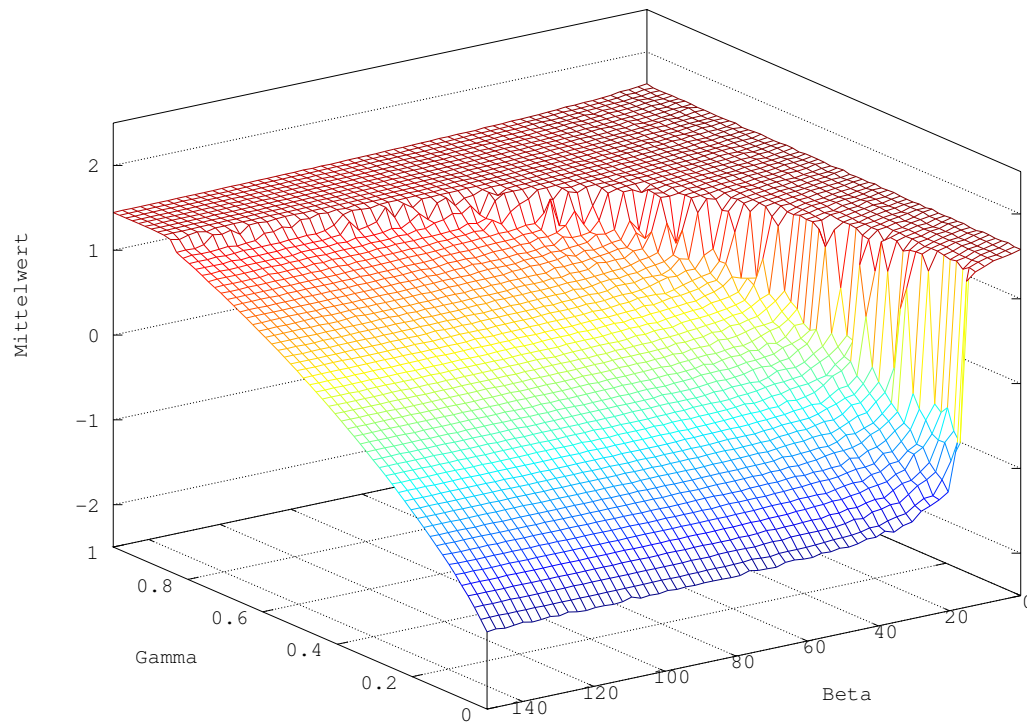


Abbildung 9: Mittlere Zufriedenheit in Abhängigkeit von β und γ

Weiters konnte das optimale γ für eine vernünftig erscheinende Bereitschaft zur Veränderung mit $\beta = 20$ und einer Population $n = 10$ bestimmt werden:

$$\gamma = 0.2642$$

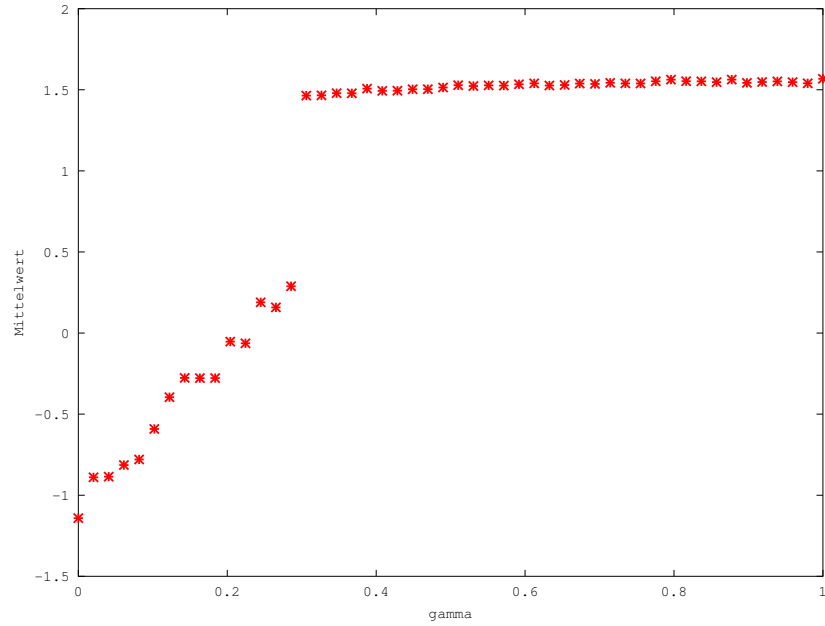


Abbildung 10: Zufriedenheit in Abhängigkeit von γ

7.3 Bruttoinlandsprodukt

Das BIP eines Landes wurde in Abhängigkeit von der Bereitschaft zur Änderung der Gesellschaft β und der Gerechtigkeit γ modelliert.

Der etwas hügelige Funktionskurvenverlauf ist darauf zurückzuführen, dass der Durchschnittswert an dieser Stelle aus zu wenigen Messungen (im konkreten Fall waren es 500), die zufällig aus allen ausgewählt wurden, errechnet wurde. Verständlicherweise wird die Funktion schöner verlaufen, wenn der Mittelwert aus mehreren Messungen errechnet wird.

$$BIP = \sum_{i=1}^n a_i \quad (10)$$

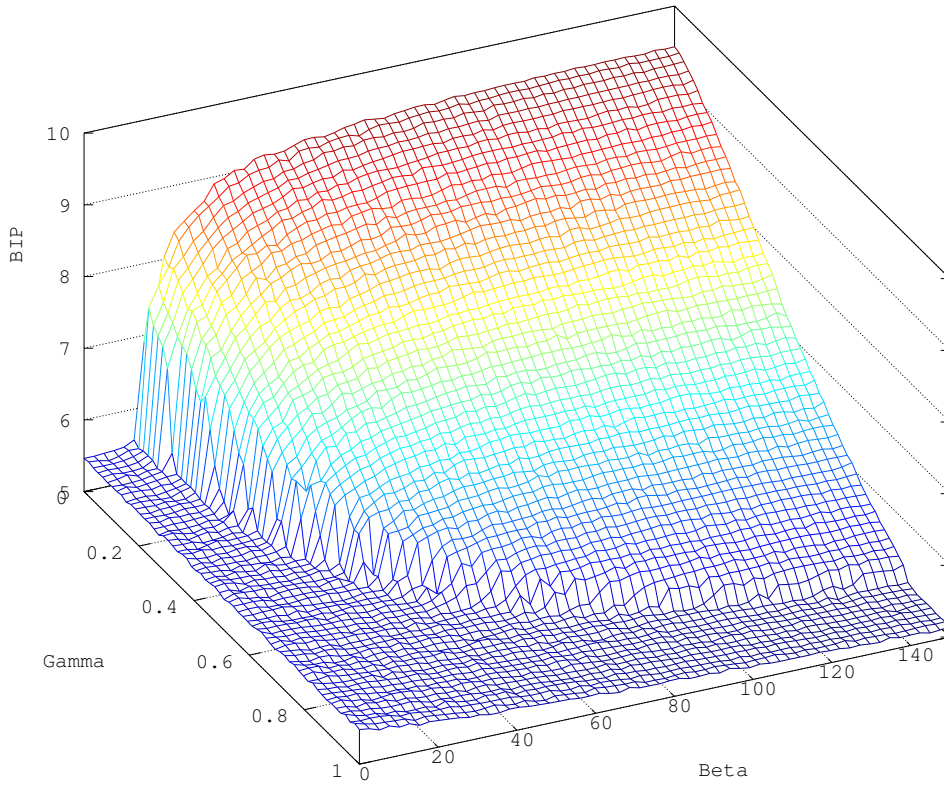


Abbildung 11: BIP in Abhängigkeit von β und γ

Das hier dargestellte BIP entspricht der Arbeitsleistung von 10 Personen ($n = 10$) zu den jeweiligen Umständen β und γ . Die Kurve zeigt, dass das höchste BIP erreicht wird, wenn gilt Gerechtigkeit $\gamma = 0$ und somit einer die ganze Macht hat. Wie schon vorher gezeigt, hat eine Person mit hoher Macht trotz wenig Arbeit einen hohen Lohn. Die Unzufriedenheit wird mit Hilfe der folgenden Formel (11) berechnet.

$$\mathcal{H} = \sum_{i=1}^n \left((1 - \gamma) \cdot m_i + \frac{\gamma}{n} \right) \cdot f(a_i, l_i) \quad (11)$$

7.4 Gini-Koeffizient

Der Gini-Koeffizient nimmt einen Wert zwischen 0 bei Gleichverteilung und 1, wenn nur eine Person das komplette Einkommen erhält, an. Zu beachten ist, dass mit Gleichverteilung eine Verteilung mit einer Varianz von 0 gemeint ist, d. h. dass jede einzelne Person genau gleich viel Einkommen erhält. Dazu werden die Löhne in die Gleichung (12) eingesetzt und der Koeffizient berechnet.

$$G_{koeff} = \frac{2 \cdot \sum_{i=1}^n i l_i}{n \cdot \sum_{i=1}^n l_i} - \frac{n+1}{n} \quad (12)$$

Entwicklung einer Fußbodenheizung

Clemens Andritsch, Christoph Heininger, Christian Kern,
Patricia Offerman, David Schöngrundner, Paul Worm
Betreuer: Mag. Martin Holler

Contents

1	Einleitung	1
2	Herleitung der Wärmeleitung	1
3	Temperaturverteilung in einem 1D-Stab	3
3.1	Verfahren zum Lösen einer Differentialgleichung	3
3.2	Programm zur Berechnung der Temperaturverteilung im Stab . .	4
3.3	Simulation für realistische Vorgänge	5
4	Wärmeveränderungen in einem Raum	7
5	Conclusio	11

1 Einleitung

Wir beschäftigen uns mit der Ausbreitung thermischer Veränderungen eines Körpers durch Wärmeleitung. Unser Ziel ist es, ein Modell zu finden und umzusetzen, das uns erlaubt, diese Ausbreitung zu beschreiben. Mit einem solchen Modell wäre es z.B. möglich, die Wärmeverteilung in einem beheizten Raum zu simulieren. So könnte man am Computer untersuchen, wie sich die Position von Heizkörpern auswirkt bzw. sogar welche Art zu heizen am effizientesten ist. Ein vereinfachtes Modell hierfür ist die sogenannte Wärmeleitungsgleichung, eine partielle Differentialgleichung. Die Lösung dieses Typs von Gleichungen ist mathematisch und physikalisch sehr interessant, da partielle Differentialgleichungen viele physikalische Vorgänge realistisch beschreiben. Wir beginnen mit einer sehr einfachen, eindimensionalen Situation, die beispielsweise erlaubt, thermische Veränderungen in einem Metallstab zu beschreiben.

2 Herleitung der Wärmeleitung

Das erste Ziel ist es, eine Formel zur Berechnung der Temperatur anhand einiger physikalischer Beziehungen herzuleiten. Dazu stellen wir uns zuerst einen Metallstab unbestimmter Länge vor. Der Stab wird in gleich große Blöcke, B_i , unterteilt. Jeder Block hat eine Temperatur u_i . Der Fluss der Temperatur von einem Block B_i zu einem benachbarten Block B_{i+1} bezeichnen wir als q_i .

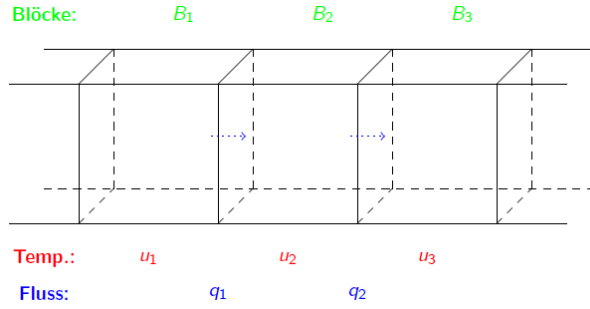


Figure 1: Temperatur im Metallstab

Bestimmt wird der Temperaturfluss durch den Temperaturgradienten zwischen den betreffenden Blöcken und einer stoffabhängigen Konstante c . In der Physik wird dieser Zusammenhang Fouriersches Gesetz genannt. Dieses beschreibt die übertragene Wärmeleistung in einem vereinfachten Modell eines festen Körpers zwischen zwei parallelen Wandflächen:

$$q_i = \frac{k}{\Delta x} A(u_i - u_{i+1}) \quad (1)$$

Nimmt man Blöcke an, deren Querschnittsfläche 1cm^2 betragen, fällt der Faktor A weg und es ergibt sich eine vereinfachte Formel:

$$\text{Fluss} = -c(\text{Temperaturgradient}) \quad (2)$$

Im besonderen Fall unseres Beispiels (Figure 1) bedeutet das:

$$q_i = -k \frac{u_i - u_{i+1}}{\Delta x} \quad (3)$$

Um die Temperatur eines Blockes zu errechnen, müssen wir desweiteren das Energieerhaltungsgesetz berücksichtigen, welches besagt, dass die Energie innerhalb eines geschlossenen Systems weder abnehmen noch zunehmen kann, also konstant bleiben muss. Daher nehmen wir auch an, dass es sich bei unserem Metallstab um ein solches geschlossenes System handelt.

Die Energie im Block B_i wird durch die spezifische Wärmekapazität des Stoffes, dessen Dichte, Temperatur und Volumen bestimmt:

$$E^i = c_p * \rho * u^i * A \quad (4)$$

Die Änderung der Energie kann man mithilfe der Temperaturflüsse berechnen:

$$E_t^i = -(q_i - q_{i-1}) * A \quad (5)$$

Kombiniert man diese Formeln, kommt man zu folgender Gleichung:

$$u_t^i = \frac{k}{c_p \rho} \left(\frac{u_{i+1} - u_i}{\Delta x} - \frac{u_i - u_{i-1}}{\Delta x} \right) \quad (6)$$

Nimmt man außerdem noch an, dass eine zusätzliche Wärmequelle vorhanden ist, wird die Formel um den Wert f_i erweitert:

$$u_t^i = \frac{k}{c_p \rho} \left(\frac{\frac{u_{i+1} - u_i}{\Delta x} - \frac{u_i - u_{i-1}}{\Delta x}}{\Delta x} \right) + f_i \quad (7)$$

3 Temperaturverteilung in einem 1D-Stab

Wir haben einen Stab, aus einem beliebigen Feststoff bestehend, welcher in einen inhomogenen Temperaturverteilungszustand versetzt wird. Wir modellieren nun die weitere Temperaturverteilung. Nach dem letzten Abschnitt kann diese wie folgt berechnet werden:

$$u_t(t, i) = \alpha \frac{u(t, i-1) - 2 * u(t, i) + u(t, i+1)}{\Delta x^2} \quad (8)$$

$u(t, i)$ beschreibt die Temperatur im Block i zum Zeitpunkt t . Um die Temperaturverteilung zu berechnen, müssen wir eine Differentialgleichung lösen. Dafür gibt es verschiedene Methoden.

3.1 Verfahren zum Lösen einer Differentialgleichung

Explizites Verfahren:

$$f(i+1) = f(i) + f'(i)\Delta t \quad (9)$$

Implizites Verfahren:

$$f(i+1) = f(i) + f'(i+1)\Delta \quad (10)$$

Crank-Nicolson Verfahren:

$$f(i+1) = \frac{f'(i) + f'(i+1)}{2} \Delta t \quad (11)$$

Ziel dieser Näherung ist es, die analytische Lösung möglichst gut zu beschreiben. Weiters müssen die Verfahren gewählt werden, dass alle Parameterkonstellationen möglich sind.

Die Einführung des Impliziten und des Crank-Nicolson-Verfahrens dient dazu, genauere Werte für unsere Approximation zu erhalten (siehe Figure 2). Da wir für weitere Funktionen bestimmte Startwerte und Konstanten variieren wollen, sollten wir zwischen verschiedenen Verfahren wählen können.

Dazu erstellen wir ein Programm, das es zulässt, zwischen den einzelnen Verfahren wechseln zu können. Dafür wurde diese Gleichung aufgestellt:

$$f(i+1) = f(i) + [f'(i) + f(1-\lambda) * f'(i+1)]\Delta t \quad (12)$$

Der Gewichtungsfaktor λ lässt sich in unserem Programm zwischen 0 und 1 variieren:

$$\begin{aligned} \lambda = 1 &\rightarrow \text{Explizites Verfahren, da } 1 - 1 = 0 \\ \lambda = 0 &\rightarrow \text{Implizites Verfahren} \\ \lambda = 0.5 &\rightarrow \text{Crank-Nicolson Verfahren} \end{aligned}$$

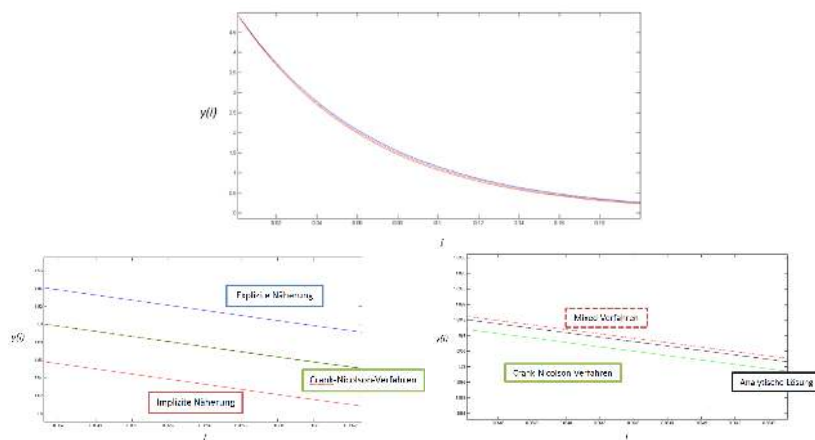


Figure 2: Vergleich der Verfahren

3.2 Programm zur Berechnung der Temperaturverteilung im Stab

Die eben beschriebene Methode setzen wir nun ein, um die Temperaturverteilung in einem Stab beliebigen Materials zu berechnen. Dazu haben wir ein Programm in Matlab implementiert, das die Differentialgleichung 8 numerisch löst. Das Programm benutzt folgende Variablen und Parameter:

Zeitparameter

T ... Länge der Simulation (in Sekunden)
 N ... Anzahl der Zeitschritte
 $\Delta t = \frac{T}{N}$... Länge der Zeitschritte

Ortsparameter

L ... Länge des Stabes
 B ... Anzahl der Blöcke (Unterteilung des Stabes in Abschnitte)
 $\Delta x = \frac{L}{B}$... Länge eines Blockes

Randtemperatur

$a(1, t)$... Temperatur am linken Rand des Stabes zum Zeitpunkt t
 $e(1, t)$... Temperatur am rechten Rand des Stabes zum Zeitpunkt t

Variablen

$u(t, i)$... Temperatur im Block i zum Zeitpunkt t

Materialparameter

α ... Temperaturleitfähigkeit des Stabes

In jedem Iterationsschritt wird folgendes Zeitupdate ausgeführt:

$$u(t+1, i) = u(t, i) + [u_t(t, i) * \lambda + (1 - \lambda) * u_t(t+1, i)] * \Delta t \quad (13)$$

wobei

$$u_t(t, i) = \frac{\alpha}{x^2} * [u(t, i-1) - 2u(t, i) + u(t, i+1)] * \Delta t \quad (14)$$

Die benötigt die Lösung eines linearen Gleichungssystems, die in Matlab effizient durchgeführt wird.

3.3 Simulation für realistische Vorgänge

In Figure 3 oben sind 2-dimensionale Plots der Temperaturentwicklung in Stäben aus Tannenholz und Eisen. Sie werden auf einer Seite mit 100° Celsius konstant erhitzt und auf der anderen mit 20° Celsius abgekühlt.

In einer weiteren Simulation (Figure 3 mitte links) wird ein Eisenstab an beiden Enden konstant mit 100° Celsius erhitzt. Die Anfangstemperatur des Stabes beträgt 20° Celsius. Nach ca. 5,5 Stunden hat der Stab an jedem Punkt eine Temperatur von 100° Celsius.

In Figure 3 mitte rechts hat ein Stab anfangs 100° Celsius in der Mitte und kühlt bei einer Temperatur von 0° Celsius stark ab.

Annahme in Figure 3 unten links: Ein 1-dimensionaler Raum wurde zu lange gelüftet und hat eine Temperatur von 0° Celsius. Die Wand gegenüber dem Fenster hat konstante 20° Celsius. Ziel ist es nun die Temperatur des Raumes auf 20° Celsius zu bringen. Ergebnis: Wir heizen nach 50 minütigem Heizen mit 50° Celsius nur mehr mit 20° Celsius und bleiben somit konstant bei 20° Celsius.

In Figure 3 unten rechts: Ein mit 10° Celsius temperierter Raum wird mit $20 \frac{W}{m^2}$ auf 24° Celsius erwärmt und letztlich mit einer Ausgangsleistung von $-20 \frac{W}{m^2}$ wieder abgekühlt.

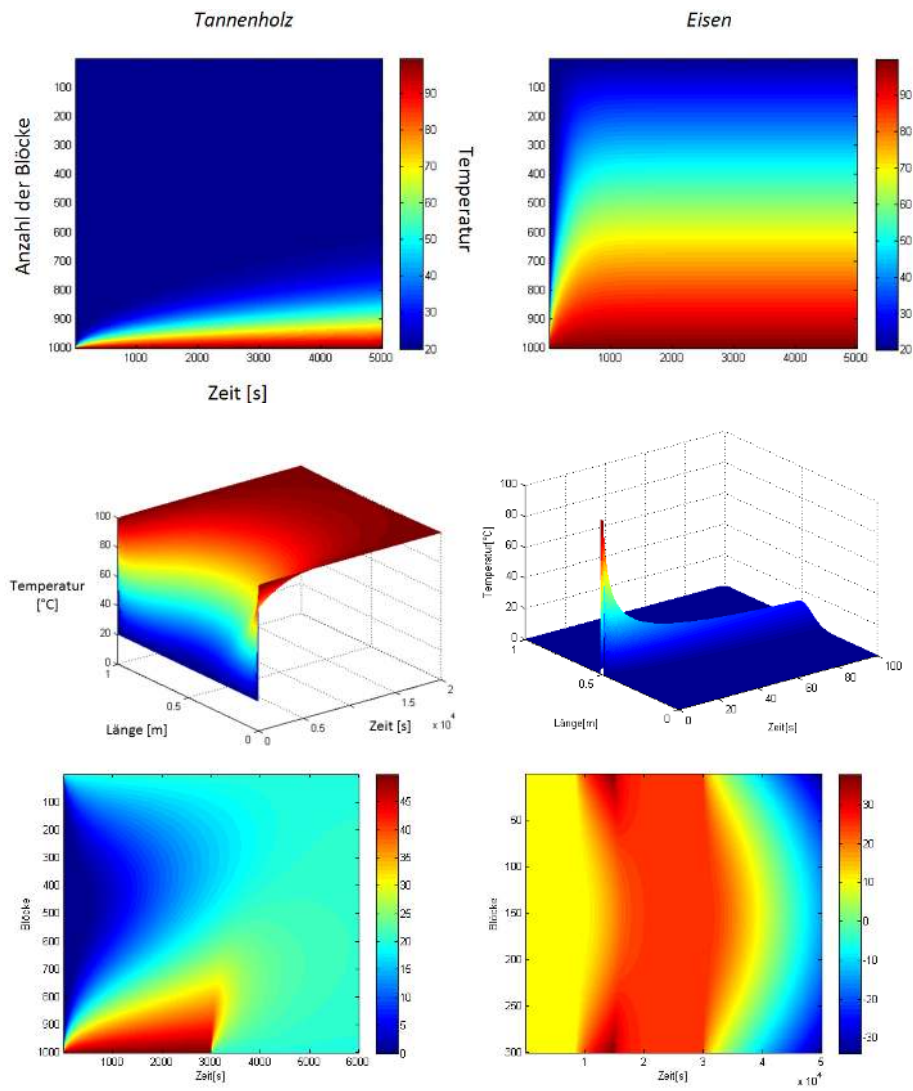


Figure 3: Oben: Temperaturentwicklung in Tannenholz- und Eisenstab. Mitte links: Ein Stab wird erwärmt. Mitte rechts: Eisenstab wird abgekühlt. Unten links: 1D Raum. Unten rechts: Erwärmen und Abkühlen

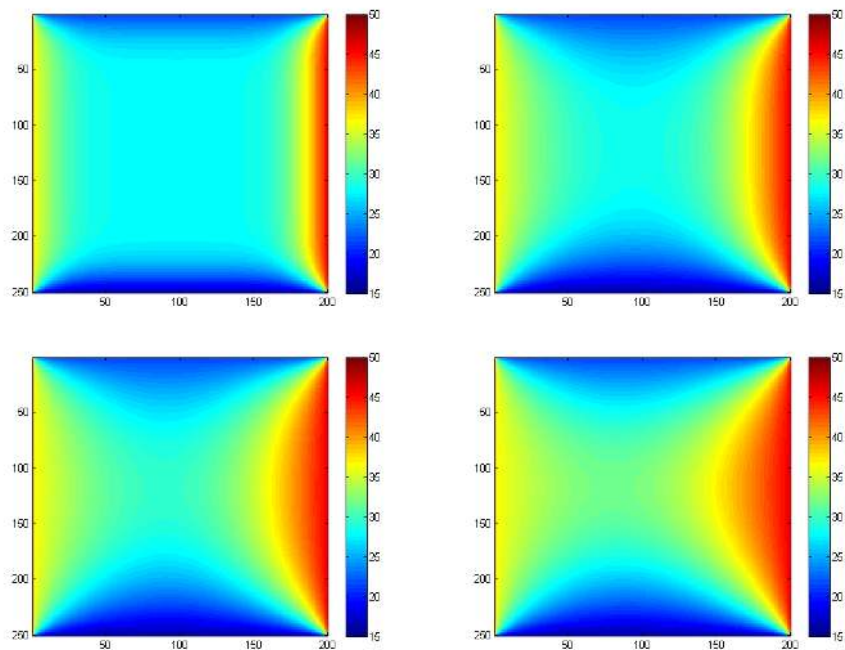


Figure 4: Veränderung der Raumtemperatur: Oben links: Nach 4 Stunden. Oben rechts: Nach 8 Stunden. Unten links: Nach 12 Stunden. Unten rechts: Nach 16 Stunden

4 Wärmeveränderungen in einem Raum

Das für den Stab verwendete System ist auch für unsere zweite Aufgabe, die Wärmeänderungen in einem 2-dimensionalen Raum darzustellen, in abgeänderter Form zu benutzen. Der gravierendste Unterschied besteht dabei darin, dass nun jeder Block vier anstatt zwei Nachbarblöcken hat, und deshalb auch mehr Wärme Flüsse unter den Blöcken existieren. Dazu kommt, dass wir nun vier verschiedene Randbedingungen erzeugen können müssen. Ein Beispiel: Die Temperaturen der Wände betragen oben 25° , rechts 50° , unten 14° und links 37° , die anfängliche Raumtemperatur ist 28° . Die Temperaturen verändern sich wie in Figure 4.

Eine weitere Aufgabe, die wir uns zum Ziel gesetzt haben war, dass wir eine gute (Boden-) Heizung schaffen, sodass die Temperatur trotz kälterer Außenwände möglichst konstant bleibt, und so gut es geht auch auf den ganzen Raum aufgeteilt ist. Die Angabe schaut dabei so aus: Wir befinden uns in einem Schlafzimmer, die obere und die linke Wand haben jeweils 20° , die untere 14° und die rechte Wand 18° . Die Grundtemperatur ist 19° , und soll so gut es geht gehalten werden. Dazu sollen wir eine Bodenheizung designen, die 10 Watt-Stunden an Energie in die Luft abgibt, und die diese Aufgabe möglichst gut löst. Wir testen zwei verschiedene Heizungen, eine typische Bodenheizung in Fall 1 und die von uns vorgeschlagene Bodenheizung in Form eines L's.

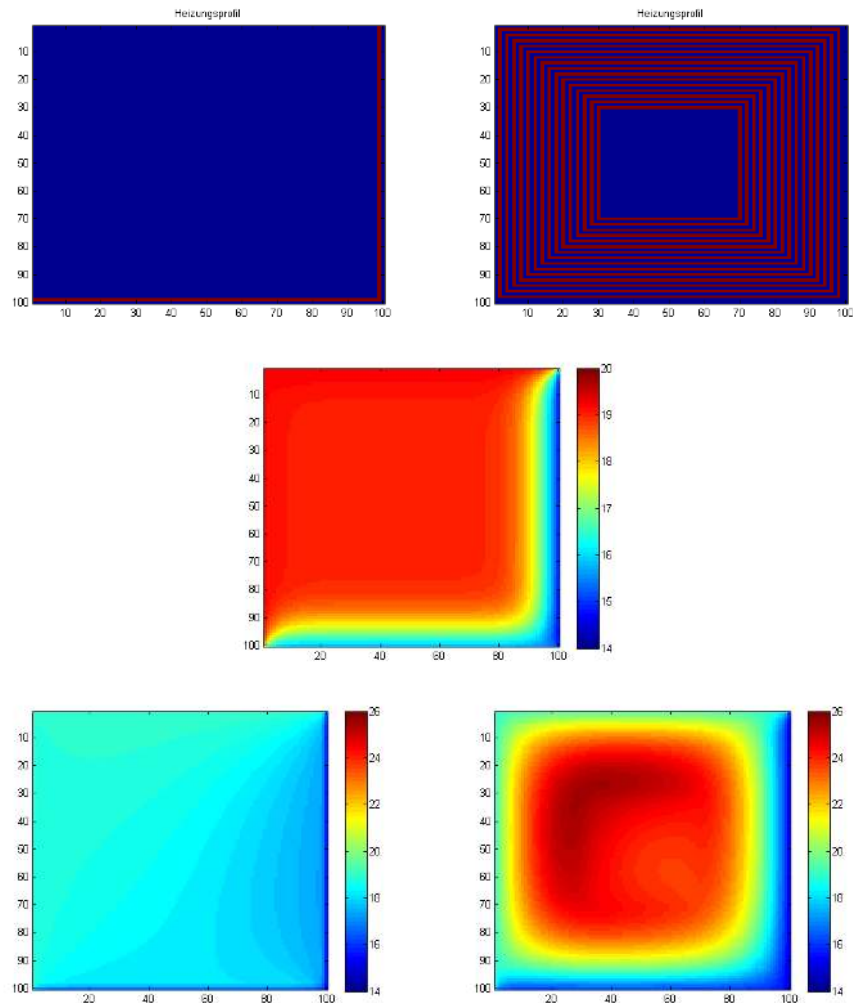


Figure 5: Vergleich zwischen L-Heizung und Standard-Heizung

Ad Figure 5: Die ersten beiden Bilder zeigen die Verlegung der Fußbodenheizung, das dritte die Ausgangslage vor dem Einschalten der Heizung, und die letzten beiden die Ergebnisse. Dabei sieht man, dass die normale Fußbodenheizung sehr unkonstant und sehr lokal das Zimmer auf eine zu hohe Temperatur anheizt. Im Gegensatz dazu schafft unsere L-Heizung eine angenehmere und konstantere Raumtemperatur.

Als letzte Herausforderung stellen wir uns, die Temperaturen abhängig zur Tageszeit zu machen, und dann auch noch eine Heizung zu erstellen, die ebenfalls davon abhängig ist. Zu Mittag, wenn die Wände wärmer sind, muss man natürlich weniger Energie in die Heizung stecken. Die Temperaturänderung der Wand haben wir mit einer Sinuskurve simuliert und die Heizung genau gegengleich mit einem Cosinus. Danach haben wir diese Heizung mit einer anderen L-Heizung verglichen, die den ganzen Tag über immer die gleiche Temperatur abgibt. Siehe Figures 6 - 8 für die Raumtemperatur zu verschiedenen

Tageszeiten. Die von uns erstellten Simulationen, in Form von Videos, haben ergeben, dass die von der Tageszeit abhängige extrem effektiv ist, da sie über den Tag verteilt für ziemlich konstante Temperatur sorgt. Im Gegensatz dazu sorgt die Heizung, die konstant heizt, zu Mittag für zu hohe Temperaturen, und am Abend kann sie die Kälte nicht so effektiv aus dem Raum halten.

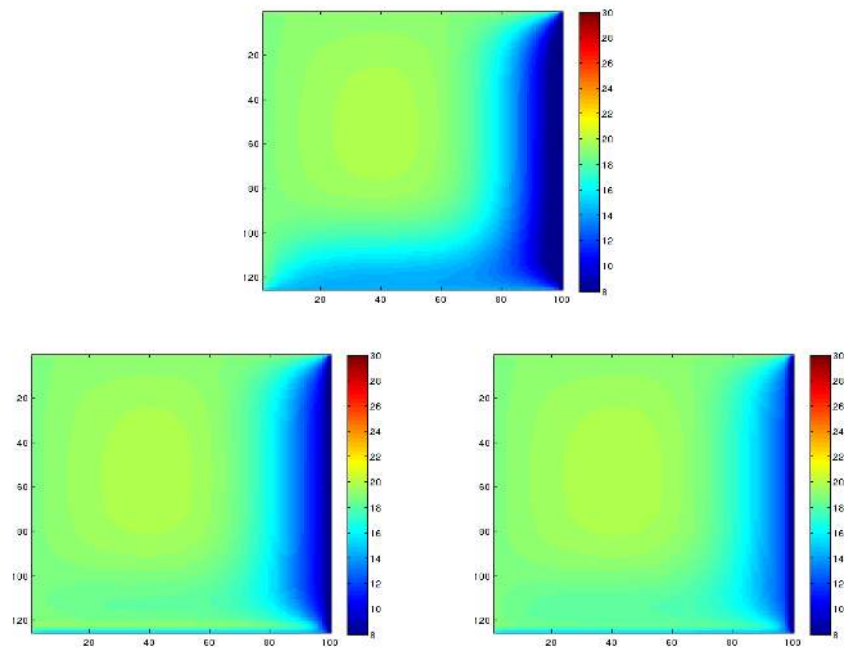


Figure 6: Temperaturverteilung um 4:48 Uhr. Oben: keine Heizung, unten links: Dauerheizung, unten rechts: Zeitheizung

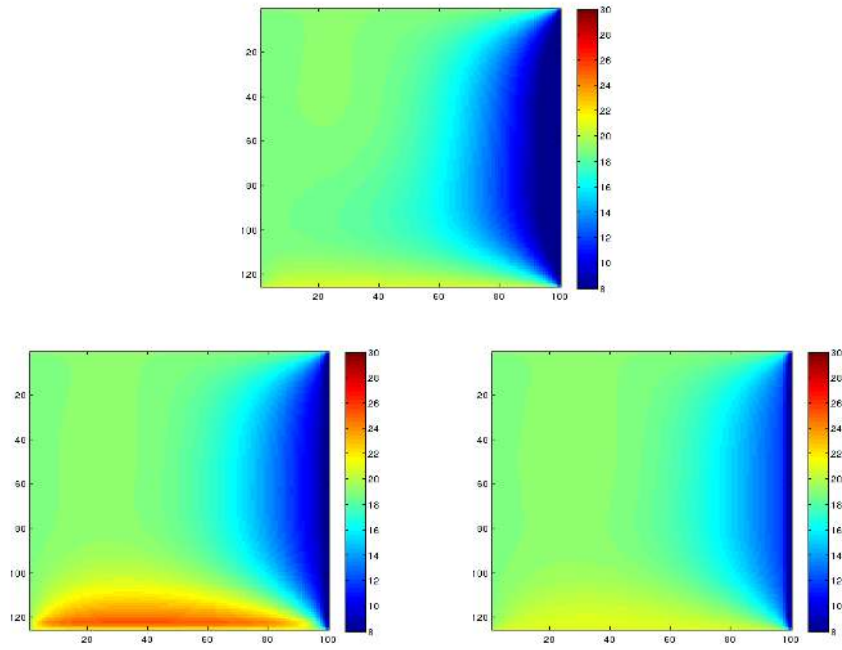


Figure 7: Temperaturverteilung um 12:00 Uhr Oben: keine Heizung, unten links: Dauerheizung, unten rechts: Zeitheizung

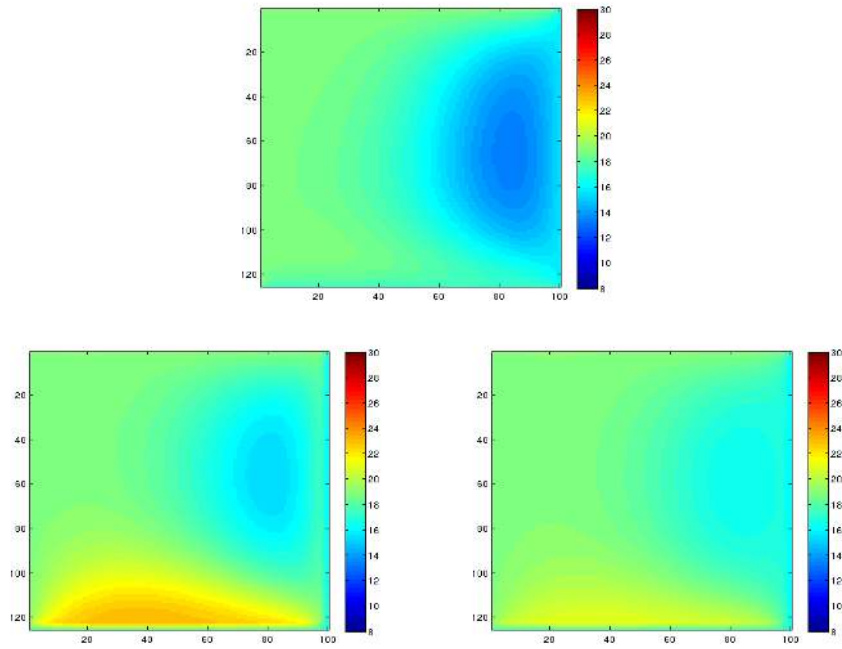


Figure 8: Temperaturverteilung um 17:24 Uhr Oben: keine Heizung, unten links: Dauerheizung, unten rechts: Zeitheizung

5 Conclusio

In dieser Woche haben wir einiges gelernt. Wir konnten ein neues, in vielen Hinsichten sehr praktisches Programm kennen lernen, und erhielten so neue Programmierfähigkeiten. Während den Simulationen haben wir einige interessante Erkenntnisse bekommen, zum Beispiel, dass man das Programm schon richtig einstellen muss, um den Raum in annehmbarer Zeit zu erwärmen beziehungsweise abzukühlen. Durch stundenlange Diskussionen, ob unser Programm nun tatsächlich realistisch ist, welche Faktoren wir vergessen haben, oder wo genau der Fehler im Programm liegt, bekamen wir einen tieferen Einblick in die Situation, als sich selbst der Gruppenleiter vorher gedacht hat. Außerdem halfen uns die verschiedenen Interessen, die wir innerhalb der Gruppe hatten, um jeden Aspekt, sowohl logisch, mathematisch, so wie auch physikalisch zu interpretieren und zu hinterfragen.

Projekt Kunst und Fotografie

Erstellung und Beurteilung eines Fotomosaiks

Dr. Stephen Keeling



TeilnehmerInnen:

Tim Sagaster

Philipp Welsch

Edin Dedic

Margit Galsterer

Alexandra Forster

Alexander Eber

Inhaltsverzeichnis:

1.Problemstellung

2.Problembearbeitung

2.1 Erste Gedanken

2.2 Ähnlichkeiten

Farbähnlichkeit

Kovarianz

2.3 Probleme

2.4 Lösungen

2.5 Grundlagen

3.Ergebnis= Code:

3.1 Zielbild

3.2 Einzelbilder

3.3 Optimierung

3.4 Ähnlichkeit

3.5 Fotomosaik

1.PROBLEMSTELLUNG:

Gegeben seien Einzelfotos und ein Zielfoto. Mit den Einzelfotos soll ein Fotomosaik erstellt werden, das in gewisser Weise dem Zielfoto ähnlich ist. Obwohl die Aufgabe leicht zu beschreiben ist, steckt der Teufel im Detail.

Wie viele Einzelfotos sind für die Aufgabe geeignet? Was bedeutet ähnlich und welches Ähnlichkeitsmaß soll verwendet werden, um den Platz für ein Einzelfoto im Mosaik zu finden? Vielleicht passt genau ein Einzelfoto am besten an einem bestimmten Platz im Mosaik, aber dieses Einzelfoto könnte an mehreren Plätzen am besten passen. Welcher Platz soll für dieses Einzelfoto ausgewählt werden? Angenommen werden alle Einzelfotos genau einmal verwendet. Wie berechnet man die optimale Verteilung der Einzelfotos im Mosaik zumindest annäherungsweise in durchführbarer Zeit? Man kann in erster Linie annehmen, dass alle Einzelfotos gleich große viereckige Plätze im Mosaik besetzen, aber was ändert sich, wenn diese viereckigen Plätze nicht unbedingt gleich groß sein müssen oder wenn die Einzelfotos deformiert werden dürfen?

Nachdem alle dieser Fragen beantwortet worden sind und ein Fotomosaik erstellt worden ist, stellt sich die Frage trotzdem, was ist eine gute oder eine schlechte Lösung? Wie beurteilt man das Ergebnis objektiv? Vielleicht könnte sogar ein nicht schönes Zielfoto eine gewisse Schönheit durch das Umwandeln in ein Fotomosaik gewinnen, aber ein solches Phänomen könnte von vielen Faktoren abhängen. Wie misst man den Gewinn an Schönheit quantitativ und ermöglicht so eine Beurteilung eines Ergebnisses?

Für dieses Projekt soll ein Computerprogramm geschrieben werden, mit dem ein Fotomosaik durch verschiedene Methoden erstellt werden kann. Zusätzlich sollen Kriterien entwickelt werden, nach denen Ergebnisse beurteilt werden können.

2.PROBLEMBEARBEITUNG

2.1 Erste Gedanken

Als Erstes muss man die Bilder einlesen, um das Zielbild aus den anderen Bildern zu erstellen. Aus diesem Grund teilen wir das Zielbild in mehrere Zellen auf. Um einen Vergleich zu bekommen müssen wir das Bild in die richtige Größe formatieren. Das gleiche müssen wir auch mit den Einzelbildern machen, wobei die Pixelgröße des Zielbildes ein natürliches Vielfaches der Einzelbilder sein muss. Danach stellen wir einen Vergleich auf, der uns erlaubt, ein wahrheitsgetreues Abbild zu erzeugen. Man muss die Ähnlichkeiten berechnen und die Einzelbilder den Zellen zuteilen. Danach sollte das Fotomosaik problemlos entstehen können. Es stellten sich schon die ersten Fragen: Nach welchem Prinzip teilen wir die Einzelbilder zu? Sollten wir durch das Bild wandern und jeder Zelle nach der Reihe dieses Bild geben, das sie benötigt? Was passiert dann jedoch mit den letzten Bildern? Weiter machten wir uns Gedanken, nach welchen Ähnlichkeiten sie verglichen werden sollten und welche der Möglichkeiten uns als wichtigste erscheint.

2.2 Ähnlichkeiten

Es gibt mehrere Möglichkeiten die Zielzellen mit den Einzelbildern zu vergleichen. Unserer Meinung nach ist das wichtigste Kriterium die Farbähnlichkeit. Wir berechnen den Durchschnitt der RGB-Werte, und stellen eine Ähnlichkeit zwischen den Zielzellen und den Einzelbildern auf. Die Zielzelle nimmt sich das Einzelbild, welches ihr am ähnlichsten ist. Ein weiteres Ähnlichkeitsmaß ist die Kovarianz, welche sich auf die Struktur der Bilder bezieht. Unserer Erfahrung nach ergibt eine Berechnung nur nach der Kovarianz kein erkennbares Fotomosaik. Daher müssen beide Kriterien miteinbezogen werden, wobei unsere Gruppe entschieden hat, dass die Farbe um einiges wichtiger ist als die Kovarianz.

Farbähnlichkeit:

Man berechnet den durchschnittlichen RGB-Wert und bekommt einen Wert, der bei höherer Ähnlichkeit geringer ist. Dieser Wert wird von allen Einzelbildern mit der zu bestimmenden Zielzelle verglichen und das Einzelbild, das den geringsten Wert hat, wird der Zielzelle zugewiesen. Danach wird dieses Bild von der Liste gelöscht, und ist für die nächste Zielzelle nicht mehr verfügbar.



Kovarianz:

Die Kovarianz bezieht sich auf die Struktur. Diese Ähnlichkeit wird eingeführt, dass strukturgleiche Bilder, mit unterschiedlichen Farben, nicht als ganz unterschiedliche Bilder angesehen werden. Diese Bilder können dann auch mit einer komplett gegengleichen Farbe (siehe Bild nebenan) ersetzt werden. Sie wird durch das skalare Produkt der einzelnen Stücke eines Bildes berechnet.



Erklärung:

0	0
1	1

0	1
0	1

Die Zahlen werden von links nach recht und oben bis unten aufgeschrieben und der Mittelwert wird subtrahiert.

Diese Werte werden multipliziert und zusammengerechnet (Skalares Produkt)

0	0	-1/2	-1/2	1/4	
0	1	-1/2	1/2	-1/4	= 0
1	0	1/2	-1/2	-1/4	
1	1	1/2	1/2	1/4	

Ist das Ergebnis 0, passen die Strukturen der beiden Bilder nicht zusammen.

0	0
1	1

1	1
0	0

0	1	-1/2	1/2	-1/4	
0	1	-1/2	1/2	-1/4	= -1
1	0	1/2	-1/2	-1/4	
1	0	1/2	-1/2	-1/4	

In diesem Fall sind sich die Strukturen ähnlicher, auch wenn die Farbe nicht ganz übereinstimmt. Die Zahl die sich ergibt darf nicht 0 sein, + 1 oder -1 spielt keine Rolle.

2.3 Probleme:

Das erste Problem ist, dass wir Unmengen (mehr als 14.000) an Einzelbildern brauchen um ein erkennbares Fotomosaik zu erstellen.

Ebenfalls mussten wir die Einzelbilder gerecht verteilen. Es könnte sein, dass mehrere Zielzellen das selbe Einzelbild in Anspruch nehmen wollen. Deshalb benötigten wir eine gerechte Aufteilung durch Erstellen von Prioritätslisten. Am Ende sollten alle Zellen zufrieden sein. Jede Zelle erstellt eine Liste mit den ähnlichsten Bildern, wobei das ähnlichste am Anfang steht. Wir brauchten eine stabile Lösung.

Ein weiteres Problem, welches uns am Erfüllen unseres Auftrages zu hindern versuchte, war die benötigte Rechenkapazität. Um einfache Bilder (1024x768) zu erzeugen brauchten wir mit unserem Code bereits 5-10 Minuten. Wir mussten etwas ändern und vereinfachen, um die Rechenzeit zu verkürzen.

Ein Problem ist auch, dass wir jedes Bild nur einmal verwenden dürfen. Dies konnten wir natürlich erst ausführen, als wir genug Bilder zusammen hatten, um ein Bild in kleine Teile aufzuspalten, damit man noch etwas erkennen kann.

2.4 Lösungen:

Um die vielen Bilder zu bekommen haben wir zwei Videos zerlegt, und sind jetzt um 14000 Bilder reicher.

Unsere erste Lösung war, dass man die erste Zeile von ersten bis zum letzten Bild durchgeht und jedes sich der Reihe nach ein Einzelbild aussucht. Leider wird so das Bild am Anfang schöner, weil die erste Zelle die meisten Einzelbilder hat, aus denen es wählen kann, und das letzte Bild nur mehr wenige oder sogar nur mehr eine Möglichkeit hat, und vielleicht gezwungen wird, ein Bild zu nehmen, welches überhaupt nicht dazu passt.

Unsere derzeitige Lösung wäre der Befehl „randperm“. Mit diesem Befehl wird zufällig eine Zelle ausgewählt, eine Liste erstellt, und dann ein Einzelbild ausgewählt, welches am besten passt. Dies hat den Vorteil, dass wenn ein Bild aus sehr großen Flächen besteht, wird eine Seite nicht schöner als die andere, weil sie vorher zum Auswählen dran ist. Somit kann man eine optimale Verteilung erzielen. Jedoch stellte sich die Frage noch immer, ob jede Zelle damit ausreichend zufriedengestellt werden kann.

Beispiel:

Wir haben unser Logo 3mal durch unseren Code laufen lassen. 1mal mit dem Code, dass die Zeilen der Reihe nach ersetzt werden, und 2mal mit dem Zufallsbefehl. Hier kann man unsere Ergebnisse sehen:



Die Lösung für die gerechte Verteilung der Bilder war das Heiratsproblem. Die Zielzellen stellen Anträge auf die gewünschten Einzelbilder wobei diese dann entscheiden sollten, in welche Zelle sie lieber gehen würden. Es gibt mehrere Runden die diese Bilder durchlaufen. Immer wieder werden Anträge gestellt und auch die eigentlich schon verkoppelten, die noch einen Antrag erhalten, müssen neu entscheiden. Es sollte am Ende eine stabile Lösung geben, das heißt, jedes Einzelbild und jede Zielzelle sollen maximal zufriedengestellt werden und es darf nicht auftreten, dass z.B.: A lieber B und B lieber A hätte, als jene, mit denen sie bereits verbunden sind.

Beispiel Heiratsproblem:

Burschen: A,B,C,D

Mädchen1,2,3,4

Jeder hat seine Prioritäten:

A 1 3 2 4

B 4 2 3 1

C 2 3 1 4

D 1 4 2 3

1 A C B D

2 B D A C

3 A B D C

4 C A D B

Runde 1:A und D stellen Anträge an Mädchen 1, B an Mädchen 4, C an Mädchen 2. Somit muss Mädchen 1 zwischen A und D auswählen. Mädchen 1 hätte lieber A, somit muss D erneut einen Antrag stellen. Dies darf aber nur bei Mädchen sein, welchen er noch keinen gestellt hat. Dies geschieht in der zweiten Runde (siehe zweite Spalte unten). D stellt somit einen Antrag an die zweite in seiner Liste: Mädchen 4. Diese hat schon einen von B in der

vorherigen Runde und muss sich nun zwischen den beiden entscheiden, wobei sie D bevorzugt. Jetzt muss Bursche B wieder einen Antrag stellen, da er wieder frei ist. Das nächste Mädchen in seiner Liste, welchen er noch keinen Antrag stellte ist Mädchen 2. Es beginnt die dritte Runde (dritte Spalte). Diesmal hat das Mädchen 2 von C und B einen Antrag, muss sich entscheiden und wählt. Nun ist Bursch C wieder alleine, stellt einen neuen Antrag (Runde vier, 4. Spalte) an Mädchen 3 und alle sind vergeben.

1 (A) D	1 (A)	1 (A)	1 (A)
2 (C)	2 (C)	2 C (B)	2 (B)
3	3	3	3 (C)
4 (B)	4 B (D)	4 (D)	4 (D)

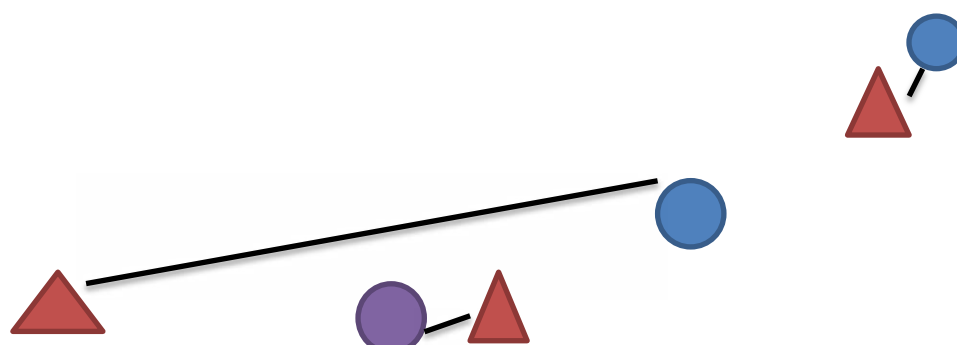
Es entsteht eine stabile Lösung was bedeutet das jeder zufriedengestellt worden ist und kein Bursche lieber ein Mädchen hätte und dieses den Burschen. Wir konnten dieses Modell in unseres einbauen indem wir die Burschen als Zielzellen und die Mädchen als Einzelbilder betrachteten. Wir erstellten in allen Gruppen sogenannte Prioritätslisten in denen die Ähnlichkeiten der Zielzellen mit den Einzelbildern aufgelistet waren. Es musste nur noch das Heiratsproblem auf die oben angeführte Weise gelöst werden. Somit dürften alle Teile des Zielbildes zufrieden sein, aber auch die jeweiligen Einzelbilder.

Das Heiratsproblem schien anfangs eine perfekte Lösung zu sein, jedoch entsprachen die Mosaikbilder nicht ganz unserer Vorstellung. Außerdem muss man sehr viele Runden durchlaufen bis jede Zelle maximal zufrieden gestellt wird, was bedeutet, dass sich die Rechenzeit noch um einiges verlängert. Es fiel uns dabei eine weitere Möglichkeit ein, die Bilder aufzuteilen.

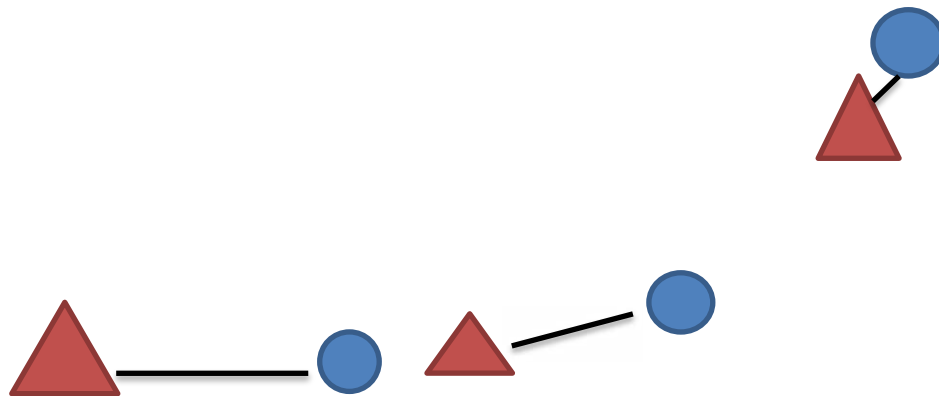
Beim Heiratsproblem werden die Zellen verbunden die am nächsten beieinander liegen. Daher könnte es sein das eine Zielzelle schlussendlich ein Bild erhält, welches sehr weit, bezüglich der Ähnlichkeit, entfernt ist.

 Zellen des Zielbildes

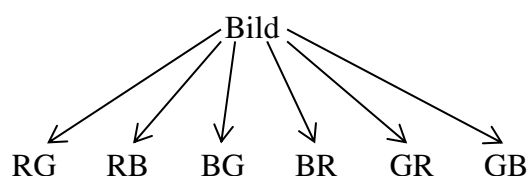
 Einzelbilder



Eine neue Idee war es, ein Bild mit globaler Zufriedenheit zu kreieren. Die Bilder, welche die größte minimale Ähnlichkeit haben, wählen zuerst ein Einzelbild. Wie man in der Darstellung darunter erkennen kann, dürften alle Zellen zufrieden sein und nicht zu weit abweichende Einzelbilder erhalten was infolge ein schöneres Bild ergibt.



Um das Problem mit der zu langen Rechenzeit zu lösen vereinfachten wir den Code. Es dauert lange, bis jede Zielzelle mit jedem Einzelbild verglichen wurde. Daher sollten die Ähnlichkeiten zwischen Zelle und Einzelbild nur noch berechnet werden, wenn sie sich auch gleichen. Unser Gedanke war es, die Einzelbilder nach den RGB-Werten in Ordner zu teilen. Wir erstellten hierfür sechs verschiedenen Ordner: RG, RB, BG, BR, GB, GR. Im Ordner RG befinden sich alle Bilder mit den meisten Rotanteilen und den zweitmeisten Grünanteilen, im Ordner RB die Bilder mit den meisten Rotanteilen und zweitmeisten Blauanteilen und so weiter. Es teilt sich sozusagen nach dem Verhältnis von Rot-, Blau- und Grünanteilen in die jeweiligen Ordner auf. Danach mussten die Zielzellen nur noch in den jeweiligen Gruppen, die ihnen ähnlich waren, nach Einzelbilder suchen.



Wir vereinfachten dieses Verfahren, indem wir die Bilder nicht in Ordner sondern in verschiedene Klassen teilten. Es hat den Vorteil dass keine Ordner mehr existieren und damit wird die Entstehung des Mosaiks beschleunigt.

Ein weitere Punkt der die Zeit verkürzen soll: Wir speicherten schon bei dem letzten Durchgang die Variablen in unserem Code und konnten bei den folgenden Durchgängen darauf zugreifen. Dies kann aber nur funktionieren, wenn wir die Maße vom Zielbild und die Größe einer Zelle immer gleich lassen. Sonst muss diese Datei gelöscht werden, und noch einmal neu vom Programm geschrieben werden muss. Deshalb ist es ratsam immer mit den gleichen Maßen zu arbeiten, um schnelle Ergebnisse zu erzielen. Natürlich können wir nicht immer nur die gleichen Bilder verwenden, aber nur um zu probieren, ob der Code funktioniert und keine Fehlermeldung ausspuckt, ist diese Lösung optimal.

Die Lösung vom Problem mit der einmaligen Verwendung ist, dass man 2 unterschiedliche Listen erstellt (z.B.: Lin und Lout). Dies ist aber nur ein Vorschlag zur besseren Orientierung. In der Liste Lin befinden sich am Anfang sämtliche Bilder, aus denen ich mein Zielbild machen will. Am Anfang ist die Liste Lout leer (also „Lout=[]“). Sobald ich ein Bild verwende wird dieses Einzelbild in die Liste Lout verschoben, bis am Ende die Liste Lin leer ist, und die Liste Lout voll ist. Diese Listen müssen aber nicht erstellt werden, denn sie dienen nur dem Programm, damit es weiß welche Bilder es noch verwenden darf.

2.5 Grundlagen:

Um das Bild einzulesen muss man den Befehl `„Z=imread('Bildname.jpg');“` ein. In diesem Codestück ist Z einfach eine Variable, mit der wir unser Bild von nun an bezeichnen.

Mit diesem Befehl `„[Nx,Ny,Nz]=size(Z);“` werden für Nx,Ny und Nz die Größe definiert. Das machen wir nur aus dem Grund, damit wir automatisch eine Änderung der Größe bei Veränderung des Bildes haben. Der `“;“` am Ende jeder Zeile bedeutet, dass das Programm das Ergebnis dieser Zeile nicht ausgibt, sondern nur für sich behält und errechnet. Dies sollte man machen, um nicht von Zahlenkolonnen überschüttet zu werden, die zum Beispiel beim Einlesen eines Bildes entstehen.

`„Zr=imresize(Z(:,1),[Nx,Ny]);“`: Mit diesem Befehl kann man ein ganzes Bild auf eine bestimmte Größe bringen. In diesem Fall wird das Bild aber nicht vergrößert oder verkleinert, sondern die rot,grün,blau Anteile getrennt. Die anderen Befehle die dazu vonnöten sind lauten `(„Zg=imresize(Z(:,2),[Nx,Ny]); Zb=imresize(Z(:,3),[Nx,Ny]);“`.

Die folgenden Befehle sind zur Einteilung des Bildes in unterschiedliche Areale. Diese können später, wie in unserem Fall, durch die Einzelbilder ersetzt werden: `„Mx=32;My=32;Mz=3; Lx=round(Nx/Mx); Ly=round(Ny/My); Nx=Lx*Mx; Ny=Ly*My;“`. In diesem Fall haben wir die Variable `“L”` als Anzahl der entstandenen Zellen eingesetzt. Dies kann aber je nach Belieben geändert werden.

Um ein Bild am Ende auszugeben, muss man den Befehl `„imagesc(F);“` verwenden. Hier verwenden wir noch einen Strichpunkt, weil das Bild normalerweise bei dieser Ausgabe noch verzerrt ist. Dem wirken wir sofort mit diesen zwei Befehlen entgegen: `„axis image; axis off“`. Man sieht den fehlenden Strichpunkt nach dem off und weiß daher, dass nach diesem Befehl etwas ausgegeben wird. Wenn man nur `„axis image“` schreibt, werden noch Achsen in die x- und die y-Richtung angezeigt.

Das einzige Problem ist jetzt noch die Einlesung der Einzelbilder. Das kann man entweder automatisch machen, oder man erstellt selber eine Liste. Wir bevorzugen die automatische Methode. Aber die Lösung zu diesem Problem werdet ihr noch später in unserem Code sehen. Für die automatische Einlesung haben wir einige Zeit benötigt, weil wir und einige passende Befehle suchen mussten.

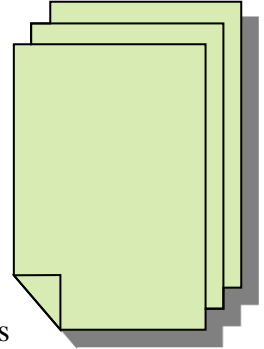
3. ERGEBNIS= CODE:

3.1 Zielbild

Im ersten Teilstück des Codes wird unser Zielbild eingelesen, formatiert und in unterschiedliche Areale eingeteilt, die sich in einen Stapel umwandeln. In diesem Stapel haben wir 4 Koordinaten: x, y, r RGB und die Stelle in diesem Stapel.

3.2 Einzelbilder

Der nächste Teil liest automatisch die alle Fotos in einem Zielordner, in diesem Fall C:\Users\Beispiel\Desktop\Fotos', erstellt eine Liste, in diesem Fall „liste.txt“, mit allen .jpg Dateien, die sich in diesem Ordner befinden, mit Ausnahme des Zielbildes, das sich ganz am Ende dieser Liste befinden muss. Sollte es das nicht sein, muss man es umbenennen (z.B.: zzzZielbild.jpg). Dann werden diese Bilder neu formatiert, und sind dann in diesem Beispiel 32*32 Pixel groß. Dies kann aber variiert werden, und muss an das Zielbild angepasst werden. Dafür verwenden wir eine „for-Schleife“, die für alle Bilder von 1 bis N. N muss vorher definiert werden.



```
system('cd C:\Users\Beispiel\Desktop\Fotos');
system('dir /b *.jpg>liste.txt');
Nx=32; Ny=32;
fid=fopen('liste.txt');
FILES=textscan(fid,'%s');
FILES=FILES{1};
fclose(fid);
FILES=FILES(1:end-1);
N=size(FILES,1);
E=zeros(Nx,Ny,3,N);
for n=1:N
    Et=imread(char(FILES(n)));
    Er=imresize(Et(:,1),[Nx,Ny]);
    Eg=imresize(Et(:,2),[Nx,Ny]);
    Eb=imresize(Et(:,3),[Nx,Ny]);
    E(:,1,n)=Er;
    E(:,2,n)=Eg;
    E(:,3,n)=Eb;
End
```

3.3 Optimierung

Um den ganzen Prozess noch schneller werden, nutzen wir das Heiratsproblem. Die ganzen Bilder werden in Ordner eingeteilt. Dies erfolgt automatisch und das Programm teilt die Bilder nach dem Farbdurchschnitt ein. Dann wird auch das Zielbild in solche Ordner eingeteilt. Da dies unter den selben Kriterien geschieht, kann sich dann eine Zielzelle ein Bild aussuchen, dann im Gegenstück des Ordners liegt. Dadurch muss nicht für jedes Bild eine Liste erstellt werden, auf der auch Bilder stehen, die ganz unterschiedliche Eigenschaften haben, sondern nur einander ähnliche Bilder. Dieser Teil kann nur dann eingesetzt werden, falls das ganze Ähnlichkeitssystem nur auf Farbähnlichkeit basiert.

3.4 Ähnlichkeit

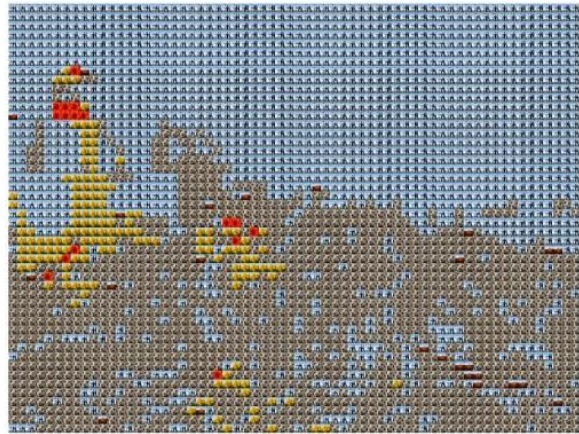
Im dritten Abschnitt wird das Ähnlichkeitsmaß berechnet, und die Zielzelle sucht sich das ihr am Ähnlichsten Einzelbild. Dieses Einzelbild wird dann aus der Liste gelöscht, und kann deshalb nicht mehr verwendet werden. Die Ähnlichkeit kann dann auch noch gewichtet werden. Man kann zum Beispiel die Farbähnlichkeit viel wichtiger als die Kovarianz (z.B.: $A1*9$ und $A2$).

3.5 Fotomosaik

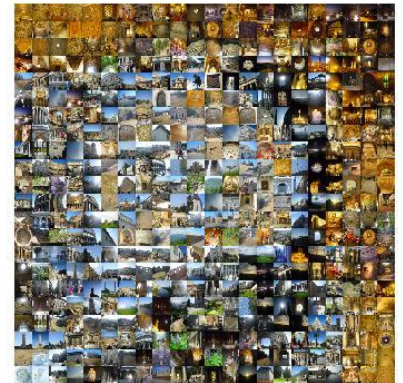
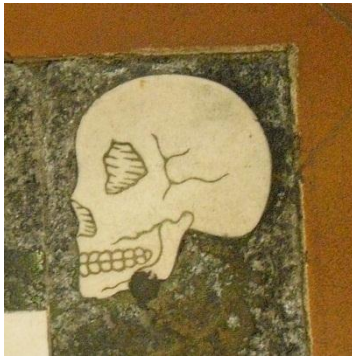
Im nächsten Abschnitt unseres Codes wird das Fotomosaik zusammengefügt, von einem Stapel in ein normales Bild formatiert, das Bild angezeigt, entzerrt und schlussendlich ausgegeben.

```
Ft=F;
F=[]; n=0;
for i=1:Lx
    Fj=[];
    for j=1:Ly
        n=n+1;
        Fj=[Fj,Ft(:, :, n)];
    end
    F=[F;Fj];
end
F=uint8(F);
imagesc(F);
axis image;
axis off
```

Erste Versuche:



Das war einer unserer ersten Versuche, in der wir einfach jedes Bild öfter verwendet haben. In diesem Bild kann man nur 5 Bilder erkennen, weil diese Bilder die ähnlichsten zu dem Zielbild waren.



Hier kann man einen anderen Versuch sehen. Das linke Mosaik wird nur über die Farbähnlichkeit bestimmt und das rechte Mosaik nur über die Kovarianz. Wie man sehen kann ist im linken Bild der Kopf dem Zielbild ähnlicher, aber im rechten ist der Umriss von der orangen Umrandung sichtbar.



Dieses Bild wurde nach dem Prinzip des Heiratsproblems gelöst, ist jedoch eines der ersten Versuche. Es widerspiegelt die Zweifel, die wir bereits hatten. Es werden Bilder an Zellen vergeben, wo sie nicht wirklich hinpassen (gelbe Bilder im Himmel), weil diese Zellen auf jene Einzelbilder zurückgreifen müssen. Es gibt zwar eine stabile Lösung, jedoch gibt es starke Abweichungen vom Originalbild.

Am Ende hatten wir viele verschiedene Möglichkeiten, ein Fotomosaik zu erstellen. Also mussten wir uns für eine, und zwar die, die uns die besten Zielbilder lieferte, entscheiden. Wir hatten 2 Codes, welche schlussendlich schöne Bilder erzeugen konnten.

Zum einen den Code, welcher zufällig Zellen im Zielbild auswählt und denen dann die gewünschten Einzelbilder, der Reihe nach, gibt. Die letzten die ausgewählt werden erhalten „schlechtere“ Bilder, was aber nicht wirklich auffällt, da sie am ganzen Bild verteilt sind.

Unsere zweite Wahl war der Code, der die Lösung durch das Heiratsproblem enthält. Mit genügend kleinen und ähnlicheren Bildern kann man das Endbild trotzdem annähernd erkennen.