

Institut für Neuro- und Bioinformatik  
Universität zu Lübeck

# Computergestützte Analyse von biologischen und bioinspirierten Signalverarbeitungs- und Wahrnehmungsprozessen

Inauguraldissertation

zur Erlangung der Doktorwürde (Dr.rer.nat.)  
der Universität zu Lübeck  
- Technisch-Naturwissenschaftliche Fakultät -

vorgelegt von

**Diplom-Informatiker**  
**Amir Madany Mamlouk**  
aus Berlin - Spandau

Lübeck, im März 2008



Betreuer: Prof. Dr.rer.nat. Thomas Martinetz  
Institut für Neuro- und Bioinformatik

1. Gutachter: Prof. Dr.rer.nat. Thomas Martinetz  
2. Gutachter: Prof. Dr.rer.nat. Bernd Fischer  
3. Gutachter: Prof. Dr. Felix Wichmann

Datum der Einreichung: 04. März 2008  
Datum der Prüfung: 16. Mai 2008

gez. Prof. Dr.rer.nat. Jürgen Prestin  
Dekan der Technisch-Naturwissenschaftlichen Fakultät

Meinen Eltern.





## Eidesstattliche Erklärung

Hiermit versichere ich, die vorliegende Dissertation selbständig und ohne unerlaubte Hilfe angefertigt zu haben. Bei der Verfassung der Dissertation wurden keine anderen als die im Text aufgeführten Hilfsmittel verwendet.

Ein Promotionsverfahren zu einem früheren Zeitpunkt an einer anderen Hochschule oder bei einem anderen Fachbereich wurde nicht beantragt.

Lübeck, den 4. März 2008

Amir Madany Mamlouk



# Zusammenfassung

Computerassistierte Systeme sind besonders in der Neurowissenschaft unverzichtbar geworden. Die vorliegende Arbeit widmet sich ganz allgemein dem maschinellen Lernen von biologischen und biologienahen Zusammenhängen, und beschäftigt sich exemplarisch mit zwei Problemen aus diesem Anwendungsfeld: Der Elektrodennavigation im Gehirn und der computergestützten Analyse von Wahrnehmungsräumen, hier am Beispiel des Farbensehens und, besonders ausführlich, des Geruchssinnes.

Es wird vorgestellt, wie mit Hilfe der Independent Component Analysis (ICA) auf simulierten Elektrodendaten das umliegende Nervengewebe gemessen und lokalisiert werden kann. Diese komplett am Computer entstandene Studie zeigt auf, wie operative Eingriffe zur Tiefenhirnstimulation zukünftig mittels einer ICA-basierten Elektrodennavigation um wesentliche Informationen zu dem elektrodennahen Gewebe ergänzt werden könnten. Die ICA zeigt hervorragende Ergebnisse auf dem simulierten Elektrodenlayout – besonders im Bereich des Spike Sorting, der Zuordnung von Spikes zu ihren erzeugenden Zellen.

Weiterhin wird ein Analyserahmen vorgestellt, mit dessen Hilfe Wahrnehmungsprozesse auf der Basis von verbalen Beschreibungen analysiert werden können. Es wird ein umfangreiches Experiment vorgestellt, in dessen Verlauf Probanden mehr als 25.000 paarweise Bewertungen von 90 Farbstimuli und 16 Farbbezeichnern abgegeben haben. Auf der Basis dieser Daten wird detailliert beschrieben, wie man aus diesen Daten mit Hilfe von Hauptkomponentenanalyse (PCA) und Multidimensional Scaling (MDS) nicht nur den Merkmals- sondern auch den zugehörigen Wahrnehmungsraum ableiten kann. Es wird am Beispiel des Farbensehens demonstriert, dass die Interpretation des Merkmalsraumes oftmals ein subjektives Verständnis des untersuchten Systems voraussetzt und wir daher besonders bei unverständenen Wahrnehmungsprozessen eine objektive Interpretation des Systems benötigen.

Es wird gezeigt, dass die auf den gleichen Daten basierende Analyse des Wahrnehmungsraumes genau diesen Ansprüchen genügt und als Ergebnis die objektiv berechnete Darstellung der Beschreibungsdaten liefert. Dieselbe Analyse wird auf den komplexen und hochdimensionalen Wahrnehmungsraum der Gerüche angewendet und die Ergebnisse mittels Kohonen-Karten in eine zweidimensionale Darstellung gebracht. Auf diesen Karten ist es erstmals möglich, Geruchsqualitäten zu quantifizieren.

Schließlich bemüht sich diese Arbeit um einen Brückenschlag zwischen dem Merkmalsraum des Riechens (also der sensorischen Aktivität des Riechkolbens) und dem Wahrnehmungsraum (also der Kategorisierung verschiedener Stoffe durch Geruchsqualitäten) durch eine Kopplung der vorgestellten Geruchskarten mit experimentellen Aktivitätsbildern des Riechkolbens. Hierbei erweisen sich mit Hilfe von linearen SVMs berechnete lineare Klassifikatoren, sogenannte Decision Images (DI), als nützliche Werkzeuge. Es wird gezeigt, wie mittels der DIs der kombinatorische Kode des Riechens stückweise in seine klassenspezifische Anteile zerlegt werden kann. Auf diese Art und Weise können einzelne Aktivitätsbilder in ihre funktionalen Module zerlegt werden. Dies entspricht klassenspezifischen „natürlichen“ Wörtern des Geruchskodes.

Es wird auf synthetischen Daten demonstriert, wie einfache kombinatorische Kodierungsmodelle mit Hilfe der Decision Images sichtbar gemacht werden können. Die Ergebnisse dieser Arbeit legen jedoch nahe, dass der kombinatorische Kode der Geruchswahrnehmung sich auf die eindeutige Kodierung jeder einzelnen Substanz – im Sinne eines Fingerabdruckes – beschränkt.

Schließlich wird dargestellt, dass es einen grundlegenden Unterschied gibt zwischen den Decision Images für chemische Substanzeigenschaften und denen für Geruchsqualität. Dies scheint mit der funktionellen Trennung des piriformen Cortex zu korrespondieren, der ebenfalls auf Basis der Bulbusaktivität sowohl feinste Unterschiede zwischen zwei Molekülen erkennen wie auch chemisch unterschiedliche Duftstoffe mit einer identischen Geruchsqualität assoziieren kann. Unabhängig von der endgültigen Erklärung für die hier beobachteten Effekte eröffnet die Datenanalyse unter Verwendung der DIs ganz neue Möglichkeiten, die Aktivität im Riechkolben zu interpretieren.

# Inhaltsverzeichnis

|                                                                         |            |
|-------------------------------------------------------------------------|------------|
| <b>Zusammenfassung</b>                                                  | <b>iii</b> |
| <b>1 Einleitung</b>                                                     | <b>1</b>   |
| <b>Teil I: Biologische Systeme und Maschinelles Lernen</b>              | <b>7</b>   |
| <b>2 Multielektroden-Navigation mittels ICA</b>                         | <b>9</b>   |
| 2.1 Einleitung . . . . .                                                | 9          |
| 2.1.1 Mehrkanal-Mikrosondenaufnahmen von Zellaktivität . . . . .        | 10         |
| 2.1.2 Eine Neural-Cocktail-Party - ICA für Spike Sorting . . . . .      | 11         |
| 2.1.3 GENESIS Simulation des Hippocampus CA-3 . . . . .                 | 14         |
| 2.2 ICA auf Sondendaten des CA-3 . . . . .                              | 17         |
| 2.2.1 PCs und ICs der simulierten Multielektrodensignale . . . . .      | 18         |
| 2.2.2 Signalzuordnung und Zelllokalisierung . . . . .                   | 18         |
| 2.2.3 Spike Sorting und Spike Detection mittels ICA . . . . .           | 24         |
| 2.3 Diskussion . . . . .                                                | 27         |
| <b>Teil II: Wahrnehmungsprozesse aus Sicht des Maschinellen Lernens</b> | <b>29</b>  |
| <b>3 Der Farbwahrnehmungsraum</b>                                       | <b>31</b>  |
| 3.1 Einleitung . . . . .                                                | 31         |
| 3.1.1 Sprache und Wahrnehmung . . . . .                                 | 33         |
| 3.1.2 Merkmalsräume vs. Wahrnehmungsräume . . . . .                     | 33         |
| 3.1.3 Grundlagen des Farbensehens . . . . .                             | 34         |
| 3.2 Das „Color Perception Project“ . . . . .                            | 37         |
| 3.2.1 Farben und Farbbeschreiber . . . . .                              | 37         |
| 3.2.2 Durchführung der Studie . . . . .                                 | 38         |
| 3.2.3 Ergebnisse . . . . .                                              | 42         |
| 3.3 Diskussion . . . . .                                                | 56         |
| <b>4 Der Geruchswahrnehmungsraum</b>                                    | <b>57</b>  |
| 4.1 Einleitung . . . . .                                                | 57         |
| 4.1.1 Grundlagen des Riechens . . . . .                                 | 58         |
| 4.1.2 Geruchskarten . . . . .                                           | 61         |
| 4.1.3 Geruchsbeschreiber . . . . .                                      | 64         |
| 4.2 Datenanalyse des Geruchswahrnehmungsraumes . . . . .                | 66         |
| 4.2.1 Dimensionen des Riechraumes . . . . .                             | 68         |

|                                                        |                                                                  |            |
|--------------------------------------------------------|------------------------------------------------------------------|------------|
| 4.2.2                                                  | Kohonen-Karten zur Projektion von 32D in 2D . . . . .            | 69         |
| 4.2.3                                                  | Karten des Geruchswahrnehmungsraumes . . . . .                   | 72         |
| 4.2.4                                                  | Visualisierung von stoffklassenspezifischen Geruchsqualitäten    | 74         |
| 4.3                                                    | Diskussion . . . . .                                             | 75         |
| <b>Teil III: Der kombinatorische Kode des Riechens</b> |                                                                  | <b>77</b>  |
| <b>5</b>                                               | <b>Rezeptororganisation im Riechkolben</b>                       | <b>79</b>  |
| 5.1                                                    | Einleitung . . . . .                                             | 79         |
| 5.1.1                                                  | Sensorische Verschaltung im Riechkolben . . . . .                | 80         |
| 5.1.2                                                  | Deoxyglucose-Mapping . . . . .                                   | 81         |
| 5.1.3                                                  | Klassen und Klassenbeschreibungen . . . . .                      | 86         |
| 5.2                                                    | Uptake-Pattern des Riechkolbens . . . . .                        | 91         |
| 5.2.1                                                  | Chemische Eigenschaften und ihre Repräsentanz im Bulbus .        | 91         |
| 5.2.2                                                  | Das Riechalphabet und natürliche Riechwörter . . . . .           | 93         |
| 5.2.3                                                  | Geruchsqualitäten und ihre Repräsentanz im Bulbus . . . .        | 98         |
| 5.2.4                                                  | Statische Signifikanz der Klassifikation mittels Decision Images | 102        |
| 5.3                                                    | Diskussion . . . . .                                             | 105        |
| <b>6</b>                                               | <b>Kodierungsmodelle für Aktivierungsmuster im Bulbus</b>        | <b>107</b> |
| 6.1                                                    | Einleitung . . . . .                                             | 107        |
| 6.1.1                                                  | Kodierungsmodelle . . . . .                                      | 108        |
| 6.1.2                                                  | Synthetische Modellierung . . . . .                              | 110        |
| 6.2                                                    | Kodierungsmodelle auf den Uptake-Geruchsdaten . . . . .          | 114        |
| 6.2.1                                                  | Multiple kombinatorische Geruchskodes . . . . .                  | 114        |
| 6.2.2                                                  | Substanzspezifische Unterräume . . . . .                         | 115        |
| 6.2.3                                                  | Chemische Eigenschaften vs. Geruchsqualitäten . . . . .          | 122        |
| 6.3                                                    | Diskussion . . . . .                                             | 124        |
| <b>7</b>                                               | <b>Zusammenfassung und Ausblick</b>                              | <b>127</b> |
| <b>A</b>                                               | <b>Methoden und Algorithmen</b>                                  | <b>135</b> |
| A.1                                                    | Multidimensional Scaling . . . . .                               | 135        |
| A.2                                                    | Self-Organizing Maps . . . . .                                   | 137        |
| A.3                                                    | DoubleMinOver Lernalgorithmus . . . . .                          | 139        |
| A.4                                                    | Stratified Cross-Validation . . . . .                            | 140        |
| <b>B</b>                                               | <b>Karten und Beschreiber</b>                                    | <b>141</b> |
| B.1                                                    | Übersetzungen . . . . .                                          | 145        |
| B.1.1                                                  | chemische Klassen . . . . .                                      | 145        |
| B.1.2                                                  | Geruchsbeschreiber . . . . .                                     | 146        |
| <b>Literaturverzeichnis</b>                            |                                                                  | <b>147</b> |

# 1 Einleitung

Wenn ein Mensch an einem Erdbeereis leckt und in der Folge eine Erdbeerempfindung hat, wie ist diese Wahrnehmung in seinem Kopf repräsentiert? Schmeckt es eventuell nach Erdbeere, wenn man nur an der richtigen Stelle seines Gehirns leckt?

*Vermutlich nicht.*

Aber dennoch ist die geistige Fähigkeit, seine Umgebung wahrzunehmen, Dinge zu erkennen und zu begreifen eine einmalige und sehr faszinierende Eigenschaft von Lebewesen, welche auch heute noch sowohl die Naturwissenschaften wie auch die Philosophie gleichermaßen beschäftigt [89]. Die Forschung versucht seit Jahrzehnten künstliche Intelligenz auf Computern zu erzeugen und verzweifelt in vielerlei Hinsicht an den selbstgesteckten Zielen. Mobile Roboter können sich zwar unterdessen tatsächlich autonom bewegen [96], dennoch liegt der Tag noch in weiter Ferne, an dem ein Roboter ins Büro des Forschers tritt und sagt:

„Ich habe in Deinem Regal ein Schachspiel gefunden und es aufgebaut.  
Wenn Du Lust hast, können wir eine Partie spielen.“

Es ist unstrittig, dass man Computer Schach spielen lassen kann [83]. Die Fiktion hierbei findet sich vielmehr in einem Roboter, der nicht nur zufällig ein Schachspiel findet, sondern dieses auch als solches erkennt und anschließend aus dem Regal holt, um es aufzubauen.

*Warum können Maschinen so etwas eigentlich nicht?*

Wer kann sich schon einen Roboter vorstellen, der morgens im Berufsverkehr z.B. auf dem Weg zur Arbeit zufällig einen Nachbarn in der überfüllten U-Bahn entdeckt? Aber es ist genau das, was wir von einem *intelligenten* System erwarten: Die *kognitive* Fähigkeit, sich ein internes Weltmodell aufzubauen und die *kommunikative* Fähigkeit, sich mit seiner Umwelt auszutauschen!

Daher hat schon Alan Turing in seinem legendären Test vorgeschlagen, künstliche Intelligenz im Gespräch mit Menschen zu testen [92]: Können Menschen nach einem fünfminütigen Gespräch (z.B. im Chat) nur mit 70%iger Zuverlässigkeit zwischen Mensch und Maschine unterscheiden, so handelt es sich um eine *intelligente* Maschine.

Turing hat prognostiziert, dass es spätestens im Jahre 2000 Computer mit einem solchen Potential geben würde. Dies hat sich nicht bewahrheitet. Der im Jahre 1990 mit 100.000 US\$ dotierte Loebner-Preis für das erfolgreiche Bestehen des Turing-Tests ist noch immer ausgeschrieben<sup>1</sup>. Dabei hatte Turing noch nicht einmal den mindestens ebenso schwierigen Teil des Problems in seine Zielformulierung aufgenommen: Die Wahrnehmung und Verarbeitung einer komplexen Umwelt.

Man kann sich natürlich auf den Standpunkt stellen, dass beide Beispiele Unsinn seien und Maschinen eh nie „zur Arbeit fahren“ müssen und vermutlich auch keine Nachbarn haben, die sie zufällig treffen werden, aber es ist interessant einzusehen, dass der Roboter dies auch nicht *könnte*. Und es ist spannend zu verstehen, wie der Mensch sich ganz grundsätzlich in überfüllten und unübersichtlichen Momenten so gut und so zuverlässig zurechtfinden kann.

Die Wissenschaft ist sich der Schwere dieser Aufgabe mit den Jahren bewusst geworden und verwendet den Terminus *Künstliche Intelligenz* (KI) mittlerweile umsichtiger. Es gibt zwei grundsätzliche Forschungszweige, die sich mit den Kernthemen *Leben* und *Lernen* der klassischen KI auseinandersetzen:

Einerseits gibt es das sogenannte *Artificial Life*, in dem nicht der Anspruch besteht, das echte biologische Leben *nachzubauen*, sondern vielmehr einfache Systeme zu entwerfen, die sich wie echtes Leben *verhalten*. Häufig verwendet man hierzu sogenannte Agentensysteme, also viele gleichartige Programme oder auch real existierende Roboter, welche durch Interaktion mit anderen Agenten und der Umwelt ihr Verhalten planen und verändern können. In solchen Simulationen ist es möglich, auch komplexere Phänomene wie evolutionäre Strategien und soziale Interaktionen nachzubilden.

Der andere Zweig versucht die existierenden *intelligenten* biologischen Systeme nachzubilden und dadurch besser zu verstehen. Die etwas weichere Variante ist die *Neuroinformatik*, in der mit künstlichen Neuronalen Netzen versucht wird, kognitive Prozesse, insbesondere das Lernen, von gehirngesteuerten Lebewesen mathematisch abzubilden. Die *Neuroinformatik* hat ein breites Anwendungsfeld in der Mustererkennung und Datenanalyse gefunden, da sich die Lernverfahren und Netze im kleinen Maßstab leicht auf Computern umsetzen lassen. Die Disziplin der *Computational Neuroscience* hingegen versucht, Neurone und biologische Systeme so detailliert wie möglich abzubilden und zu simulieren, mit dem hehren Ziel, irgendwann einmal das Gehirn in seiner Funktionalität komplett zu verstehen.

Die *Neuroinformatik* versucht also mit ihren Methoden die komplexe Musterverarbeitung des Menschen zu formalisieren und auf dem Computer nachzubilden, während es im Bereich der *Computational Neuroscience* eher darum geht, biologische Prozesse exakt nachzubauen, um sie am Computer funktional zu zerlegen und besser zu begreifen.

---

<sup>1</sup><http://www.loebner.net/Prizel/loebner-prize.html>



---

Im Rahmen dieser Arbeit werden immer wieder Methoden und Ansätze aus der *Neuroinformatik* verwendet. Dennoch verlieren wir die biologischen Systeme nie ganz aus dem Blick. Die Daten, die hier mit bio-inspirierten Methoden analysiert werden, stammen in der Mehrheit aus Ansätzen der *Computational Neuroscience*.

Inhaltlich lässt sich diese Arbeit in drei Teile gliedern:

### ***Teil I: Biologische Systeme und Maschinenlernen***

In dem ersten Teil möchte ich ein Projekt vorstellen, in dem es in erster Linie um das Maschinenlernen – also die *Neuroinformatik* – geht. Da hier methodische Grundlagen gelegt werden, die auch für die weiteren Kapitel relevant sind, habe ich mich dazu entschlossen, dieses an den Anfang dieser Arbeit zu stellen.

Kapitel 2 behandelt die Analyse von Mikroelektrodenableitungen, bei denen Neuronenaktivität mit hochauflösenden Mehrkanalsystemen in einem lebenden Tier oder Menschen erfasst werden. Ein solcher Datenstrom umfasst schnell mehrere Megabyte, die umgehend durch automatisierte Erkennungs- und Analysesoftware handhab- und speicherbar gemacht werden müssen.

Da wir im lebenden Gehirn jedoch keine Aussage darüber treffen können, welches Neuron wann feuern wird, ist eine Evaluation solcher Vorverarbeitungssoftware sehr schwierig. Ein Lösungsansatz, den wir hier vorstellen werden, ist die realistische Simulation eines überschaubaren Hirnareals, welche es ermöglicht, alle typischerweise verwendeten Verfahren [66] sowie auch neue Verfahren zur Datenanalyse daran zu testen. [59]. Die Neuronensimulation wurde mit GENESIS [9] realisiert, einer Pioniersoftware im Bereich der *Computational Neuroscience*. Bei GENESIS wird eine Nervenzelle in verschiedene diskrete funktionale Abschnitte, Kompartimente, eingeteilt, die sich jeweils wie der zugehörige Zellabschnitt verhalten. Die Simulation eines solchen Zellmodells erfolgt dann Kompartiment für Kompartiment.

Die zentrale Analysemethode, die wir zur Datenauswertung verwenden werden, ist die sogenannte Independent Component Analysis (ICA) [17]. Diese Methode wiederum stammt aus dem klassischen Repertoire der *Neuroinformatik*, handelt es sich doch um die mathematische Formulierung einer Lösung des sogenannten *Cocktail-Party-Problems*: Wie kann man inmitten von lautem Stimmengewirr vieler Menschen dennoch einer einzelnen Person zuzuhören. Ein derartiges Problem wird nicht nur durch das Gehör sondern auch z.B. durch den Geruchssinn tagtäglich gelöst. Bei der ICA kann man also durchaus von einem System sprechen, welches biologie-inspiriert ist. Zwar ist die Ursprungsformulierung mathematisch gewesen, dennoch gibt es unterdessen auch einen sog. neural-ICA-Ansatz, der zumindest demonstriert, dass auch biologische Systeme mit Hilfe von rekurrenten neuronalen Netzen eine ICA-ähnliche Entmischung von überlagerten Eingaben erreichen könnten [6].

Die hier vorgestellte Multielektrodenanalyse mittels ICA wurde in [59] publiziert.

### ***Teil II: Wahrnehmungsprozesse aus Sicht des Maschinellen Lernens***

Der zweite Teil beschäftigt sich mit menschlicher Wahrnehmung. Wahrnehmung beschreibt im weitesten Sinne die Fähigkeit eines Lebewesens, mittels seiner Sinne Information aus seiner Umwelt aufzunehmen und zu verarbeiten. In vielen Fällen können wir nun auf das Ergebnis eines solchen Prozesses zugreifen und müssen uns auf die Beschreibung der Probanden bzw. der Reaktion eines Tieres auf bestimmte Reize berufen, um etwas über die Struktur und Anordnung der Reize im menschlichen Wahrnehmungsraum zu lernen. Wir verwenden in diesen Kapiteln ebenfalls Methoden der *Neuroinformatik*, um die Daten zu analysieren, das Grundmotiv dieser Arbeiten liegt allerdings im Bereich der allgemeinen Neurowissenschaften, versuchen wir doch grundlegende Wahrnehmungsprozesse besser zu verstehen.

Beschreiben Probanden ihre Wahrnehmungen zu einem bestimmten Reiz, sei es eine Farbe, ein Geruch oder ein Gesicht, so ist eine rein verbale Beschreibung meist ungenau und nicht eindeutig zu reproduzieren. Wesentlich präziser kann ein Proband jedoch entscheiden, ob ein vorgegebener Beschreiber mit dem Reiz korrespondiert oder nicht. Dies können Worte sein, es können aber auch Reproduktionen des Reizes sein, wie z.B. Phantombilder, die die Polizei einsetzt, um die Beschreibung von Gesichtern zu präzisieren. Versucht man nun diese Beschreibungen zu quantifizieren, bedarf es einer geeigneten Metrik, um die diskreten Profile miteinander vergleichen zu können. Außerdem verwenden wir in beiden Kapiteln dieses Teiles der Arbeit das sogenannte Multidimensional Scaling (MDS), mit dem wir beliebige Ähnlichkeiten in den kleinstmöglichen euklidischen Raum einbetten. Dies ist unbedingt notwendig, will man zugängliche Wahrnehmungskarten entwerfen, die menschenlesbar sind.

Wir werden in Kapitel 3 einen grundsätzlichen Analyserahmen vorstellen und auf das Beispiel des *Farbensehens* anwenden. In einem von uns durchgeführten Experiment haben wir Probanden verschiedene Farbstimuli mit Farbbeschreibern vergleichen lassen. Ausschließlich basierend auf diesen Daten beschreiben wir detailliert, wie und was wir über das Farbensehen ohne jedes weitere Vorwissen lernen könnten und vergleichen dies mit dem, was wir im Falle des Farbensehens bereits an systemischem Wissen haben.

Die hier vorgestellten Ergebnisse zu dem Analyserahmen wurden mit Bezug auf die Geruchswahrnehmung bereits in [58] publiziert. Die Publikation zu der exemplarischen Anwendung auf das Farbensehen befindet sich noch in der Vorbereitung.

Anknüpfend daran führen wir in Kapitel 4 dieselbe Analyse auch für die Beschreibungen von Gerüchen durch. In einer Kooperation mit dem *California Institute of Technology, USA* wurde uns eine umfangreiche Datenbank mit solchen Beschreibungen zur Verfügung gestellt. Die Ergebnisse dieses Projektes sind ungleich span-

---

nender, da wir zwar eine Grundvorstellung davon haben, wie sich Farbbeschreiber wie *gelblich*, *grünlich* und *bläulich* zueinander verhalten, bei Geruchsqualitäten wie *apfelig*, *kirschig* oder *bananig* sind wir uns bei weitem nicht so sicher über eine passende Anordnung im Riechwahrnehmungsraum. Wir setzen uns intensiv mit der Frage auseinander, wie viele Faktoren man aus den gegebenen Daten ableiten kann, die potentiell den Geruchswahrnehmungsraum aufspannen. Entgegen vergleichbarer Arbeiten arbeiten wir Argumente heraus, dass der Geruchsraum hoch-dimensional ist. Die Topologie dieses hoch-dimensionalen Raumes projizieren wir mittels einer selbst-organisierenden Karte in nur zwei Dimensionen.

Diese Karte untersuchen wir – soweit dies für einen nicht vollständig verstandenen Sinn möglich ist – auf ihre Plausibilität und wir versuchen, auch neue Erkenntnisse bezüglich der Quantifizierung von Geruchswahrnehmungen aus der Karte abzulesen. Es wird deutlich, dass die Kartendarstellung zwar neue Türen aufstößt, wir aber dennoch noch mehr über die hierarchischen Anordnungen der Bezeichner untereinander verstehen müssten.

Die Erzeugung von Geruchswahrnehmungskarten wurde in [57] publiziert.

### ***Teil III: Der kombinatorische Kode des Riechens***

Um die in Kapitel 4 aufgezeigten Probleme zu überwinden und potentielle Geruchsbezeichner besser kategorisieren zu können, erscheint es unverzichtbar, auch experimentelle Daten mit in die Analyse einfließen zu lassen. In einer engen Kooperation mit der *University of California in Irvine (USA)* wurden uns solche Daten in Form von sog. Uptake-Bildern (Bilder neuronaler Aktivität, [84]) des Rattenriechkolben zur Verfügung gestellt.

Diese Bilder zeigen die Aktivitäten des gesamten Bulbus unter der Stimulation mit einer Einzelsubstanz. Wir stellen in Kapitel 5 einen Ansatz vor, die Aktivitäten von Einzelsubstanzen in ihre funktionalen Bestandteile zu zerlegen und diese substanzunabhängig zu begreifen. Wird eine Substanz also als *fruchtig* beschrieben, mit einer leichten Apfelnote, so versuchen wir dem zugehörigen Aktivitätsmuster der Substanz zu entnehmen, welcher Anteil der Substanz die Assoziation *fruchtig* auslöst und welche für den *apfelig*-Anteil verantwortlich ist.

Da wir hier nach wie vor das Problem eines biologischen Systems mit nur sehr wenig Referenzdaten haben, beginnen wir unsere Betrachtungen mit den chemischen Eigenschaften der Einzelsubstanzen und versuchen deren Anteil am Substanzmuster abzuleiten. Die Annahme, dass die Riechrezeptoren entlang biochemischer Bindungseigenschaften der Moleküle organisiert sind und sich diesbezüglich durch einen kombinatorischen Kode auszeichnen, gilt als nachgewiesen [13, 69, 75], die genaue Struktur und Organisation dieses Kodes ist die Wissenschaft bisher jedoch schuldig geblieben.

Wir berechnen hier mit Hilfe von linearen Klassifizierern sogenannte Decision Images, also Prototypen, die ein gegebenes Uptake-Bild als zu einer Substanzklasse zugehörig erkennen können. Es wird an diesen Decision Images mittels einiger Fallbeispiele erläutert, wie damit etwas über den Geruchskode zu lernen ist.

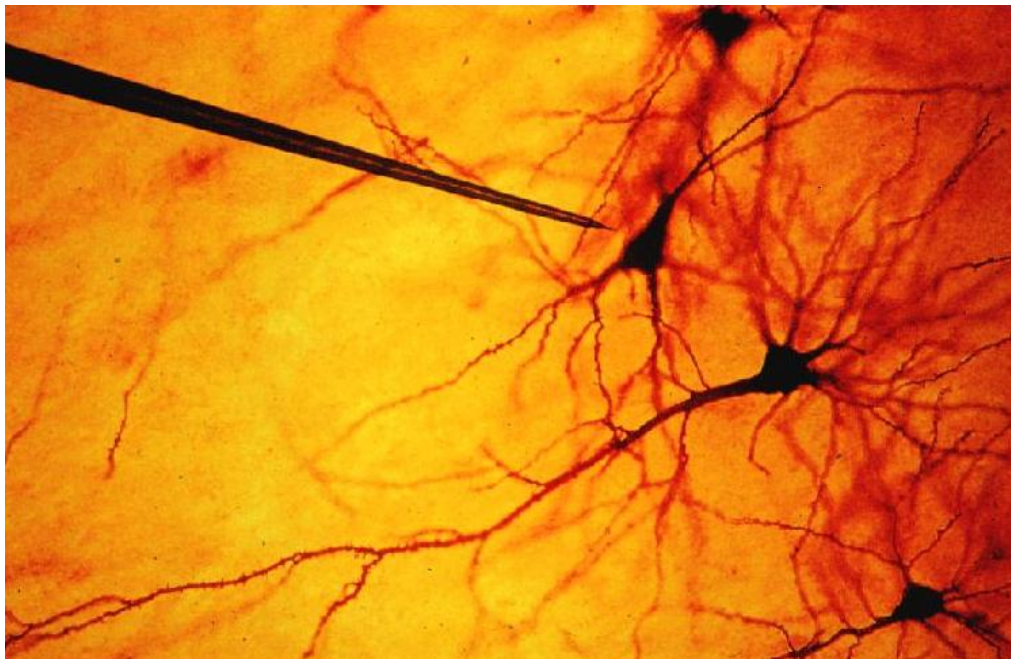
Ob sich dieser Kode nun auch formalisieren und somit entschlüsseln lässt, damit setzt sich das letzte Kapitel (Kapitel 6) dieser Arbeit auseinander. Da die Aktivierungsbilder der Einzelsubstanzen sich als eine Mischung von unterschiedlichen Eigenschaften und Qualitäten auffassen lassen und diese Einzelbilder sich dann auch noch in jeweils anderen Unterräumen befinden, ist es schwer, hier ein geschlossenes Kodierungsmodell abzuleiten.

Wir stellen hier vor, wie wir in Referenz zu einem synthetischen Beispiel versucht haben, zumindest bestimmte mögliche Kodierungsmodelle auszuschließen. Weiterhin haben wir versucht, den Zusammenhang zwischen chemischen Eigenschaften und der zugehörigen Geruchsqualität näher miteinander zu vergleichen und sind auf überraschende Beobachtungen gestoßen: Beide Klassen scheinen sich bezüglich ihrer signifikanten Region grundsätzlich voneinander zu unterscheiden.

Diese Beobachtung lässt sich motivieren durch die funktionale Trennung, die in der Riechrinde vermutet wird. Hier werden, ebenso wie wir die DIs trainieren, auf Basis der Riechkolbenaktivitäten sowohl feinste Unterschiede zwischen zwei Molekülen wie auch chemisch unterschiedliche Duftstoffe mit identischer Geruchsqualität assoziiert. Unabhängig von der endgültigen Erklärung für diese hier beobachteten Effekte eröffnet die Faktorisierung der Riechkolbenaktivitäten mittels der DIs eine ganz neue Interpretation der Aktivitäten im Riechkolben.

Ein zu diesem Themenkreis gehöriger Konferenzbeitrag [60] fand in einem PNAS-Beitrag von Leon und Johnson [52] zwar bereits Erwähnung, die zugehörigen Publikationen befinden sich aber noch in der Vorbereitung.

# Teil I: Biologische Systeme und Maschinelles Lernen



*Glaube denen, die die Wahrheit suchen,  
und zweifle an denen, die sie gefunden haben.*

André Gide

## 2 Multielektroden-Navigation mittels ICA

### 2.1 Einleitung

Das menschliche Gehirn besteht aus einer enormen Anzahl von kleinsten Steuereinheiten, den sogenannten Neuronen. Diese 100 Milliarden Zellen sind eigentlich reine Reizleiter, d.h. sie werden von einem spezifischen Reiz stimuliert. Übersteigt dieser Reiz eine gewisse Schwelle, so leitet das Neuron einen Impuls an alle nachgeschalteten Nervenzellen weiter. Dies sind in der Regel mehrere zehntausend Neurone.

Das Gehirn löst mittels dieser einfachen Bauteile und einer massiv parallelen, komplexen Verschaltung der Neurone erstaunliche Aufgaben und Probleme. Nicht nur Wahrnehmungsaufgaben, wie z.B. das Erkennen von Gesichtern, sondern auch komplizierteste motorische Aufgaben, wie z.B. Klavierspielen, kann das Gehirn bewältigen.

Es gibt aber leider auch einige Krankheiten, bei denen das so beeindruckende Zusammenspiel der Neurone immer weiter gestört wird, wie z.B. bei der sog. parkinsonschen Krankheit. Bei der von James Parkinson erforschten Krankheit, die er selbst „Schüttellähmung“ nannte [74], leiden viele der Betroffenen an einem sogenannten Tremor. Eine solche unkontrollierte Muskelbewegung mit einer Frequenz von 5 bis 7 Hz kann die Verrichtung von alltäglichen Dingen erheblich erschweren.

Die Ursache für die parkinsonsche Krankheit findet sich in einer nur wenige Millimeter großen Region, in der sich – ähnlich einem Kurzschluss – Neuronen unkontrolliert gegenseitig aktivieren und so zu den unbeherrschbaren motorischen Bewegungen führen. Der Tremor ist in vielen Fällen operativ zu beseitigen, indem diese Region mittels Tiefenstimulation durch eine eingeführte Elektrode behandelt wird [50]. Eine ungenau positionierte Elektrode würde nicht nur den Behandlungserfolg gefährden, sondern könnte durch Stimulation gesunder Gehirnregionen gravierende Nebenwirkungen zur Folge haben. Vor dem Eingriff aufgenommene Schnittbilder des Gehirns sind nur bedingt brauchbar, da sich durch das Öffnen der Schädeldecke eine Verformung und Verlagerung des Gehirns ergibt, die zu großen Ungenauigkeiten auf dem vorhandenen Bildmaterial führt.

Im Rahmen des BMBF-Projektes *navEgate* [67] wurde versucht, durch elektrophysiologische Signableitung mit Mikroelektrodensonden sozusagen „blind“ durch das Gehirn zu navigieren und auf der Basis der erfassten Zellsignale Rückschlüsse über das erreichte Gehirnareal zu gewinnen.

Ob und in welcher Güte aus solchen Elektrodendaten Informationen aus dem gerade durchfahrenen Zellgewebe mittels sogenannter *Blind Source Separation*-Methoden gewonnen werden können, war die Aufgabenstellung, die in diesem Zusammenhang von uns bearbeitet wurde [59]. In diesem Kapitel sollen sowohl die verwendeten Methoden als auch die Ergebnisse auf den simulierten Gehirnsignalen vorgestellt werden.

### 2.1.1 Mehrkanal-Mikrosondenaufnahmen von Zellaktivität

Mit Hilfe von Mikroelektroden kann man auch in lebenden Organismen die Aktivität von einzelnen Neuronen messen bzw. teilweise sogar manipulieren. Sticht man mit der Elektrode jedoch in die Zelle (die sogenannte intrazelluläre Messung), wird das Zellsignal unter Umständen stark beeinflusst, mitunter sogar die Zelle komplett zerstört. Daher wird die Messelektrode häufig extrazellulär, also außerhalb der Zelle, positioniert, um dort die negative Depolarisationswelle eines Aktionspotentials (eines sogenannten Spikes) zu messen.

Da die Neurone typischerweise sehr dicht gestaffelt sind (im menschlichen Gehirn z.B. etwa 150.000 Neurone pro  $\text{cm}^3$ ), wird bei einer solchen Messung der Zellaktivität viel Rauschen bzw. ebenfalls die Aktivität der umliegenden Zellen aufgezeichnet. Das Signal- zu Rauschverhältnis kann dabei sehr schlecht sein, so dass die Gefahr groß ist, nicht nur Spikes von einer Zelle sondern auch die der Nachbarzellen zu messen.

Soll eine einzelne Zelle untersucht werden, so ist es aber unerlässlich, zuverlässig alle Spikes dieser Zelle zu erkennen und möglichst wenig *False Positives* - das sind Spikes von umliegenden Zellen - zu erfassen. Dieses sogenannte *Spike Detection* ist eine zentrale Aufgabe bei der Erfassung von Zellaktivitäten, da das Datenvolumen einer Langzeitmessung schnell auf einige Megabyte anschwillt und daher in der Praxis oft nur Spikes in binärer Form abgespeichert werden. Allerdings ist die Validierung von Spike-Detection-Methoden durchaus eine Herausforderung, da oftmals kein Wissen über die tatsächliche Zellaktivität verfügbar ist. In einem vorangegangenen Projekt der Gruppe wurden basierend auf einer Simulation von biologisch realistischen Neuronen mittels GENESIS [9] verschiedene Spike-Detektoren auf ihre Güte untersucht, mit dem Ergebnis, dass keines der getesteten Verfahren (von einfachen Schwellwertmethoden bis zu diskreten Wavelets) eine zufriedenstellende Güte auf den simulierten Daten erreicht. Um die sog. *False Acceptance Rate* (FAR) nicht über 10% steigen zu lassen, stieg im Gegenzug die *False Rejection Ra-*



---

te (FRR) auf dramatische 60% [65, 66]. Mit anderen Worten kann man zwar dafür sorgen, dass nur 10% der erkannten Spikes fälschlicherweise erkannt werden, allerdings muss man dann umgekehrt damit rechnen, dass potentiell 60% der *echten* Spikes gar nicht erkannt werden. Dies ist ein ernüchterndes Ergebnis, in Anbetracht der Tatsache, dass oftmals die Datenanalyse von Elektrodenableitungen aus der Analyse und Interpretation von aufgezeichneten Folgen von Spikes besteht.

Weiterhin sind die erkannten Spikes oftmals nicht Spikes der einen Zelle, die man eigentlich erforschen will, sondern es sind die Spikes von umliegenden Zellen, die kaum unterscheidbar das eigentliche Zellsignal überlagern. Hieraus entstand der Versuch, die gefundenen Spikes anhand der zellspezifischen Spikeform einzelnen Neuronen zuzuordnen, das sogenannte *Spike Sorting*.

Da das vorliegende Zell-Ensemble durch Mehrkanal-Mikrosonden, also durch mehrere Elektroden aus unterschiedlichen Perspektiven aufgezeichnet wurde, bot es sich an, unter Verwendung einer sogenannten „Independent Component Analysis“ (ICA) [30] diese beiden Schritte umzukehren und erst auf den Mehrkanaldaten das Sorting vorzunehmen um hinterher auf den entwirrten Einzelneuron-Aktivitäten die Spike-Detektion durchzuführen.

### 2.1.2 Eine Neural-Cocktail-Party - ICA für Spike Sorting

Die Paradeanwendung für ICA ist die Lösung des sogenannten *Cocktail-Party-Problems*. In diesem Szenario nehmen mehrere Mikrofone von unterschiedlichen Positionen aus die Geräusche auf einer Cocktail-Party auf: Sprache von unterschiedlichen Leuten, ein Glas, das gerade zu Boden fällt und eine Band, die im Hintergrund spielt. Wie kann nun das einzelne gesprochene Wort oder auch einfach nur die Musik aus dieser Partymischung herausgerechnet werden?

Man legt zugrunde, dass die  $n$  gesuchten Signale  $\mathbf{s} \in \mathbb{R}^n$  über eine lineare Mischmatrix  $A \in \mathbb{R}^{n \times n}$  wie folgt zu  $n$  aufgenommen Mixsignalen  $\mathbf{x} \in \mathbb{R}^n$  vermischt wurden:

$$\mathbf{x} = A\mathbf{s}$$

Dann ist eine „Entmischungsmatrix“  $E \in \mathbb{R}^{n \times n}$  gesucht mit

$$E = A^{-1},$$

da dann gilt

$$\mathbf{s} = E\mathbf{x}.$$

Mit anderen Worten, man könnte mit Hilfe von  $E$  aus dem gegebenen Rauschsignal  $\mathbf{x}$  das gesuchte Eingangssignal  $\mathbf{s}$  wieder herausrechnen.

Die Grundannahme für den hier verwendeten Lösungsansatz ist nun, dass die  $n$  Ausgangssignale  $s$  statistisch unabhängig und einer nicht-gaußförmigen Verteilung unterliegen. Da nach dem zentralen Grenzwertsatz der Statistik die Mischung von Signalen gegen eine Gaußverteilung strebt, kann durch das Minimieren dieser Gaußheit eine Entmischung von  $\mathbf{x}$  erreicht werden [30].

Independent Component Analysis (ICA) ist eine Methode, die unter den genannten Voraussetzungen genau dies zu leisten vermag. Man kann die ICA als eine Erweiterung der Hauptkomponentenanalyse, der sogenannten Principal Component Analysis (PCA) betrachten [44]. Bei beiden Verfahren werden die Ausgangsdaten  $\mathbf{x}$  typischerweise erst einmal vom ersten statistischen Moment, dem Mittelwert  $\bar{x}$ , befreit:

$$\tilde{x} = \mathbf{x} - \bar{x}, \quad \text{mit} \quad \bar{x} = \frac{1}{n} \sum_{i=1}^n x(i)$$

Im nächsten Schritt wird dann die Kovarianzmatrix  $C \in \mathbb{R}^{n \times n}$  verwendet, um die Dekorrelationsmatrix  $U \in \mathbb{R}^{n \times n}$  und die Varianzmatrix  $S \in \mathbb{R}^{n \times n}$  abzuleiten:

$$C = \tilde{x} \cdot \tilde{x}^T, \quad UCU^T = S$$

Dabei enthält  $U$  die sogenannten Eigenvektoren  $e_i \in \mathbb{R}^n$  und  $S$  die zugehörigen Eigenwerte  $\lambda_i \in \mathbb{R}$  mit

$$U = \begin{pmatrix} \mathbf{e}_1 \\ \vdots \\ \mathbf{e}_n \end{pmatrix} \quad S = \begin{pmatrix} \lambda_1 & & \mathbf{0} \\ & \ddots & \\ \mathbf{0} & & \lambda_n \end{pmatrix}$$

Mit Hilfe von  $U$  kann jetzt das eigentliche Ergebnis der PCA berechnet werden, ein dekorreliertes Signal  $\tilde{x}_U$  mit

$$\tilde{x}_U = U\tilde{x}.$$

Die Datenpunkte  $x$  werden mit Hilfe von  $U$  so im Datenraum rotiert, dass die neuen Hauptachsen entlang der Eigenvektoren von  $U$  liegen. Die Datenvarianz in Richtung Eigenvektor  $\mathbf{e}_i$  entspricht dem Eigenwert  $\lambda_i$ . Häufig wird nun  $U$  und  $S$  von der höchsten Varianz beginnend sortiert, so dass in der Tat die ersten Komponenten von  $\tilde{x}_U$  den „Hauptkomponenten“ im Sinne der größten Varianz entsprechen.

Hier erkennt man auch sofort, weshalb die PCA auch zur Datenkompression geeignet ist: Besitzt die Datenmenge eine intrinsische Dimension  $n_i$ , die unter der extrinsischen Datendimension  $n$  liegt, so werden  $n - n_i$  Eigenwerte null sein. Oftmals werden zusätzlich aber auch Dimensionen mit relativ niedrigen Dimensionen weggelassen. Diese Problematik wird in Kapitel 4 noch ein wenig tiefer beleuchtet.

An dieser Stelle werden wir die Varianz - das zweite statistische Moment - nicht

weiter betrachten, im Gegenteil, wir werden die entlang der Eigenvektoren dekorrelierten Daten nun „weiß“ machen, d.h. die Varianz wird mit Hilfe von  $S$  auf 1 normiert:

$$W = S^{-1} = \begin{pmatrix} 1/\lambda_1 & & \mathbf{0} \\ & \ddots & \\ \mathbf{0} & & 1/\lambda_n \end{pmatrix}$$

D.h. die Daten  $\tilde{x}_{WU}$  mit

$$\tilde{x}_{WU} = W\tilde{x}_U = WU\tilde{x}$$

sind nun mittelwertfrei, dekorreliert und varianzbefreit. Handelt es sich bei den vermischten Signalen  $\mathbf{s}$  um gaußverteilte Daten, so haben wir nun nur noch eine homogene, runde Datenwolke.

Ist allerdings nur höchstens eines der Signale  $\mathbf{s}$  gaußförmig, so kann nun eine ICA durchgeführt werden. Es wird hier, wie eben schon kurz angedeutet, ausgenutzt, dass durch eine Mischung immer eine gewisse „Gaußheit“ entsteht. Wir suchen deshalb nach einer Rotation  $V$ , die die Achsen entlang der größten „Nicht-Gaußheit“ ausrichtet, in der Hoffnung, so die gesuchten Ausgangssignale wieder aus dem Mischsignal herausrechnen zu können.

$V$  ist dann eine Schätzung der „Entmischungsmatrix“  $E$ . Dies zwar nur auf den weißen Daten  $\tilde{x}_{WU}$ , aber wir können damit trotzdem  $\mathbf{x}$  entmischen, da ja gilt:

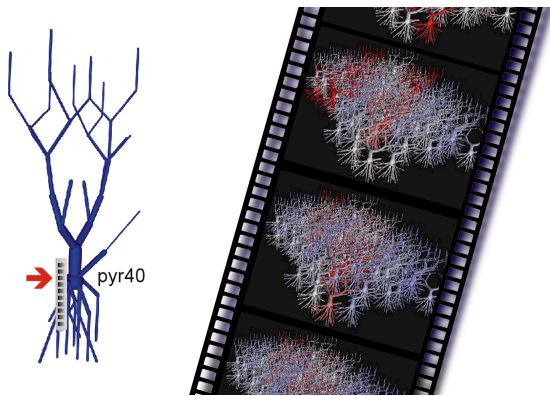
$$\tilde{\mathbf{s}} = V\tilde{x}_{WU} = VW\tilde{x}_U = VWU\tilde{x} = VWU(\mathbf{x} - \bar{x})$$

Ein klassisches Maß für die Nicht-Gaußheit, d.h. um  $V$  möglichst optimal für diese Aufgabe zu wählen, ist die Kurtosis (das vierte statistische Moment) oder auch andere Maße wie die Negentropy [30, 31]. Im weiteren Verlauf der Arbeit haben wir das fastICA-Paket für Matlab<sup>1</sup> von Hyvärinen et al. [31] verwendet, eine sehr robuste und schnell konvergierende Implementation der ICA.

Ebenso wie auf der klassischen Cocktail-Party das Gesagte mehrerer Personen aufgrund ihrer unterschiedlichen Positionen zu den aufnehmenden Mikrofonen aus der Mischung aller aufgenommenen Stimmen herausgerechnet werden kann, kann man auch bei Mehrkanal-Mikrosondenaufnahmen versuchen, die Aktivitäten einzelner Nervenzellen aufgrund ihrer unterschiedlichen Positionen zu den einzelnen Elektroden zu trennen.

Brown et al. [11] haben bereits ICA für Spike Sorting verwendet. Mittels eines 448-Kanal Fotodetektors haben sie Nervenzellaktivitäten bei der Meeresschnecke *Tritonia diomedea* gemessen. Die Güte des Sortings wurde mittels künstlich erzeugter Aktivitätsbursts abgeschätzt, was natürlich nur eine grobe Schätzung der

<sup>1</sup><http://www.cis.hut.fi/projects/ica/fastica/>



**Abbildung 2.1:** Pyramidalzellen im Kompartimentmodell nach Traub. Pyramidalzellen werden in eine bestimmte Anzahl von Abschnitten zerlegt und dadurch abschnittsweise simuliert. Alle Zellen liegen in einem dichten Netz beisammen, in dem induzierte Eingangssignale von Zelle zu Zelle weitergeleitet werden.

tatsächlichen Performance sein kann.

Dies stellt ein fundamentales Problem solcher Experimente dar, denn Methoden wie die ICA können nur grob basierend auf dem erwarteten Verhalten des Systems oder basierend auf der subjektiven Schätzung von Experten evaluiert werden [99].

Dieser Problematik wird hier mit der biologisch realistischen Simulation mittels GENESIS begegnet. Die ursprünglich nur für die Simulation von Einzelneuronen ausgelegte Software wurde erweitert, um auch einen ganzen Gewebeabschnitt mit vielen Zellen und ihre Interaktionen untereinander zu untersuchen – mit dem Vorteil, nun die völlige Kontrolle über die intrazellulären Zustände der Zellen zu haben.

### 2.1.3 GENESIS Simulation des Hippocampus CA-3

Die mit GENESIS simulierten Nervenzellen wurden in eine Konfiguration gebracht, die in etwa dem Gehirnnareal CA-3 des Hippocampus entspricht. GENESIS Neuronen bestehen aus einem sogenannten Kompartiment-Modell, bei dem diskrete Abschnitte der Zelle durch physiologisch realistisches Verhalten simuliert werden. Die Idee eines solchen Modells ist es, mit steigender Anzahl an Kompartments eine immer realistischere Simulation der Einzelzelle zu bekommen.

In Abbildung 2.1 ist ein solches Kompartimentmodell in einer räumlichen Darstellung zu sehen. Ebenfalls zu erkennen ist die relative Position der simulierten Elektrode und der Ausgangsposition der simulierten Zellen. Vor der zufälligen Dislokalisierung liegen die Zellsomata zwischen der 5. und 6. Elektrode. Das Modell kann man am Besten als eine Diskretisierung des kontinuierlichen Potentialraumes einer natürlichen Zelle verstehen. Je größer die Anzahl der Kompartments wird, desto realistischer simuliert das Modell eine tatsächliche Zelle.

---

Für die Simulation des Hippocampus wurden Pyramidenzellen aus 66 Kompartimenten zusammengesetzt, die Interneurone aus 48. Die genauen Parametrisierungen der einzelnen Kompartimente wurden von den Zellmodellen von Traub et al. übernommen [90, 91], das extrazelluläre Potential an den simulierten Elektroden wurde nach Nunez gemäß folgender Gleichung modelliert:

$$F = \frac{1}{4\pi s} \sum_{i=1}^n \frac{I_i}{r_i}$$

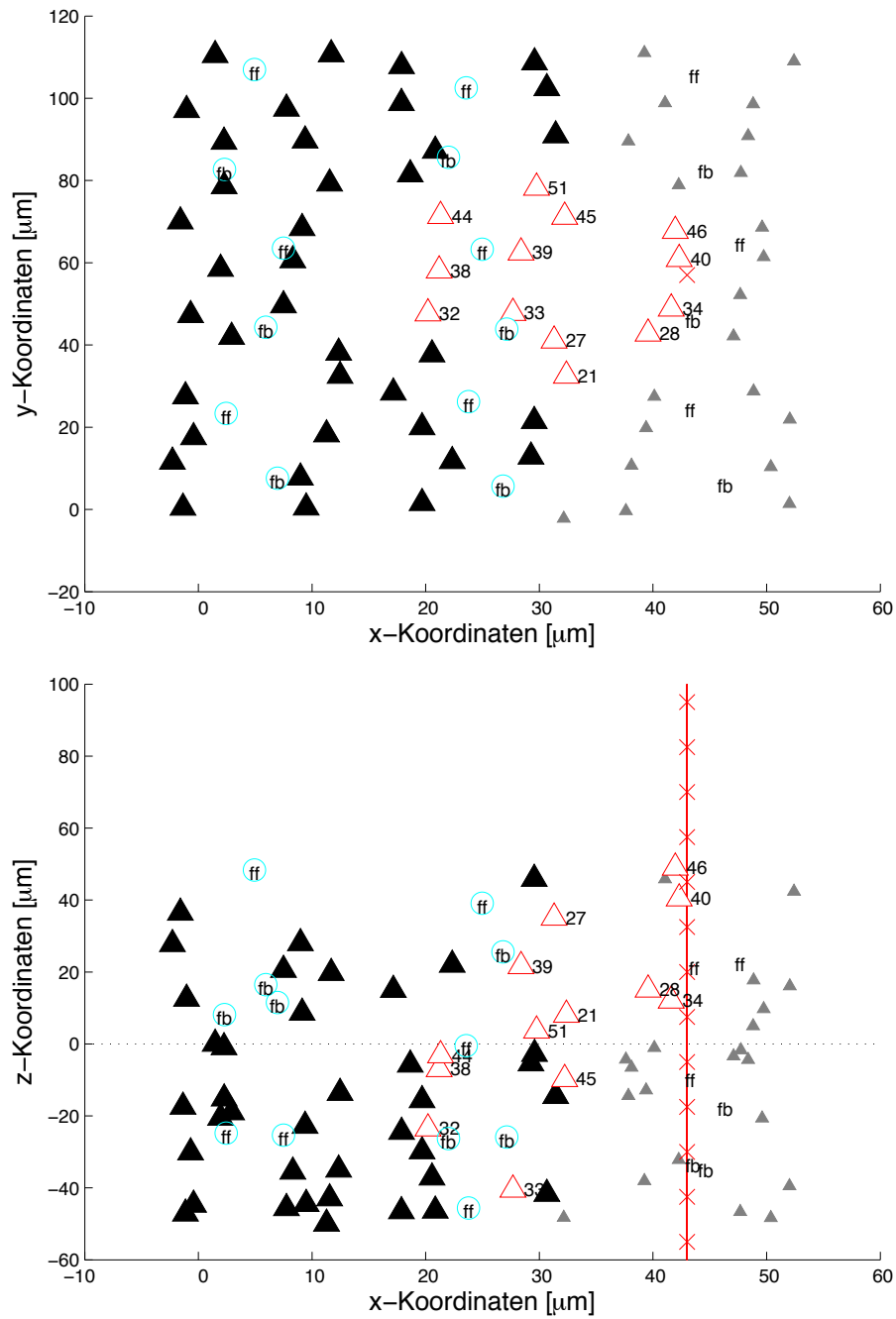
Das Potential ergibt sich also aus der Summe der Transmembranströme  $I_i$  der  $n$  Kompartimente in Bezug auf ihren jeweiligen Abstand  $r_i$  vom simulierten Ableitungspunkt. Die Konstante  $s$  beschreibt die elektrische Leitfähigkeit im Zellumfeld. Dies ist ein stark vereinfachtes, aber durchaus übliches Modell, um das extrazelluläre Umfeld abzubilden.

Im folgenden betrachten wir eine Simulation von 72 Pyramiden-Zellen (identische Klone) in einer zweidimensionalen Anordnung ( $6 \times 12$  Zellen). Dazwischen wurden 18 Interneurone (9 feedforward- und 9 feedback-Interneurone) verteilt, die ebenfalls zu dem abgeleiteten Signal beigetragen haben. Die Zellen wurden durch Zufallswerte in ihrer  $z$ -Position variiert ( $\pm 50\mu m$ ) und um die  $z$ -Achse rotiert, um die morphologische Variabilität noch weiter zu erhöhen. Es gibt unterdessen auch Modelle, bei denen der Dendritenbaum mit realistischen Varianzen wächst, was dann zu einer noch realistischeren Simulation führt [82].

In Abbildung 2.2 ist dieses Netzwerk abgebildet. Die Pyramidenzellen sind durch Dreiecke, die Interneurone durch Kreise symbolisiert. Der virtuelle Messpunkt ist durch das Kreuz an den Koordinaten  $[40, 60]$  in Abbildung 2.2 zu erkennen. Es handelt sich hier um eine Multielektrode mit 13 Mikroelektroden in linearer Anordnung. Die Elektroden liegen übereinander mit einem Abstand von  $12.5\mu m$  und sind mit einem Öffnungswinkel versehen, der ihre Lage auf einer Seite einer isolierten Nadel abbilden soll.

Zellaktivität, die  $15\mu m$  von einer Elektrode gemessen wird, hat bereits  $2/3$  seiner Intensität verloren, d.h. ab dieser Distanz ist nicht zu erwarten, Zellen noch zuverlässig messen zu können. Zusammen mit dem Öffnungswinkel kann man bereits mutmaßen, welche Zellen in den simulierten Aufnahmen überhaupt nur gefunden werden können. In Abbildung 2.2 sind die Zellen hinter dem Öffnungswinkel der Elektroden durch kleine Pyramiden gekennzeichnet, die „sichtbaren“ durch größere. Rot und unausgefüllt sind die 13 Zellen, die sich zusätzlich noch im Radius von  $15\mu m$  vor der Elektrode befinden [65].

In Abbildung 2.3 sind die simulierten intrazellulären Aktivitäten dieser 13 Zellen zu sehen, ebenso das Messsignal an den 13 simulierten Elektroden. Einige der Spikes sind noch zu erkennen, viele werden allerdings offensichtlich durch die Depolarisationen anderer Zellen ausgelöscht. Es ist leicht zu erkennen, dass es grundsätzlich



**Abbildung 2.2:** Die Ausgangssituation der GENESIS Simulation. Pyramidenzellen sind durch Dreiecke gekennzeichnet, Interneurone durch *ff* (feedforward) und *fb* (feedback). Die Pyramidenzellen im Sichtbereich der Elektrode sind groß dargestellt, Zellen innerhalb des  $15\mu m$ -Radius zusätzlich hohl. Die Elektrode befindet sich senkrecht zu dem Kreuz.

---

keine einfache Aufgabe ist, aus den Elektrodensignalen die Spikes der einzelnen Zellen (dem Ausgangssignal) herauszurechnen.

Wir haben in diesem Szenario genauso viele Elektroden wie Zellen, die über dem Rauschlevel Signale beisteuern. Für eine „Entmischung“ braucht man mindestens so viele Meßsignale, wie es Signalquellen gibt. Sowohl die Anzahl der Zellen im Sichtbarkeitsbereich wie auch die Anzahl der Mikroelektroden auf der Sonde bewegen sich in realistischen Größenordnungen, d.h., es ist sehr wahrscheinlich, auch in einem echten Setup durchaus Sonden zu finden, die in der Größenordnung einer solchen (1 : 1) -Abdeckung Signale liefern.

Unter der Annahme, dass die aufgenommenen Zellaktivitäten eine lineare Mischung von statistisch unabhängig feuernden Neuronen darstellen, können wir auf diese Daten die in 2.1.2 bereits vorgestellte ICA anwenden. Streng genommen sind vermutlich beide Annahmen nicht ganz korrekt, da sich weder das Verhalten der extrazellulären Spannungen linear exakt modellieren lässt noch die Zellen in diesem Verbund wirklich statistisch unabhängig feuern.

## 2.2 ICA auf Sondendaten des CA-3

In Abbildung 2.3 ist noch einmal das Ausgangsszenario dargestellt. Im oberen Teil der Abbildung sind die intrazellulären Potentiale der mit GENESIS simulierten Pyramidalzellen zu sehen. Man kann deutlich erkennen, wie sich innerhalb der Zelle die Spannung aufbaut, welche sich beim Überschreiten der Reizschwelle in einem Aktionspotential entlädt. Diese sog. Spikes haben eine signifikant höhere Amplitude und wären durch eine einfache Detektionsschwelle leicht von den Restsignalen zu segmentieren. Wie wir bereits erwähnt haben, ist es aus technischen Gründen jedoch nicht möglich, intrazelluläre Ableitungen aus einem ganzen Zellverbund zu messen.

Im unteren Teil der Abbildung 2.3 ist das Elektrodensignal dargestellt, wie es in der Simulation bezüglich der Zellanordnung aus Abb. 2.2 an den (links dargestellten) 13 Elektroden gemessen wird. Es ist deutlich zu sehen, dass zwar einige wenige Spikes noch zu erkennen sind, die meisten der Signale in der extrazellulären Aufnahme durch die Addition der zusammenlaufenden Potentiale eliminiert bzw. sehr stark verfälscht werden.

### 2.2.1 PCs und ICs der simulierten Multielektrodensignale

Abbildung 2.3 (unten) zeigt genau die Signale, welche in einem realen Experiment als Eingabe in die nachfolgende Datenanalyse gegeben werden. Oftmals wird zwar durch zeitgenaue Reize eine Reaktion des System induziert, um Marken in den Signalen zu haben, die man für eine weitere Analyse verwenden kann, diese sind jedoch unpräzise und setzt voraus, dass das beobachtete System auch genau wie erwartet auf den Reiz reagiert.

Wir wollen im folgenden versuchen, nur auf der Basis der gegebenen Elektrodenableitungen die intrazellulären Zellaktivitäten so gut wie möglich vorherzusagen. Dies geht über die reine Vorhersage von Spikes im System hinaus, da wir diese Vorhersagen *zellspezifisch* machen wollen.

Bevor wir uns mit dieser Zuordnung auseinandersetzen, wollen wir uns die Ergebnisse der beiden Verfahren – PCA und ICA – einmal vergleichend ansehen. Wie bereits ausführlich erläutert, erhalten wir mit Hilfe der PCA ein dekorreliertes Signal, welches bei der ICA zusätzlich noch zu einer statistisch unabhängigen Lösung transformiert wird. Betrachten wir nun Abbildung 2.4 (oben), so müssen wir feststellen, dass die PCA das Problem augenscheinlich nicht wirklich lösen kann. Einzig die Hauptkomponente 4 (PC04) weist eine gewisse Ähnlichkeit zu einem Elektrodensignal auf, viele andere lassen zwar so etwas wie Spikes erkennen, oftmals zeigen diese Spikes allerdings in die falsche Richtung oder das Signal hat Spikes auf beiden Seiten des Signals.

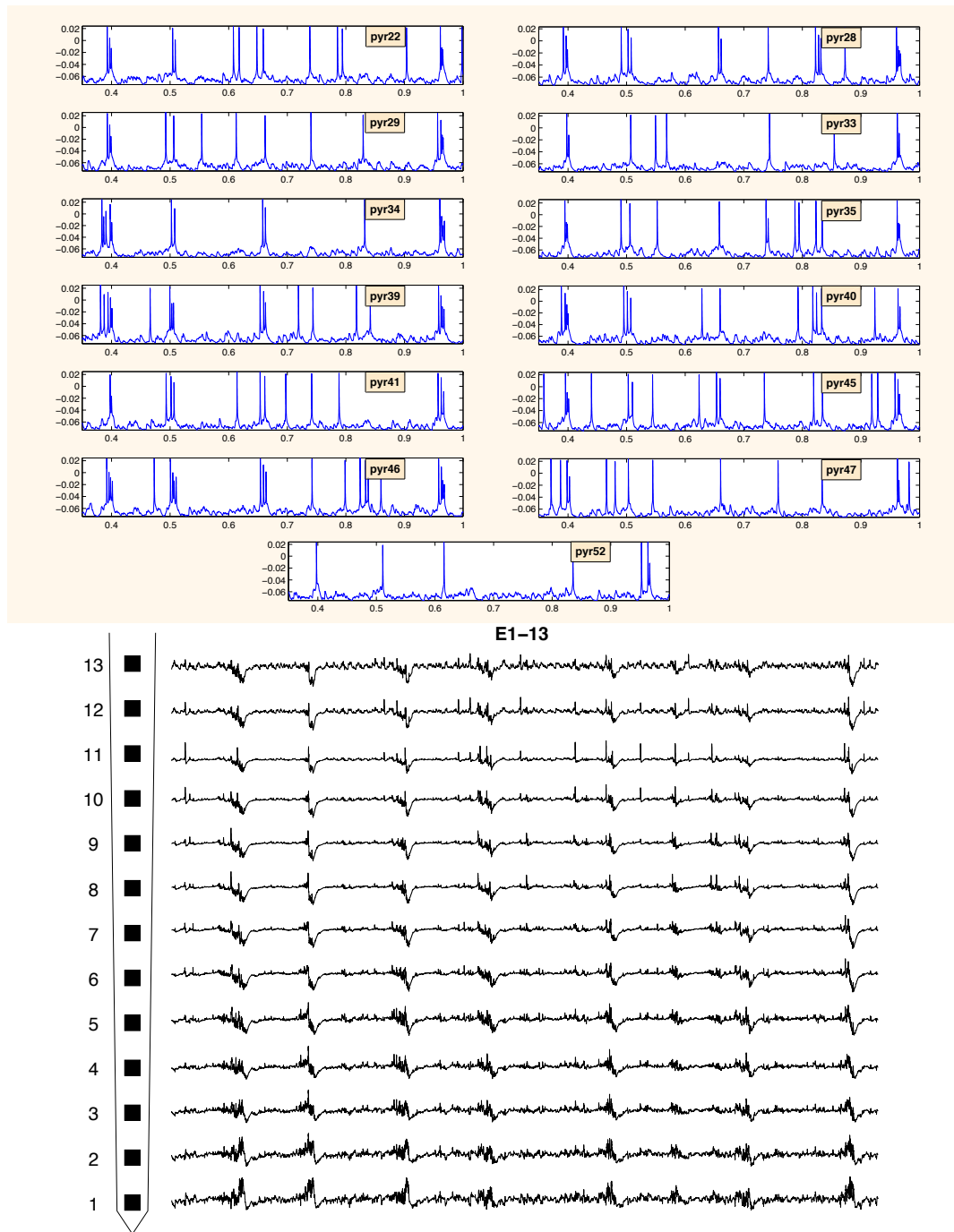
Bei den Ergebnissen der ICA (Abb. 2.4 unten) sieht es bereits etwas besser aus, visuelle Inspektion lässt zumindest bereits bei den ICs 1, 2, 3, 5 und 7 ein Spikesignal vermuten. Dann jedoch gibt es wieder einige Komponenten wie zum Beispiel die IC 4, bei der es sich eher um amplifiziertes Rauschen zu handeln scheint.

### 2.2.2 Signalzuordnung und Zelllokalisation

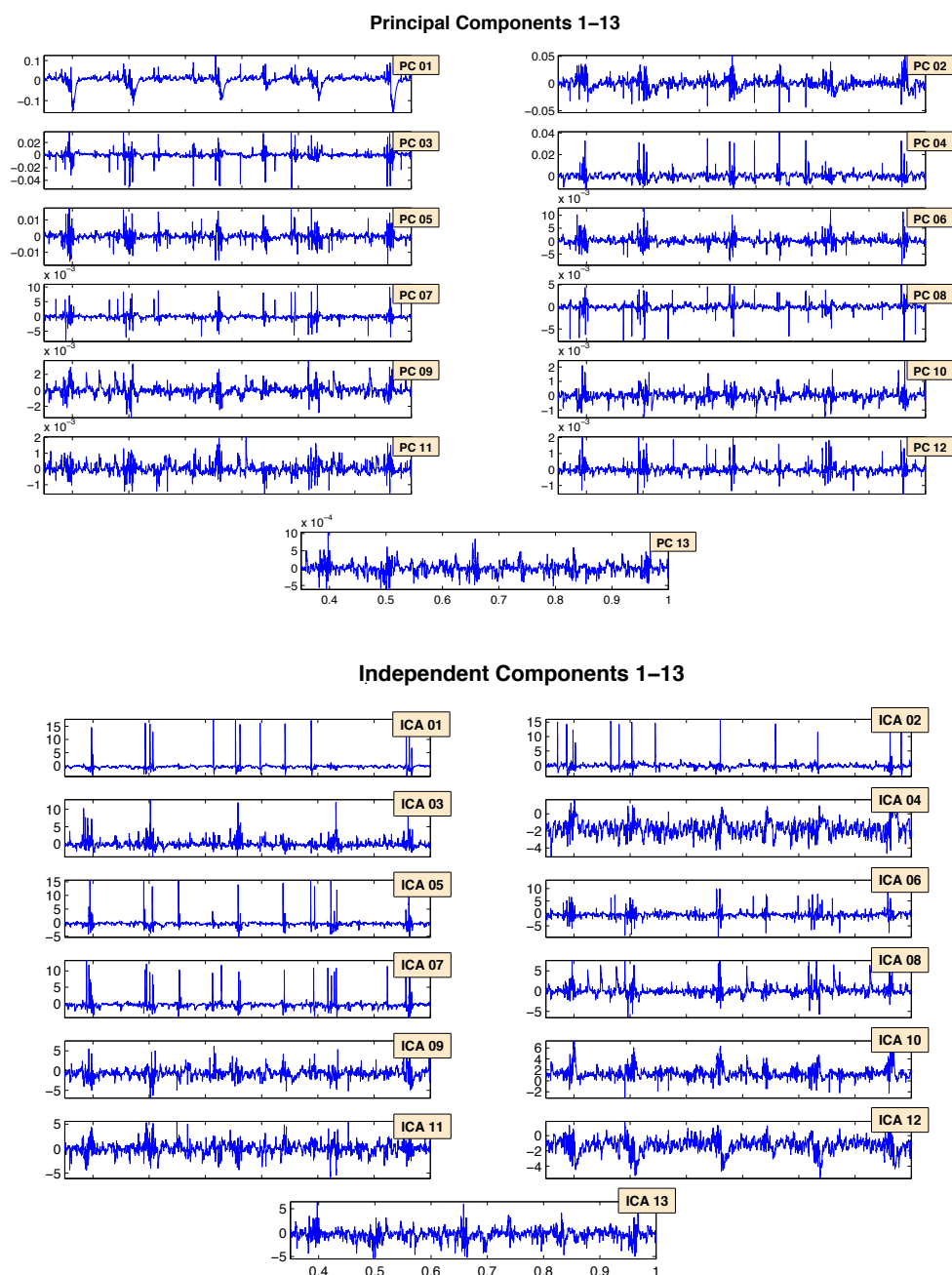
An dieser Stelle sei noch einmal ein entscheidender Unterschied in der Auswertung von PCA und ICA genannt: Bei der PCA können wir anhand der Varianz der einzelnen Komponenten, also der Größe ihrer zugehörigen Eigenwerte, auch direkt die Relevanz der einzelnen Komponenten einschätzen. In unserer Abbildung zum Beispiel sind die PCs bereits nach absteigendem Eigenwert sortiert, d.h. da die ersten 5 Hauptkomponenten bereits mehr als 99% der Gesamtvarianz der Daten ausmachen, würde man die letzten PCs als Rauschen abtun.

Bei der ICA hingegen sind die Komponenten zueinander ungewichtet. Durch das Whitening ist sämtliche Varianz aus dem System entfernt worden und wir können





**Abbildung 2.3:** In der oberen Abbildung ist die intrazelluläre Aktivität der 13 sichtbaren Pyramidenzellen für eine Dauer von 1.5s aufgetragen. Die Kennungen entsprechen denen aus Abb. 2.2. In der unteren Abbildung dargestellt sind die aufgezeichneten Signale der im Zellverband simulierten 13 Mikroelektroden.



**Abbildung 2.4:** Oben: Abgebildet sind die 13 resultierenden Komponenten der PCA. Auch wenn einige der Spikes etwas verstärkt sind, ist für die meisten Komponenten kein klares Signal zu erkennen. Unten: Abgebildet sind die 13 Komponenten nach der ICA. Für einige der Komponenten ist das Signal-zu-Rausch-Verhältnis wesentlich besser geworden.

---

keine Aussage über die Relevanz einer Komponente treffen. Dies ist grundsätzlich natürlich kein Nachteil, aber dennoch müssen wir einen Weg finden, um auch auf Echtwelt-Szenarien eine Entscheidung treffen zu können, welche Komponenten diejenigen sind, in der wir gesuchte Information finden und welche nicht.

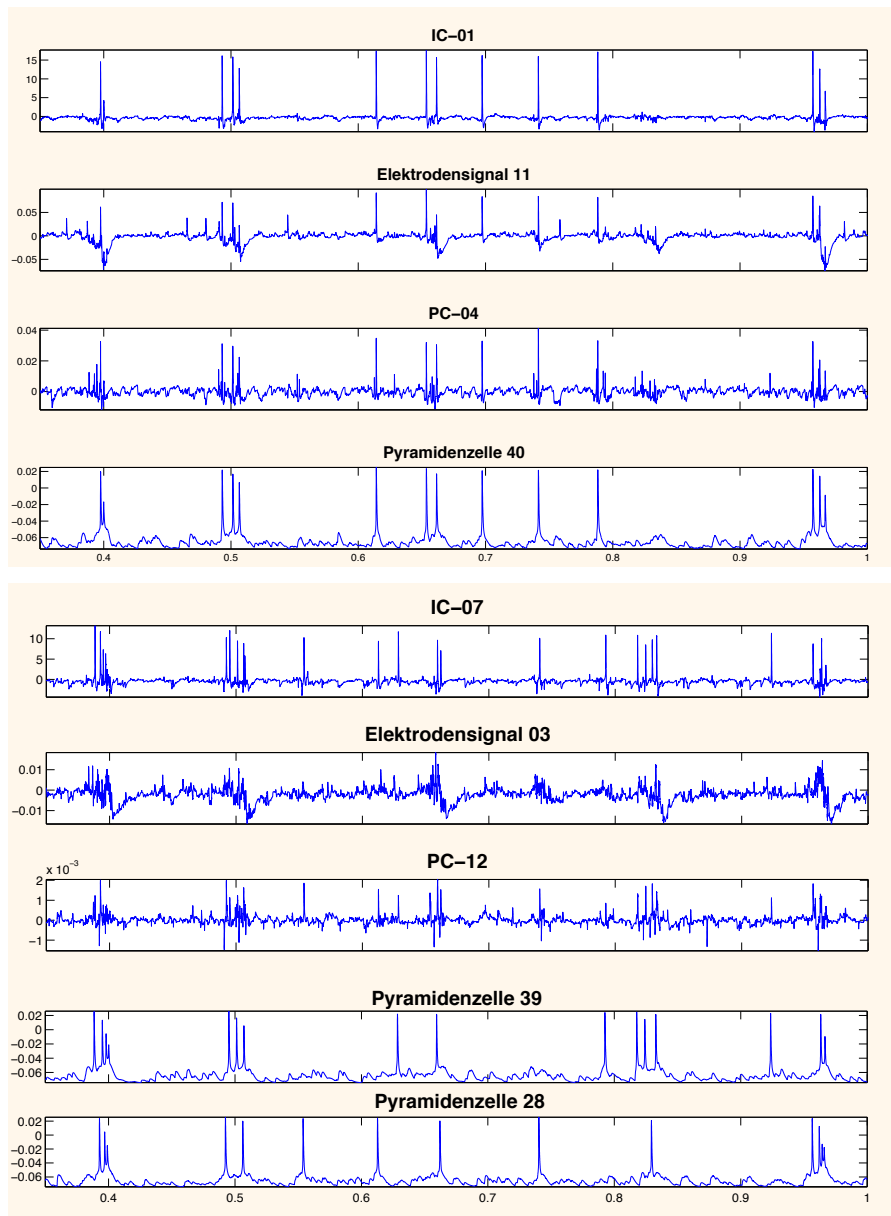
Konsequenterweise müssen wir diese Information vermutlich aus dem Elektrodensignal auslesen. Ist eine Zelle gut „sichtbar“ vor einer der Elektroden positioniert, so sollte die Wahrscheinlichkeit recht hoch sein, dass diese Zelle auch in dem Signal sichtbare Spikes hinterlässt. Somit haben wir uns entschieden, erst einmal die ICs, die mit den sichtbaren Spikes im Elektrodensignal am stärksten korreliert sind, als relevant zu betrachten. Mit anderen Worten, wir versuchen die Spikes von Zellen, die besonders gut von einer der Elektroden aufgenommen wurde, als solche in den ICs wiederzufinden.

Um eine solche Zuordnung zu erhalten, haben wir das Skalarprodukt zwischen den Elektrodensignalen und den ICs berechnet und so den 13 ICs ein abschätzendes Gütemaß zugeordnet. Je höher die Korrelation zwischen der IC und dem Elektrodensignal ist, desto wahrscheinlicher handelt es sich dabei um ein Zellsignal.

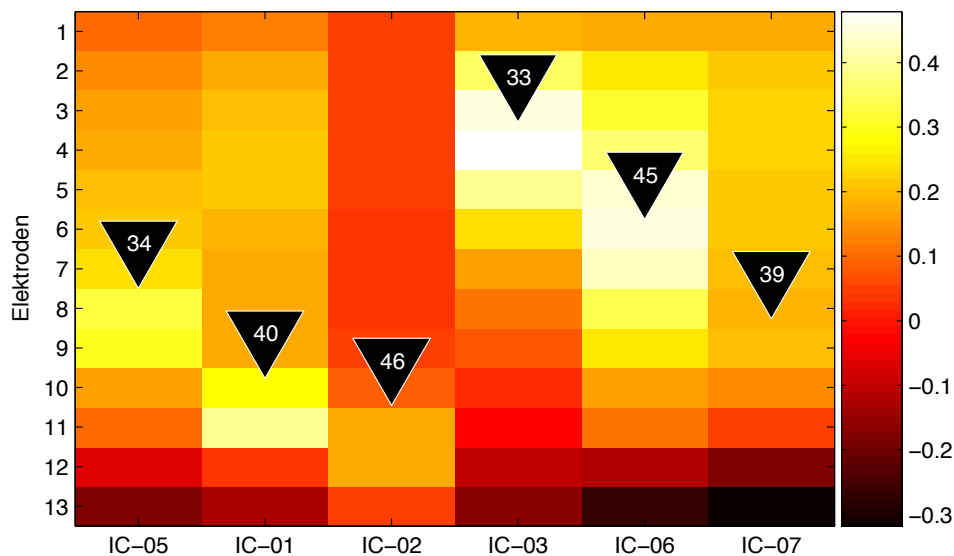
Tatsächlich erhalten wir 6 IC-Komponenten, die eine signifikant höhere Korrelation mit den Elektrodendaten haben und somit vermutlich Signale von gut sichtbaren Zellen rekonstruieren. Im folgenden werden wir uns auf die Analyse dieser 6 Komponenten konzentrieren, im Rahmen des hier vorgestellten Projektes haben wir kein weiteres Kriterium gefunden, welches sich ausschließlich nur auf in Laborszenarien verfügbare Informationen beschränkt und ebenfalls etwas über die Güte der IC-Signale aussagt.

In Abbildung 2.5 sind zwei Komponenten der ICA mit den Elektrodendaten, mit denen sie am stärksten korreliert sind, abgebildet. Ebenfalls abgebildet sind die zugeordneten Hauptkomponenten und die Pyramidenzellen, deren Spikes sie vermutlich rekonstruiert haben. Die Komponente *IC-01* ist am stärksten korreliert mit der Pyramidenzelle 40 (zu sehen in der untersten Zeile in der Abbildung), welche sich direkt vor der Elektrode befindet und entsprechend prominent auf den Signalen vertreten ist (siehe auch Abb. 2.2). Bei der ICA-Komponente *IC-07* handelt es sich hingegen um das schwächste Signal, welches wir mit hinzugenommen haben. *IC-07* ist am stärksten mit Pyramidenzelle 39 korreliert, welche zwar nicht verdeckt aber doch in einem Abstand in Höhe der 7. Elektrode liegt (siehe Abb. 2.2).

Das am besten zur *IC-01* passende Elektrodensignal *e-11* (2.Zeile) zeigt ebenfalls viele Spikes der Zelle 40, es ist aber auch zu erkennen, wie einige der Spikes bereits in Wellentälern verschwinden bzw. andere Zellen mit weiteren Spikes das Signal überlagern. Dies ist beim Elektrodensignal *e-03* (6.Zeile) noch wesentlich stärker ausgeprägt. Solche Täler entstehen häufig, wenn umliegende Zellen kurz vorher gefeuert haben und durch ihre Depolarisation den Spike der beobachteten Zelle neutralisieren. Entsprechend viele Fremdspikes sind auch im Elektrodensignal vor



**Abbildung 2.5:** Der oberste Plot zeigt die ICA-Komponente mit der besten Entsprechung auf den Elektrodensignalen. In den Zeilen darunter sieht man das mit Skalarprodukt ermittelte ähnlichste Elektrodensignal (2.Zeile), die ähnlichste Hauptkomponente (3.Zeile) und die Pyramidenzelle mit dem ähnlichsten Signal (4.Zeile). In den Zeilen 5–8 ist dasselbe für die ICA-Komponente mit der zweitbesten Entsprechung zu sehen. Hier gibt es zwei Pyramidenzellen (28+39), die der IC-07 ähnlich sind.



**Abbildung 2.6:** Die Helligkeiten zeigen das normalisierte Skalarprodukt zwischen den Elektrodendaten und den ICs. Dargestellt sind spaltenweise die 6 ICs mit der stärksten Korrelation zu den Elektroden und ihr Skalarprodukt zu allen 13 Elektroden. Weiterhin sieht man die Position des Zellsoma für die der jeweiligen IC ähnlichste Pyramidalzelle.

einem solchen Spannungstal zu erkennen.

Die zugeordneten Hauptkomponenten (3.+7.Zeile) schaffen hier auf dem dekorrelierten Signal bereits eine wesentliche Verbesserung. Die Spikes werden klarer segmentiert, die von anderen Zellen werden abgeschwächt. Dies klappt bei der PC für die Zelle 40 allerdings wesentlich besser als für Zelle 39. Bei der *PC-12* (7.Zeile) sind einige Spikes komplett eliminiert, einige *False Positives*, also zellfremde Spikes, erscheinen in dem Signal.

Betrachtet man vergleichend hierzu nun die ICs, so fällt auf, dass die Spikes deutlich besser rekonstruiert werden. Vor allem gehen kaum Spikes verloren. Die Rekonstruktion der Spikes von Pyramidenzelle 40 durch *IC-01* ist nahezu komplett, bei der Rekonstruktion von Zelle 39 durch *IC-07* scheinen im wesentlichen auch alle Spikes herausgerechnet worden zu sein, allerdings gibt es einige *False Positives*. In dieser IC scheinen zwei Zellen vermengt zu sein. Um diesem Verdacht nachzugehen ist auch das intrazelluläre Signal der Pyramidenzelle 28 in Abbildung 2.5 (Zeile 9) gezeigt. Diese Zelle ist am zweitbesten mit *IC-07* korreliert und tatsächlich sind auch deren Spikes der IC zugeordnet.

Betrachtet man in Abbildung 2.2 die Positionen innerhalb der Simulation, befinden sich die beiden Zellen zwar an unterschiedlichen Positionen, liegen jedoch in nahezu derselben Höhe zu der Elektrode. Bedingt durch die senkrechte Anordnung der Mikroelektroden scheinen die beiden Signale nicht zu trennen zu sein, zumindest für die ICA erscheinen sie als eine einzige Quelle.

In Abbildung 2.6 ist das Skalarprodukt der 13 Elektroden jeweils mit den sechs ICs dargestellt, die die stärkste Korrelation zwischen Elektrodensignalen und IC haben. Jede Zeile in dieser Grauwertmatrix entspricht einer Elektrode, jede Spalte einer der ICs. Je heller der Grauton desto besser passt die IC zu dem dort aufgezeichneten Signal. Es fällt auf, dass die besten Treffer auch tatsächlich in Abhängigkeit stehen mit der z-Position der Elektrode im simulierten Umfeld.

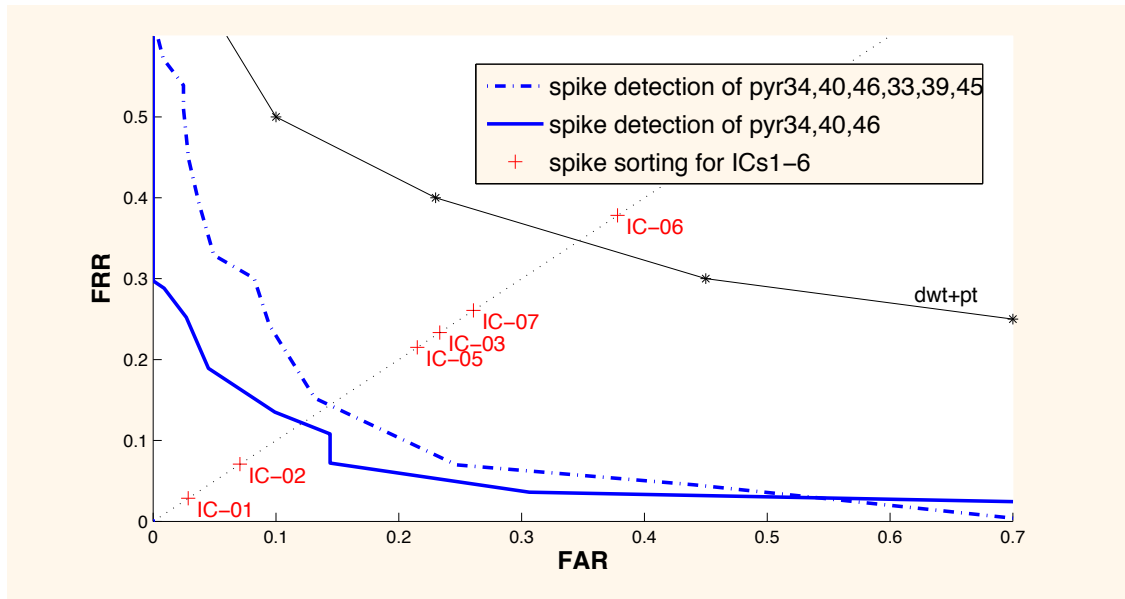
Diesen Zusammenhang haben wir bereits eben bei der Überdeckung von zwei Zellen in einer IC-Komponente angedeutet. Dies ist nicht erstaunlich, da ja genau die relative Position zu den Elektroden der Faktor ist, der die Einträge in der initialen Mischmatrix  $S$  maßgeblich bestimmt. Angemerkt sei hier auch noch die Tatsache, dass sich ein extrazelluläres Signal kurz unterhalb des Somas der Zelle am besten ableiten lässt. Um dies ebenfalls zu überprüfen und zur vollständigen Analyse von Abb. 2.6 haben wir auch die am stärksten zu den ICs korrelierten Zellen berechnet und deren Soma-Position mit in die Matrix eingetragen.

Wir haben für das Skalarprodukt zwischen den ICs und den intrazellulären Potentialen natürlich Wissen verwendet, das in der realen Laborumgebung nicht vorhanden ist. Dennoch ist es interessant zu sehen, wie gut sich die Zellen hier in ihrer z-Koordinate lokalisieren ließen. Dies hat einen positiven Nebeneffekt: Haben wir bisher festgestellt, dass es eine IC stören kann, wenn zwei Zellen relativ zu den Elektroden nicht unterscheidbar zum Mischsignal beitragen, so könnten wir im umgekehrten Fall – also bei idealer Anordnung der Elektroden – vermutlich nicht nur in noch besserer Qualität die Signale trennen, sondern wir könnten auch eine Aussage darüber treffen, *wo* sich diese Zellen relativ zu den Elektroden befinden.

Hier haben wir durch die Ausrichtung entlang der z-Achse eine sehr einfache Elektrodenanordnung verwendet, sollte sich diese Anordnung in zukünftigen Entwürfen allerdings auch entlang der anderen Achsen erstrecken, so ist das Ziel einer Lokalisierung durchaus zu erreichen. Es gibt in der ursprünglichen Formulierung des Cocktail-Party-Problems bereits Ansätze, die gefundenen Quellen auch zuverlässig im Raum zu lokalisieren [5, 8]. Auch Tiere und Menschen sind mittels zweier „Mikrofone“ in der Lage, sich auf einen Sprecher zu konzentrieren bzw. eine isolierte Quelle zu lokalisieren.

### 2.2.3 Spike Sorting und Spike Detection mittels ICA

Das eigentlich hier zu lösende Problem ist das Detektieren von Spikes und das Zuordnen selbiger zu den ursächlichen Zellen. In Abbildung 2.5 wurde schon deutlich, dass wir mittels Skalarprodukt zumindest für die gut auf den Elektroden sichtbaren Zellen eine brauchbare Kandidatenauswahl erhalten, die wir gegen bisherige Ansätze zur Zuordnung der Spikes zu den Zellen (dem sog. Spike Sorting) vergleichen können.



**Abbildung 2.7:** Die schwarze Linie mit Sternen zeigt die ROC-Kurve für die Spike Detection mit dem diskreten Wavelet-Ansatz von Menne et al [65]. Die anderen beiden Kurven zeigen die ROC-Kurven für die Spike Detection bzgl. der besten sechs ICs (strichpunktiert) und die besten drei ICs (durchgezogen). Es wurden jeweils die zugeordneten Pyramidenzellen als Ground Truth verwendet (siehe Legende). Die Kreuze markieren die EER für das Spike Sorting für jede IC bzgl. ihrer zugeordneten Zelle (siehe auch Abb. 2.6).

Wie bereits erwähnt, ist der herkömmliche Weg das Detektieren der Spikes und dann – nachgeschaltet – das Zuordnen zu einzelnen Zellen. Die ICA vollzieht diese Zuordnung bereits im ersten Schritt und erlaubt es, hinterher die Spikes zu zählen. Die Detektion ist einfach, schließlich haben wir Signale, die vergleichbar sind mit intrazellulären Potentialen, d.h. eine einfache Schwellwertoperation reicht bereits aus, um die Spikes vom Rest des Signals zu trennen.

Mit diesem Schwellwert kann man den Klassifikator regulieren und zwei Messgrößen ableiten, die üblicherweise verwendet werden, um einen Klassifikator zu evaluieren. In den Extrempositionen werden entweder alle Signale, egal ob Spike oder nicht, unterhalb der Schwelle liegen, im anderen Extrem liegen immer alle Signale oberhalb und werden als Spikes „erkannt“. Man hofft, dass es irgendwo dazwischen eine Schwelle gibt, die die beiden Klassen perfekt trennt, so dass es keine einzigen *false positives* (Nicht-Spikes über der Schwelle) und keine *false negatives* (Spikes unter der Schwelle) mehr gibt. Darüber definieren sich die Meßgrößen der *False Accep-*

tance Rate (FAR) und der False Rejection Rate (FRR) als

$$\mathbf{FAR} = \frac{\#\text{false positives}}{\#\text{all negatives}} \quad (2.1)$$

$$\mathbf{FRR} = \frac{\#\text{false negatives}}{\#\text{all positives}} \quad (2.2)$$

Bewegt man nun  $R$  sukzessive in dem Intervall  $[0; 1]$ , so erhält man eine charakteristische Kurve der Klassifikationsleistung des Klassifikators, die sogenannte ROC-Kurve (receiver-operating-curve oder auch specificity-sensitivity-curve). Ein optimaler Klassifikator erreicht eine gemeinsame Fehlerrate ( $\mathbf{FAR} = \mathbf{FRR}$ ), die sog. Equal-Error-Rate (EER) von 0%. Dies entspricht der gesuchten optimalen Trennung der beiden Klassen.

Um die Qualität des Sortings zu evaluieren, haben wir für jede der sechs ICs die EER berechnet. Dies ist, wie eben erwähnt, der Punkt einer ROC-Kurve, bei der FAR und FRR gleich sind. Je weiter dieser Punkt an den Ursprung rückt, desto besser ist die Gesamtleistung des Klassifikators. In Abbildung 2.7 sind die EER-Punkte für die ICs eingetragen. Ganz offensichtlich zeigen zwei ICs ein hervorragendes Ergebnis, wohingegen die anderen ICs nicht perfekt funktionieren. Die Rate ist jedoch noch immer besser als die allgemeine Detektionsrate eines Spike Detektors, der eigentlich das vermeintlich leichtere Problem löst, die bloße Detektion von Spikes. Weiterhin muss man berücksichtigen, dass wir für die betrachteten ICs jeweils nur eine einzige Zelle berücksichtigen. Z.B. für *IC-07* haben wir bereits weiter oben motiviert, warum teilweise auch zwei Zellen zu einem Signal zusammenfallen können.

Abschließend haben wir uns auch die Ergebnisse für die Spike Detection angeschaut. Wie erwähnt wird dieses Problem ja sozusagen nebenbei gelöst: Sind die Signale der einzelnen Zellen erst einmal rekonstruiert, ist die Detektion darauf sehr leicht. Zum Vergleich zu anderen Spike-Detektions-Verfahren wurden auch für die Detektion ROC-Kurven für den ICA-Ansatz berechnet.

Wir betrachten im folgenden zwei Zielmengen: Alle intrazellulären Spikes der Pyramidenzellen, die unseren IC-Komponenten zugeordnet sind (34, 40, 46, 33, 39 und 45) sowie die drei Zellen für die ICs, die die sichtbarsten Zellen auf dem Elektrodensignal repräsentieren. Die zugehörigen ROC-Kurven sind in Abbildung 2.7 eingetragen. Wie man es erwarten würde, verläuft die Kurve für die drei besten ICs unterhalb der Kurve für die sechs ICs. Im Vergleich dazu ist eine ROC-Kurve für ein Verfahren gezeigt, welches auf denselben Elektrodendaten versucht, mittels einer diskreten Wavelettransformation spiketypische Verläufe zu finden und so Spikes zu detektieren [65]. Diese Kurve verläuft weit abgeschlagen oberhalb der IC-basierten ROC-Kurven. Deutlich wird hier, wie schwer die Detektion von den Spikes auf den Elektrodendaten ist. Auch wenn die Kurven nicht direkt miteinander vergleichbar



---

sind, ist dennoch zu erwarten, dass auch die Detektionsrate bei einem geeigneten Elektrodenlayout wesentlich verbessert werden kann.

Abschließend sei noch erwähnt, dass wir dasselbe Experiment auch auf einer Simulation von 5 Sekunden Länge durchgeführt haben, um die ICA-Analyse auf Robustheit und Zeitunabhängigkeit zu untersuchen [59]. Hierzu haben wir die ICA auf den gesamten 5 Sekunden angewendet und auf einem zufälligen Fenster mit 1.5 Sekunden Länge. Dann haben wir beide Lösungen auf weiteren Fenstern mit 1.5 Sekunden Länge getestet. Für alle betrachteten Fenster ergaben sich vergleichbare Ergebnisse für beide IC-Lösungen.

## 2.3 Diskussion

In diesem Kapitel sind zwei Methoden zur unüberwachten Zerlegung von Mischsignalen vorgestellt worden, PCA und ICA. PCA dekorreliert ein gegebenes Mischsignal, ICA hingegen versucht die dekorrelierten Signale zusätzlich noch statistisch unabhängig zu machen.

Wir haben an dem konkreten Beispiel der Multielektroden-Navigation motiviert, wie die ICA hier hervorragend verwendet werden kann, um Einzelzellaktivität aus einem Mischsignal herauszurechnen. Dies wurde mit einer GENESIS-Simulation des Hippocampus mit 100 Pyramidenzellen demonstriert. Es wurde gezeigt, wie man mit Hilfe der ICA nur auf der Basis von Multielektroden Daten bereits viele brauchbare Informationen aus dem Zellverbund rekonstruieren kann.

Die ICA löst hierbei initial bereits ein sehr schwieriges Problem, das sogenannte Spike Sorting, bei dem jedes Aktionspotential in den Ableitungen der zugehörigen Zelle zugeordnet werden soll. Die Zuordnung erfolgt dabei so gut, dass auch die Spike Detection auf dem rekonstruierten Signal kein Problem mehr ist.

Wir haben ebenfalls festgestellt, dass die Spikes nicht nur einzelnen Zellen zugeordnet werden können, sondern dass für diese Zellen auch deren relative Position zu den Elektroden ermittelt werden kann. In unserer Simulation haben wir nur eine sehr einfache lineare Anordnung der Elektroden entlang der  $z$ -Achse betrachtet. Durch eine Weiterentwicklung des GENESIS-Paketes, die eine Nutzung von Parallelrechner-Architekturen ermöglicht, sollte es in weiteren Projekten möglich sein, die simulierten Feldaktivitäten mit vielen unterschiedlichen Elektrodenarchitekturen auszumessen und zu evaluieren. Eine Erweiterung auf echte räumliche Lokalisierung wäre dann nicht nur möglich, sondern würde vermutlich sogar die Güte der Signal-Separierung verbessern.

Wir haben hier die ICA in einer sog. *Batch*-Variante betrachtet, d.h. es stehen alle

Daten zur Analysezeit bereits zur Verfügung. Es gibt allerdings auch eine *Pattern-by-Pattern*-Variante, bei der schrittweise die ICs immer besser gelernt werden. Somit wäre es möglich, die ICA in ein echtzeitfähiges Analysesystem zu integrieren, welches nicht nur aufschlüsseln kann, welche Zelle feuert, sondern diese Zellen auch auf eine Karte der Elektrodenumgebung eintragen könnte. Dies könnte zu einer ganz neuartigen und wesentlich effizienteren Elektrodennavigation führen verglichen zu dem, was derzeit möglich ist.

Abschließend bietet sich auch eine genauere Analyse der verbleibenden Signalkanäle an, die wir im Rahmen dieses Projektes nicht weiter untersucht haben. Dies waren die Signale die keine starke Korrelation mit dem Elektrodensignal aufgewiesen haben. Welche Information sie tragen, ob sie evtl. Theta-Bänder des Zellverbundes repräsentieren oder andere Kenngrößen, die für die weitere Analyse des Systems interessant sein könnten, wird in folgenden Projekten sicher auch eine zentrale Rolle spielen.

## Teil II: Wahrnehmungsprozesse aus Sicht des Maschinellen Lernens



*You don't need eyes to see,  
You need vision!*

Faithless - *Reverence*

## 3 Der Farbwahrnehmungsraum

### 3.1 Einleitung

Das fundamentale Verständnis für die menschlichen Wahrnehmungsprozesse ist schon immer ein faszinierendes Forschungsfeld gewesen. Da der Mensch in seiner Entwicklung das visuelle System als seinen primären Wahrnehmungssinn etabliert hat, ist es eine ganz besondere Herausforderung, die auf diesen Sinn spezialisierten Prozesse zu verstehen und durch maschinelle Systeme mit möglichst ähnlicher Leistungsfähigkeit nachzubilden.

Die meisten der Wahrnehmungsprozesse sind zwar in ihrem sensorischen Ablauf, also der Physiologie des Sinnes, recht gut verstanden, wie sich die hierarchische Verarbeitung bis zur finalen Wahrnehmungsebene exakt zusammensetzt und vielleicht sogar nachzubilden ist, ist noch immer ein offenes Kapitel und Forschungsinhalt vieler ambitionierter Projekte. Als Beispiel sei hier die Gesichtserkennung genannt: Menschen sind sehr gut darin, bereits auf kleinsten Bildausschnitten Personen und deren Gesichter zu erkennen und diese dann sogar noch bestimmten Personen zuzuordnen. Bei Fernsehbildern aus Fußballstadien beginnen wir, nach den Personen zu suchen, von denen wir wissen, dass sie sich derzeit im Stadion befinden. Für ein computergestütztes Erkennungssystem wäre das alles andere als ein unterhaltsames Spielchen nebenher.

Ein weiterer, weitgehend unverstandener Wahrnehmungsprozeß ist das Riechen. Auch diese Sinneswahrnehmung ist für viele Säugetiere noch immer ein primärer Sinn und hat schon die ersten Philosophen beschäftigt. Dennoch wissen wir über die genaue Funktionsweise dieses Sinns noch wenig. Vor allem können wir die Duftstoffe selbst nur sehr wenig quantifizieren. Ist eine Substanz unbekannt, ist über ihren potentiellen Geruch kaum etwas wenn nicht sogar gar nichts zu sagen. Für viele Moleküle ist es noch nicht mal möglich, auch nur vorherzusagen, ob sie überhaupt nach irgendetwas riechen.

Um das Riechen rankt sich der Rest dieser Arbeit, in diesem Kapitel werden wir uns jedoch zunächst mit dem Farbsehen beschäftigen. Die Farbwahrnehmung ist ein Prozess, welcher zumindest auf der sensorischen Seite sehr gut verstanden ist. Wir werden im folgenden den Farbwahrnehmungsraum untersuchen und motivieren,

weshalb diese Art der Analyse eventuell auch dem Verständnis von komplexeren und vor allem weniger gut verstandenen Wahrnehmungsprozessen neue Erkenntnisse liefern kann.

Zur Zeit der alten Griechen war über die Physiologie des Auges und den lichtsensitiven Nervenzellen, mit deren Hilfe wir Licht unterschiedlicher Wellenlänge differenzieren können, nur wenig bekannt. Dennoch oder gerade deshalb waren auch sie schon fasziniert von den Augen und dem, was der Mensch mit ihnen von seiner Umwelt wahrnehmen kann.

Sie konnten jedoch nur beobachten und genau analysieren, was die Beschreibungen der Menschen über das Sehen preisgeben würde. Die ältesten überlieferten Studien über das menschliche Sehsystem lassen sich bis in die Zeit von Aristoteles – er lebte im 4. Jahrhundert v. Chr. – zurückdatieren. Seine Interpretation über das Farbensehen basierte auf der Annahme, ein Objekt verändert auf eine gewisse Art und Weise das Medium zwischen sich selbst und dem Auge des Betrachters. Im Mittelalter dann hat man diese Idee verworfen und angenommen, dass das Auge bestimmte Emissionen erzeugt, die dann z.B. durch Reflexion an den Objekten selbige sichtbar machen.

Wir stellen uns jetzt einmal vor, wir wissen ebenfalls nichts über das Farbensehen, es ist uns gänzlich unbekannt, dass Farben etwas mit der Wellenlänge des Lichtes zu tun haben und dass wir drei verschiedenartig sensitive Rezeptoren in der Retina besitzen, die es uns ermöglichen, Farben zu sehen. Unser Kenntnisstand würde also ungefähr dem von Aristoteles entsprechen. Die einzige zuverlässige und reproduzierbare Quelle wäre die Beschreibung von Farben durch Menschen. Genaugenommen sind es die Worte, die sie verwenden, um ihre persönlichen Farbeindrücke zu artikulieren. Immerhin sind diese Beschreibungen als Expertenwissen zu interpretieren, schließlich haben die Personen zu dem System, das uns so interessiert, ja im wahrsten Sinne des Wortes einen direkten Draht. Auf der Basis der durch Sprache beschriebenen Farben können wir uns nun also Gedanken über die mögliche Struktur und Organisation dieses Sinnes machen.

Wir wollen diesen Weg noch ein wenig weiter gehen und versuchen, einen Analyserahmen zu entwerfen, der es uns ganz allgemein ermöglicht, derartige Daten zu verstehen.

Das Farbensehen erscheint hierbei als ein gutes Referenzsystem, da wir an den Farben viele Zusammenhänge intuitiv illustrieren können. In diesem Kontext wird auch ein Experiment zum Farbensehen vorgestellt, welches wir als eine Validierung des in diesem Kapitel vorgestellten Analyserahmens betrachten.

---

### 3.1.1 Sprache und Wahrnehmung

Grundsätzlich sind Worte natürlich vage und die Beschreibungen vielleicht nicht immer von Person zu Person übertragbar. Worte, die zur Beschreibung von Wahrnehmungen verwendet werden, sind oftmals kulturell geprägte Metabeschreibungen, getroffen aus einem bestimmten Kontext heraus.

Aber immerhin gilt für verbale Beschreibungen von Wahrnehmungen, sowohl für Farben wie auch für Gerüche: Meist sind die Beschreibungen konsistent, robust und reproduzierbar [18, 78]. Es bleibt die Frage, was man aus den Beschreibungen der Menschen über die Struktur und die Komplexität des Farbensehens lernen kann.

Ein intuitives Experiment wäre es, vielen Probanden Farben zu zeigen und diese durch Worte beschreiben zu lassen. Sinnvollerweise schlägt man den Probanden vielleicht sogar typische Farbwörter vor, aus denen sie sich die am besten zum Farbeindruck passenden Worte herausuchen können. Letztlich haben wir dann eine große Datenbank voller Informationen über den Zusammenhang zwischen den Farbstimuli und den verbalen Beschreibungen der zugehörigen Sinneseindrücke der Probanden.

Bis zu dieser Stelle werden wir in diesem Kapitel den Weg des neugierigen Philosophen gehen, ihn dann jedoch verlassen und mit den Vorzügen der heutigen Zeit mittels einer computergestützten Analyse die Daten verarbeiten und interpretieren.

### 3.1.2 Merkmalsräume vs. Wahrnehmungsräume

Wir wollen im folgenden den Farbwahrnehmungsraum analysieren und dabei möglichst wenig Vorwissen verwenden. Bei der Interpretation der Ergebnisse werden wir natürlich das Wissen über den Farbraum bemühen, um vergleichend vorhersagen zu können, was einem unbekannten Wahrnehmungssystem an Information zu entlocken ist.

Die Kernfrage, die wir in diesem Kontext beleuchten wollen, ist die Frage danach, wie die Struktur des Wahrnehmungsraumes aussieht. Beschreibungen von Stimuli jedweder Art sehen meist wie folgt aus: Wir erhalten typischerweise eine Beschreibungsmatrix  $P$  mit

$$P = (p)_{ij} \in R^{m \times n}, \quad 0 \leq p_{ij} \leq 1$$

wobei wir  $m$  Stimuli haben, die mittels  $n$  Merkmalen beschrieben werden. Diese Merkmale können Messwerte, Beschreibungen, Bewertungen oder ähnliches sein, grundsätzlich also alles, was an Informationen zum einzelnen Stimulus verfügbar ist.

Mathematisch betrachtet haben wir also  $m$  Punkte – für jeden Stimulus einen – die sich in einem  $n$ -dimensionalen Raum befinden, der durch die  $n$  vorliegenden Merkmale aufgespannt ist. Dies ist der sogenannte *Merkmalsraum*. In diesem Raum kann man analysieren, wie die Stimuli sich bzgl. der gegebenen Merkmale verhalten, z.B. welche Stimuli zu einem Cluster zusammenfallen.

Wir werden im weiteren am Farbsehen illustrieren, weshalb man nur mit gewissen bekannten Metainformationen Merkmale begreifen und über die Struktur des *Wahrnehmungsraumes* etwas lernen kann. Wir werden in unserer weiteren Betrachtung argumentieren, dass in Entsprechung zum Merkmalsraum der Wahrnehmungsraum umgekehrt genau der Raum ist, der von den  $n$  Stimuli aufgespannt wird und in welchem sich die  $m$  Bezeichner als Punkte befinden.

#### 3.1.3 Grundlagen des Farbsehens

Als Einstieg in die Problematik werden wir im folgenden jedoch erst einmal die Grundlagen des Farbsehens vorstellen, damit wir darauf bei der Interpretation unserer Ergebnisse zurückgreifen können. Ganz knapp sollen hier die Grundlagen über das sensorische System vorgestellt werden, mit dem der Mensch in der Lage ist, Licht als etwas wahrzunehmen, was er als Farbe beschreibt. Für eine ausführlichere Darstellung sei auf entsprechende Fachliteratur verwiesen (z.B. [7] und [53]).

#### Farben

Zuerst einmal wollen wir Licht als solches betrachten. Lichtwellen haben eine definierte Wellenlänge, der eine bestimmte Farbe zugeordnet werden kann. Der Mensch ist in der Lage, Licht mit einer Wellenlänge von ca.  $450nm$  bis  $700nm$  wahrzunehmen. Licht niedriger Wellenlänge ( $450nm$  und größer) wandelt mit ansteigender Wellenlänge seinen Farbeindruck aus dem noch nicht sichtbaren ultravioletten Licht zu einer blauen Farbe. Die Farbe verändert sich dann über grün und gelb bis zu rot, bevor sie jenseits der  $700nm$  als infrarotes Licht das Intervall des wahrnehmbaren Lichtes wieder verlässt.

Basierend auf dieser Erkenntnis könnte man jetzt auf die Idee kommen, der Wahrnehmungsraum des Farbsehens sei eindimensional. Dies stimmt aber nicht. Das, was wir letztlich als Farbe wahrnehmen, ist ein komplexerer Eindruck. Dies kann man sich ganz einfach bewusst machen: Jeder Wellenlänge kann zwar eindeutig eine Farbe zugeordnet werden, jedoch gibt es nicht für jede Farbe eine Wellenlänge: Für Farben wie *rosa* oder *weiß* finden wir keine Modulation, in der das Licht in dieser Farbe schimmert.



---

Wie können wir aber Farben sehen, wenn es gar kein Licht mit dieser Farbe gibt? Tatsächlich resultieren unsere Farbeindrücke aus einer *Mischung* von Licht unterschiedlicher Wellenlänge. Das Auge scheint also in der Lage zu sein, Mischlicht in sein Wellenspektrum zu zerlegen und daraus einen kognitiven Farbeindruck zu erzeugen.

## Das menschliche Auge

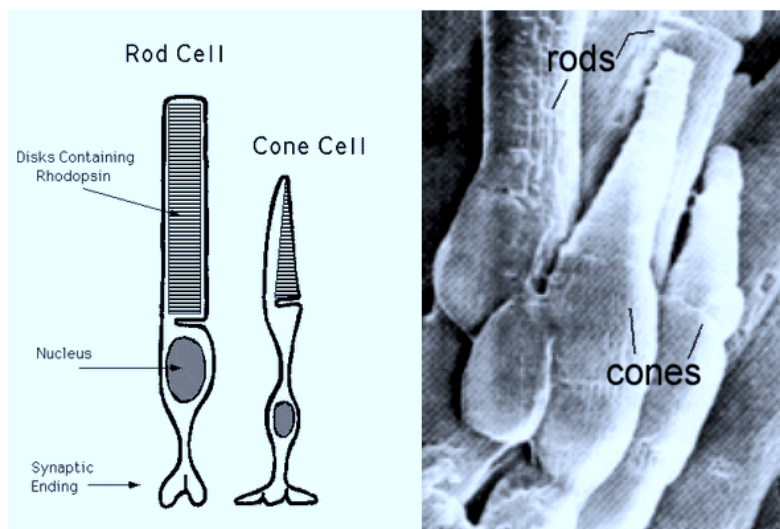
Licht wird im menschlichen Auge auf der Retina absorbiert. Dort befindet sich ein ganzes Mosaik an lichtsensitiven Nervenzellen, die sich in zwei grundsätzliche Gruppen unterteilen lassen:

1. Helligkeitssensitive Nervenzellen (Stäbchen)
2. Wellenlängensensitive Nervenzellen (Zapfen)

Die Namen Stäbchen und Zapfen leiten sich direkt von der Form dieser Zellen ab (siehe auch Abbildung 3.1). Die Stäbchen spielen eine zentrale Rolle beim strukturierten Sehen und beim Sehen im Dunklen. Sie haben keine Möglichkeit, sensitiv auf Licht bestimmter Wellenlänge zu reagieren. Die Zapfen dagegen verfügen über Pigmente, welche sie sensitiv machen für bestimmte Wellenlängen. Es gibt drei verschiedene Typen von Zapfen:

1. S-Zapfen (*short wavelength receptor*):  
Der auch Blaurezeptor genannte S-Zapfen deckt den Blau-Bereich des sichtbaren Farbspektrums ab, bei einem Absorptionsmaximum von etwa  $440nm$  (blauviolett). Er reagiert sensitiv in dem Bereich von  $350nm$  bis  $550nm$ .
2. M-Zapfen (*medium wavelength receptor*):  
Der auch Grünrezeptor genannte M-Zapfen hat ein Absorptionsmaximum von etwa  $540nm$  (smaragdgrün) und reagiert sensitiv in dem Bereich von  $400nm$  bis  $600nm$ .
3. L-Zapfen (*long wavelength receptor*): Der L-Zapfen hat ein Absorptionsmaximum von etwa  $560nm$  (gelbgrün) und reagiert sensitiv in dem Bereich von  $400nm$  bis  $700nm$ . Trotz seiner Farbantwort wird er auch Rotrezeptor genannt, da er die Hauptleistung in der Wahrnehmung des Rotbereichs übernimmt.

Stäbchen und Zapfen sind auf der Retina nicht gleichverteilt. Dort wo die optische Achse der Linse die Retina schneidet, befindet sich die sogenannte Fovea, das Zentrum des Sehfeldes. Dort findet man ausschließlich Zapfen. In dem umliegenden, wesentlich größeren Feld, der sogenannten Macula Lutea, finden sich Stäbchen und Zapfen, noch weiter außerhalb in der Peripherie finden sich nur noch Stäbchen.



**Abbildung 3.1:** Stäbchen und Zapfen sind lichtensitive Nervenzellen, welche sich in der Retina befinden. Mit Zapfen lassen sich bei guten Lichtverhältnissen Farben sehen, die Stäbchen hingegen nehmen lediglich Helligkeiten wahr, unabhängig von der Farbe des Lichtes. **Links:** Schematische Darstellung eines Stäbchen (Rod Cell) und eines Zapfen (Cone Cell). **Rechts:** Mikroskopische Aufnahme von Zapfen und Stäbchen in der Retina (Bild stammt aus [7]).

Das bedeutet, dass wir nur in einem ziemlich kleinen Bereich eine scharfe Farbwahrnehmung haben und ansonsten Farben nur sehr verschwommen oder gar nicht wahrgenommen werden können. Dies fällt aber kaum auf, da das Auge beim normalen Sehen eigentlich in konstanter Bewegung ist und sämtliche Bereiche, die uns interessieren, fokussiert und somit in den Bereich der scharfen Farbwahrnehmung bringt.

Das erfreuliche an den fundierten Kenntnissen über das Tuning der Zapfen in der Retina ist, dass wir aus dem Energiespektrum, also der Wellenlängen-Mischung, für jedes Licht vorhersagen können, wie der zugehörige Rezeptorinput in das visuelle System aussehen wird. Dieses physikalische Kontinuum liefert uns eine sogenannte Ground Truth – es können sich daran andere Verfahren messen.

Basierend auf diesem Kontinuum entsteht ein Farbmodell, wie wir es kennen: Alle Farbtöne können in einem zweidimensionalen Farbkreis angeordnet werden und zusätzlich gibt es noch eine Dimension für die Farbintensität (die sog. Sättigung) [53]. Dies entspricht der Annahme, dass man durch eine Linearkombination aus drei unterschiedlichen Rezeptoren höchstens einen dreidimensionalen Raum aufspannen kann.

Kann man aber nun alleine auf der Basis von verbalen Beschreibungen auf ähnliche Erkenntnisse kommen? Dazu haben wir ein Experiment durchgeführt, in welchem wir Farben von Probanden anhand von vorgegebenen Worten beschreiben lassen.

---

Die Ergebnisse dieses Experimentes und die zugehörige Analyse der Daten – beide im Rahmen einer von mir betreuten Arbeit entstanden [25, 26] – werde ich im verbleibenden Teil dieses Kapitels vorstellen.

## 3.2 Das „Color Perception Project“

Wie zuvor beschrieben, haben Menschen ein grundlegendes Verständnis darüber, wie Farbsehen funktioniert und welches die Faktoren sind, die die Funktion des Systems prägen. Wir wollen im folgenden testen, ob ein sensorisch weitgehend unbekanntes Wahrnehmungssystem sinnvoll zu charakterisieren und zu quantifizieren ist. Im Rahmen dieses Kapitels betrachten wir die Analyse des Farbsehens im Detail, um uns dann in Kapitel 4 an eine Analyse des Geruchswahrnehmungsraumes zu wagen.

### 3.2.1 Farben und Farbbeschreiber

Als Ausgangsbasis haben wir 90 Farbstimuli zufällig im RGB-Raum generiert. Hierzu wurde für jede Farbe jeweils ein Tripel von Zufallswerten im Intervall  $[0, 1] \in \mathbb{R}^3$  erzeugt. In Abbildung 3.2 ist die Verteilung der Farbwerte auf dem Farbkreis (links) und die Verteilung über die Helligkeitswerte (rechts) aufgetragen. Für den linken Plot wurde die Helligkeit der Stimuli auf einen konstanten Wert gesetzt, für den rechten sind Farbton und Sättigung konstant gehalten. Man kann deutlich erkennen, dass die meisten Farbtöne recht gut repräsentiert sind. Aus der zufälligen Generierung der Stimuli folgt direkt, dass wir eher satte und helle Farben in unserem Testset haben. Dies ist am ehesten noch in der Helligkeitsverteilung (rechts) zu erkennen, die Helligkeitswerte verteilen sich jedoch subjektiv recht homogen.

Wesentlich schwieriger war es, eine Gruppe von Beschreibern zu finden, die die gesamte Vielfalt an Farbeindrücken beschreiben können. Außerdem sollten es Worte sein, die Menschen typischerweise verwenden, um ihre Farbeindrücke zu beschreiben. Vermutlich die beste – leider aber auch die aufwändigste – wäre es nun gewesen, eine Reihe von Probanden Farbtafeln systematisch beschreiben zu lassen und aus der Menge der verwendeten Beschreibern langsam eine Grundmenge von typischen Farbbeschreibungen zu entwickeln.

Im Rahmen des hier vorgestellten Projektes war dieser Evaluationsschritt leider nicht umsetzbar. Wir haben stattdessen Listen aufgesetzt, in denen wir systematisch Adjektive gesammelt haben, mit denen Farben beschrieben wurden. Diese Listen haben wir nach und nach immer weiter reduziert um eine handhabbare Menge

an Begriffen mit möglichst wenig redundanten Beschreibern zu erhalten, was sich im späteren Verlauf des Experimentes in den Daten auch bestätigen sollte.

Am Ende erhielten wir eine Farbbeschreibungsliste mit 16 Farbbezeichnern, wovon sechs tatsächlich auch die Farbqualität beschreiben:

- rötlich
- grünlich
- bläulich
- gelblich
- bräunlich
- rosig

Sieben der verbleibenden Beschreiber widmen sich der Farbtemperatur:

- farblos
- dunkel
- grell
- klar
- warm
- kalt
- frisch

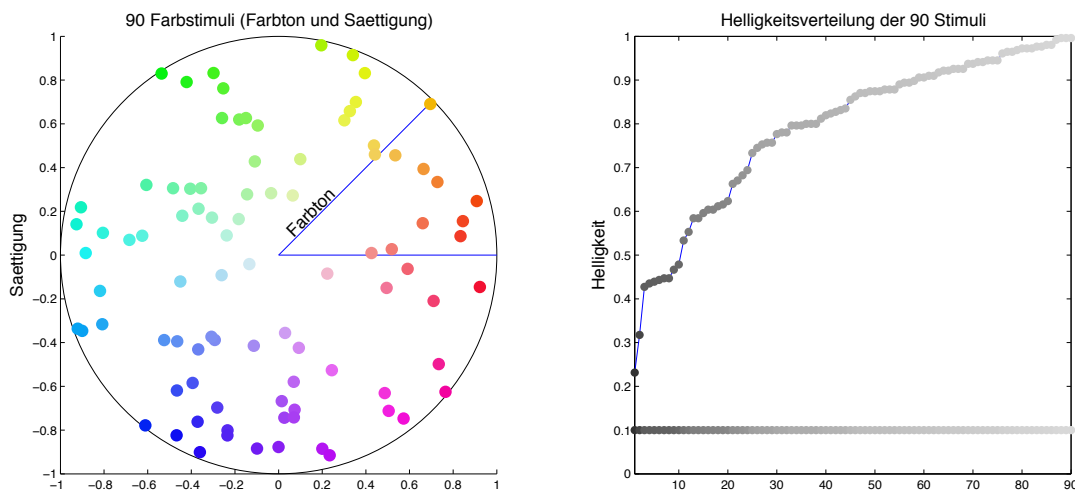
Die verbleibenden drei Beschreiber beschreiben die subjektive Wirkung der Farbe:

- angenehm
- unangenehm
- sinnlich

Ein schöner (und nicht ganz zufälliger) Nebeneffekt dieser Beschreiberliste ist, dass wir ungefähr dasselbe Verhältnis zwischen Stimuli und Beschreibern erhalten wie es auch bei den Geruchsdaten aus dem folgenden Kapitel vorliegt (851 Chemikalien, 171 Beschreibern).

#### 3.2.2 Durchführung der Studie

In der Datenerfassungsphase der Studie sollten nun die gegebenen 90 Farbstimuli von möglichst vielen Probanden mittels der eben vorgestellten 16 Farbbeschreiber



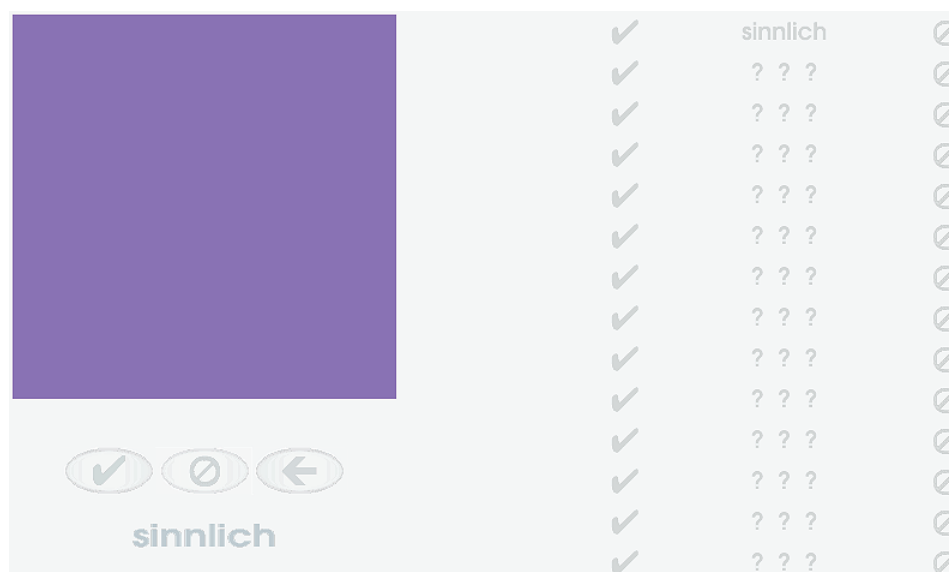
**Abbildung 3.2:** Links: Verteilung der 90 Farbstimuli entlang des Farbkreises mit dem Farbton (hue) als Winkel- und der Sättigung (saturation) als Radialkoordinate. Rechts: Helligkeitsverteilung der Farbstimuli.

charakterisiert werden. Hierbei sollten sie – ähnlich einem Phantombild – einen Stimulus mit einem Beschreiber vergleichen und sagen, ob sie die beiden als zugehörig empfinden.

Um möglichst viele Beschreibungen zu erhalten, haben wir das Experiment über das Internet verfügbar gemacht und Probanden nach einer ausführlichen Beschreibung des Protokolls die Möglichkeit gegeben, online Farben und Beschreiber miteinander in Verbindung zu bringen. Dies hat allerdings den Nachteil, dass wir keine Kontrolle über die Rahmenbedingungen des Experimentes (z.B. Farbeinstellung des Monitors) oder die Sorgfalt der Probanden bei der Durchführung der Befragung hatten.

Die meisten Teilnehmer der Studie haben wir persönlich zu dem Experiment eingeladen und in einem Aufklärungsgespräch zumindest Sehdefizite, wie z.B. Rot-Grün-Blindheit, ausgeschlossen. Dennoch mussten wir versuchen, zumindest durch die Breite an Teilnehmern und viele Beschreibungen der Farben solche Eventualitäten zu eliminieren.

Um die Probanden nicht zu überlasten und die Gefahr zu senken, dass sie vorzeitig die Befragung abbrechen, haben wir für jeden Teilnehmer die Studie in mehrere Sitzungen unterteilt. Im Vorfeld haben wir festgestellt, dass die Probanden etwa eine Minute benötigen, um eine Farbe mit allen 16 Beschreibern zu vergleichen. Weiterhin setzten nach etwa 15 Minuten erste Ermüdungserscheinungen bei den Testprobanden ein, sie begannen unkonzentriert und überhastet durch die Vergleichslisten zu klicken.



**Abbildung 3.3:** Screenshot der Benutzerschnittstelle. Oben links wird der Stimulus angezeigt, darunter befinden sich drei Schaltflächen mit den Wahlmöglichkeiten (v.L.n.R.): *Zutreffend* (Häkchen), *Nicht zutreffend* (Verbotsschild) und *Zurück* (Pfeil nach rechts). Direkt darunter ist der aktuelle Bezeichner zu lesen. Auf der rechten Seite des Bildschirms sind die bereits evaluierten Bezeichner zu sehen und die getroffene Entscheidung. Noch nicht evaluierte Bezeichner sind durch drei Fragezeichen unkenntlich gemacht.

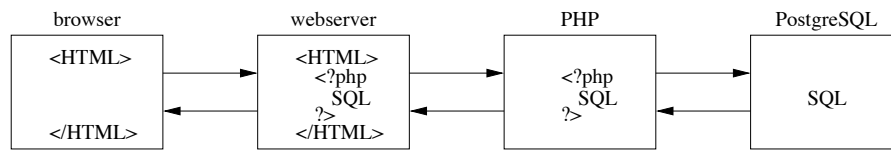
Folglich hatte jeder Teilnehmer die Möglichkeit, in insgesamt sechs 15-minütigen Sitzungen sämtliche 90 Stimuli mit allen 16 Beschreibern zu vergleichen. Da sowohl die Stimuli wie auch für jeden Stimulus die Beschreiber in zufälliger Reihenfolge dargeboten wurden, war es wenig kritisch, wenn Probanden nur an einem Teil der möglichen Sitzungen teilgenommen haben.

Für jeden präsentierten Farbstimulus wurde die Menge der Beschreiber ebenfalls permutiert, um einen möglichen systematischen Fehler aufgrund der Beschreiberreihenfolge auszuschließen.

### Technische Umsetzung

Die Benutzerschnittstelle (Abb. 3.3) bestand aus einer PHP-gestützten HTML-Seite, welche an eine PostgreSQL-Datenbank angebunden wurde. PHP ist eine Skriptsprache, die in HTML-Seiten eingebunden werden kann und sowohl die Experimentdaten aus der Datenbank holen, wie auch die Benutzereingaben direkt wieder in der Datenbank zur Weiterverarbeitung speichern konnte.

Abbildung 3.4 gibt noch einmal eine Übersicht über den hier vorliegenden Datenfluss. Weitere technische Details zu der Implementation finden sich in der zugehö-



**Abbildung 3.4:** Schematische Darstellung des Datenfluss für die Benutzerschnittstelle. Greift ein Nutzer über seinen Browser auf die Webseite zu, erfolgt eine Anfrage an den Webserver. PHP-Inhalte reicht der Webserver weiter an den PHP-Interpreter. Dieser kann dynamischen HTML-Code generieren, einschließlich Daten aus der Datenbank. Ebenfalls kann PHP Nutzerdaten in die Datenbank speichern. Der so generierte HTML-Code wird dann vom Webserver an den Browser des Nutzers gesendet und dort angezeigt.

rigen Studienarbeit von Martin Haker [25].

Eine Sitzung beginnt mit der Startseite, auf der sich die Probanden mit ihren Zugangsdaten anmelden müssen, um eine bereits begonnene Sitzung fortzusetzen oder - bei der ersten Anmeldung - mit einer Registrierungsseite, auf der erst noch ein paar persönliche Daten wie Geschlecht und Alter der Probanden und vor allem eine email-Adresse abgefragt werden. Für jeden neuen Nutzer erzeugt das System anschließend eine Permutation der Farbstimuli, die die Reihenfolge festlegt, in der der entsprechende Nutzer später die Farben zur Bewertung vorgelegt bekommt.

Diese Liste wird unter der persönlichen ID des Probanden in der Datenbank abgelegt, um ihm bei weiteren Sitzungen keine Farben mehrfach zu präsentieren. Diese ID ist gleichzeitig auch Teil der Zugangsdaten und wird an die angegebene email-Adresse versendet.

## Die Benutzerschnittstelle

Nach dem ersten Einloggen meldet sich das System mit einer kurzen Einführung in die Studie und erklärt die Benutzung der Schnittstelle. Sie muss zwei Anforderungen erfüllen: Erstens soll sie die Benutzer einfach und möglichst schnell durch die Vergleiche leiten, so dass die Probanden möglichst viele Farben bewerten können ohne frühzeitig zu ermüden. Das zweite Kriterium ist, dass das Interface die Nutzer nicht ablenken oder gar den Ausgang des Experimentes beeinflussen darf.

Entsprechend haben wir versucht, die Farben möglichst neutral zu wählen. Ein helles Grau für den Hintergrund und ein etwas dunkleres Grau für die Auswahlknöpfe und Informationstexte mit möglichst wenig Kontrast erschienen uns hier als gut geeignet.

Abbildung 3.3 zeigt die Benutzerschnittstelle. Dominiert wird das Fenster natürlich

durch das (320x320) Pixel große Viereck in der linken oberen Ecke. Dies ist der zu bewertende Stimulus. Direkt darunter finden sich die eigentlichen Bedienelemente, drei Auswahlknöpfe sowie direkt darunter der Beschreiber (hier: *sinnlich*, mit dem der Stimulus verglichen werden sollte. Die Auswahlknöpfe waren wie folgt belegt:

1. *Zutreffend*: Mit einem Häkchen ist der Bestätigungsknopf versehen. Diesen Knopf soll der Proband drücken, falls er den angezeigten Begriff als stimmig mit der gezeigten Farbe empfindet.
2. *Nicht zutreffend*: Das Verbotsschild signalisiert, dass man den Beschreiber als nicht zu dem Stimulus passend empfindet.
3. *Zurück*: Der Pfeil nach rechts signalisiert eine Möglichkeit, zu der letzten Bewertung zurückzugelangen. Dies dient zur Korrektur von Falscheingaben. Man kann genau einen Schritt rückgängig machen.

Auf der rechten Seite des Bildschirms ist der Fortschritt der aktuellen Sitzung zu sehen. Man kann dort den aktuellen Bezeichner sehen und die zuvor abgegebene Beurteilung. Die noch nicht bewerteten Bezeichner sowie die noch weiter zurückliegenden Bezeichner sind durch drei Fragezeichen ersetzt und können nicht mehr eingesehen werden.

Die letzte Entscheidung ist noch sichtbar, um eventuelle Bedienfehler oder Fehlentscheidungen direkt erkennen zu können und durch die undo-Funktion zu korrigieren. Alles darüber hinaus ist verdeckt, um einen möglichen Bias durch kausale Entscheidungen zu verhindern. Ein Proband könnte zum Beispiel einen Bezeichner als zutreffend eintragen, da er zuvor schon einmal ein ähnliches Wort ebenfalls als zutreffend bewertet hat.

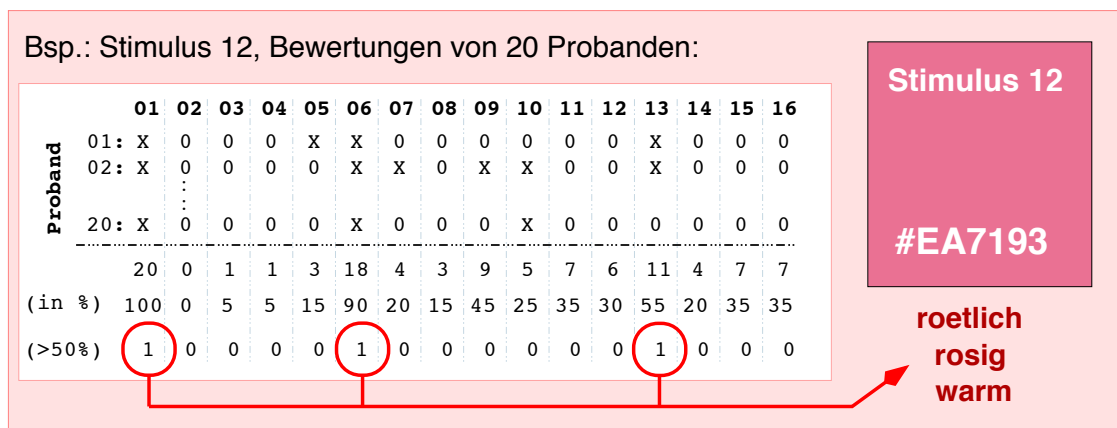
Hat der Proband eine Runde (also die Charakterisierung einer Farbe mit allen 16 Bezeichnern) abgeschlossen, wird der entsprechende Datensatz an die Datenbank geschickt und die nächste Farbe ausgelesen. Auf diese Art und Weise wird sichergestellt, dass keine unvollständigen Beschreibungsrunden in die Daten aufgenommen werden.

Nach 15 Farbstimuli – also nach 240 Vergleichen – beendet das System die Befragung und empfiehlt dem Nutzer, eine Pause zu machen.

#### 3.2.3 Ergebnisse

Im Rahmen dieses Projektes haben sich insgesamt 75 Personen in das CPP-System eingetragen und Farben evaluiert. Die meisten haben nicht alle 6 Sitzungen absolviert, sondern meist nur eine oder zwei Sitzungen abgeschlossen. Im Mittel haben die Teilnehmer etwa die Hälfte der 90 Farbstimuli gegen die 16 Bezeichner verglichen. Damit haben wir unser ursprüngliches Ziel, ungefähr 40 Teilnehmern in





**Abbildung 3.5:** Datenaufarbeitung am Beispiel von Farbstimulus 12. Der Kasten rechts oben zeigt den Stimulus 12 aus unserer Studie. 20 Probanden haben die Farbe mit den 16 Bezeichnern verglichen. Somit lagen 20 Einträge wie die hier exemplarisch gezeigten Zeilen vor. Zutreffende Bezeichner sind mit X vermerkt. Liegt die relative Häufigkeit eines Bezeichners über 50%, so wird der Bezeichner nach der Binarisierung mit einer 1 vermerkt, ansonsten mit einer 0.

jeweils 3 Sitzungen etwa 20 Beschreibungen pro Stimulus durchführen zu lassen, dennoch erreicht. Insgesamt wurde jede der 90 Farben im Schnitt von ca. 20.4 Personen beschrieben. Die Farbe mit den wenigsten Beschreibungen wurden von 18 Personen bearbeitet, die mit den meisten von 23.

Mit Hilfe dieser Beschreibungen wollen wir nun abschätzen, wie zutreffend ein gegebener Bezeichner für eine der zufälligen Farben ist. In Abbildung 3.6 ist ein Histogramm zu sehen, welches zeigt, wie viele Merkmale pro Farbstimulus von mehr als 50% der Probanden als *zutreffend* erachtet wurden. Den meisten unserer Stimuli wurden von den Probanden etwa 3-4 Bezeichner als zutreffend zugeordnet, nur für einige wenige gab es nur einen einzigen zutreffenden Bezeichner. Dies spricht auch dafür, dass wir eine vertretbare Auswahl an Bezeichnern getroffen haben.

## Hauptkomponentenanalyse

Ganz allgemein liefert uns eine Studie, wie die eben vorgestellte, eine große Datenmatrix  $\mathbf{R} \in \mathbb{R}^{n \times m}$ , wobei  $n$  der Anzahl der verwendeten Stimuli entspricht (hier also  $N = 90$  Farben) und  $m$  der Anzahl der möglichen verbalen Deskriptoren (hier also  $m = 16$  Farbbeschreiber) mit

$$\mathbf{R} = \begin{pmatrix} \mathbf{r}_1 \\ \vdots \\ \mathbf{r}_{90} \end{pmatrix}$$

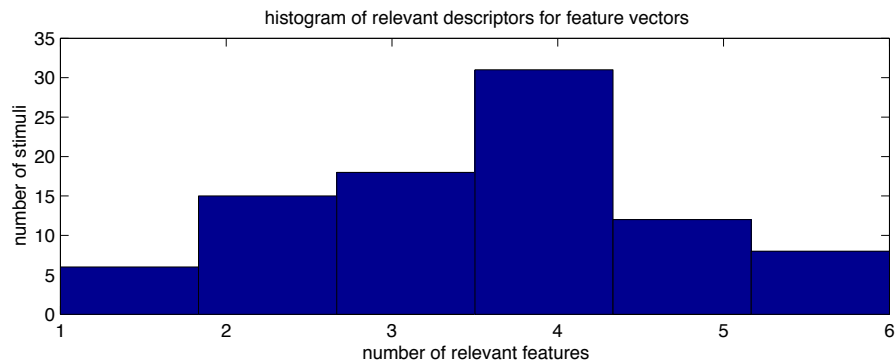
in der jeder Farbstimulus  $i$  durch den  $m$ -dimensionalen Zeilenvektor  $\mathbf{r}_i \in \mathbb{R}^m$  beschrieben wird. Die Einträge ergeben sich direkt aus den relativen Häufigkeiten der Probandenbewertungen. Sei  $\mathbf{P} \in \{0, 1\}^{p \times m}$  die Bewertungsmatrix des Stimulus  $i$  nach der Bewertung durch  $p$  Probanden, so ergibt sich – wie in Abbildung 3.5 illustriert – für jeden Eintrag  $j$  von  $\mathbf{r}_i$

$$\mathbf{r}_i(j) = \frac{1}{p} \sum_{k=1}^p \mathbf{P}_i(k, j)$$

In einer geometrischen Interpretation hat jeder Farbstimulus  $i$  durch den Vektor  $\mathbf{r}_i$  einen Repräsentanten im  $\mathbb{R}^m$ . Wir wollen im folgenden versuchen, etwas über die Struktur und vor allem über die Komplexität dieses Merkmalsraumes zu erfahren. Wie viele Faktoren gibt es, die diesen Raum aufspannen? Dies können natürlich nicht mehr als  $m = 16$  sein, da die intrinsische Dimension, also die Dimension des Raumes, den die Datenpunkte tatsächlich aufspannen, durch die extrinsische Dimension, also die Ausgangsdimension des Raumes beschränkt ist.

Ein typisches Werkzeug, um die Komplexität von Daten zu untersuchen und ob sie eventuell in einem niederdimensionalen Unterraum liegen, ist die sogenannte Hauptkomponenten-Analyse (Principal Component Analysis, PCA [44]). Bei der PCA werden die 16 Dimensionen des extrinsischen Raumes  $\mathbf{R}$  auf ihre Korrelation untersucht, in dem man die Eigenwerte der Kovarianzmatrix  $\mathbf{R}^T \mathbf{R}$  berechnet. Anhand der Eigenwerte, die signifikant größer als Null sind, kann abgeschätzt werden, wie viele Dimensionen die Punkte in  $\mathbf{R}$  tatsächlich aufspannen. Die zugehörigen Eigenvektoren nennt man auch die *Hauptkomponenten* der Daten. Diese Vektoren zeigen in Richtung größter Varianzen im Ausgangsraum.

Wir haben die PCA bereits genauer im Rahmen der ICA-Analyse in Kapitel 2 beleuchtet. Daher werden wir hier nicht tiefer in die Eigenwert-Zerlegung der Kovarianzmatrix eingehen. Wichtig ist für uns, dass die Größe der Eigenwerte direkt abhängt von der Varianz in Richtung des zugehörigen Eigenvektors. D.h., ist der Eigenwert Null, so unterscheiden sich die Daten nicht in Richtung dieses Eigenvektors. Entsprechend ist dies dann vermutlich auch kein Merkmal, welches helfen wird, Eigenschaften und Unterschiede der Punkte herauszuarbeiten. Deshalb sind die Eigenwerte ebenfalls ein Maß für die Relevanz des zugehörigen Eigenvektors. In der Praxis werden die Eigenwerte typischerweise ihrer Größe nach sortiert. Als relevante Faktoren betrachtet man die Vektoren, deren Eigenwerte aufsummiert



**Abbildung 3.6:** Histogramm über die Anzahl der zutreffenden Bezeichner je Stimulus. Ein Bezeichner ist zutreffend für einen Stimulus, wenn mehr als 50% der befragten Personen ihn zutreffend fanden. Durchschnittlich wurden die meisten Farbstimuli mit Hilfe von vier Bezeichnern beschrieben.

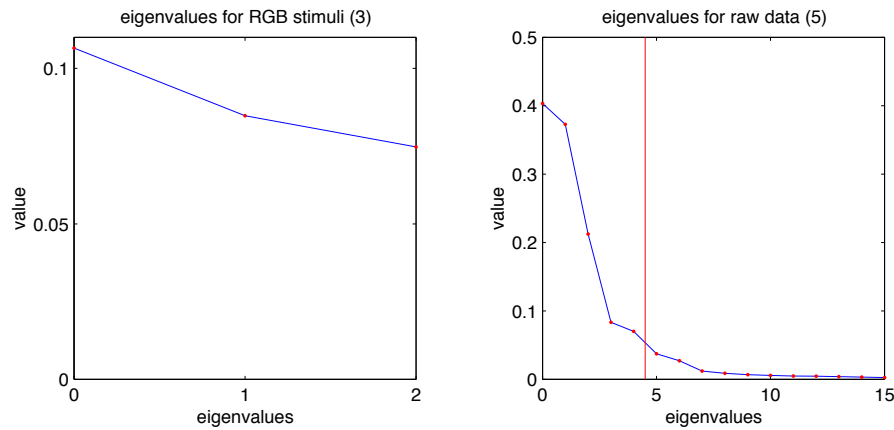
(beginnend beim größten) etwa 90% der Summe aller Eigenwerte ergeben.

Abbildung 3.7 zeigt nun die Eigenwerte sowohl für unsere 90 Farbstimuli in ihrer dreidimensionalen Originalform (links) und die Eigenwerte für die Rohdaten unseres Experimentes (rechts). Da wir die Ausgangsdaten basierend auf der gleichen Zufallsverteilung für alle drei Farbkanäle ausgewählt haben, sehen wir in der linken Abbildung, was wir auch erwarten würden, drei nahezu gleichberechtigte Eigenwerte für alle drei Dimensionen. Da wir durch unsere diskrete Stichprobe nicht in alle Richtungen genau dieselbe Varianz erhalten, sind natürlich auch die Eigenwerte nicht genau gleich.

Dies sieht für unsere experimentellen Daten (rechts) anders aus: Hier erkennt man einen ganz deutlichen Abfall der Eigenwerte. Sind die ersten beiden Eigenwerte noch sehr groß, werden die folgenden sehr schnell kleiner. Bereits der 8. Eigenwert ist – verglichen mit den ersten – quasi Null. Die Summe der ersten 5 Eigenwerte (0-4) übersteigt bereits 90% der Gesamtsumme. Man würde hier also eine intrinsische Dimension von etwa 5 erwarten, bei einer deutlichen Dominanz der ersten beiden Faktoren.

## Binarisierung der Daten

Bei psychophysikalischen Experimenten muss man sich immer mit dem Problem auseinandersetzen, dass die Beurteilungen der Probanden durchaus differieren können. Nicht immer bzw. in den wenigsten Fällen sind die Ergebnisse ganz eindeutig. Darum ist es wichtig, dass sich die Ergebnisse auf das Urteil von möglichst vielen Probanden beziehen, um einzelne abweichende Meinungen und Empfindungen herauszufiltern und im Mittel die verallgemeinerbare Information zu finden.



**Abbildung 3.7:** Links: Eigenwerte der zufällig gewählten RGB Farbstimuli. Die Daten liegen in einem dreidimensionalen Ausgangsraum, daher gibt es nach der PCA auch drei Eigenwerte. Direkt aus der Konstruktion der Daten folgend erhält man drei nahezu gleichgroße Eigenwerte. Rechts: Eigenwerte nach der PCA auf den Rohdaten aus der Datenmatrix  $\mathbf{R}$ . Deutlich sieht man den Abfall der Eigenwerte, sie sind für Dimension 8 bereits nahezu Null. 90% der Gesamtsumme ergeben sich durch Aufsummieren der ersten 5 Dimensionen.

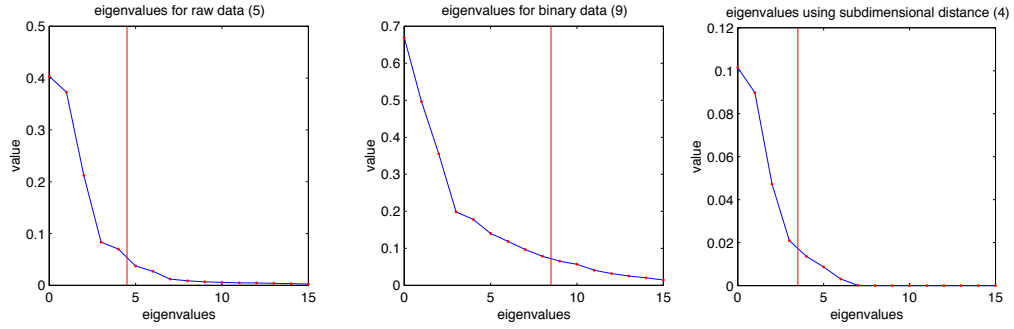
Für diese Studie haben wir an die relativen Häufigkeiten der Bewertungen einen Schwellwert von 50% angelegt und damit eine binarisierte Datenmatrix  $\mathbf{D} \in \mathbb{R}^{n \times m}$  erzeugt mit

$$\mathbf{d}_i(j) = \begin{cases} 1 & \text{falls } \mathbf{r}_i(j) \geq 0.5 \\ 0 & \text{sonst.} \end{cases}$$

D.h. haben für einen gegebenen Stimulus mehr als die Hälfte aller Probanden einen Bezeichner als zutreffend beschrieben, so haben wir diesem Bezeichner eine 1 zugeordnet, ansonsten eine 0. Dies ist auch in Abbildung 3.5 bereits angedeutet.

Betrachten wir nun auch die Struktur dieser Datenmatrix  $\mathbf{D}$ , so können wir, wie in Abbildung 3.8 zu sehen, folgendes beobachten: Die Eigenwerte für die binarisierten Daten bleiben länger groß und die Komplexität der Daten hat sich somit offensichtlich erhöht. Auch wenn dies auf den ersten Blick erstaunen mag, so kann man es sich an folgendem Beispiel leicht verdeutlichen: Punkte, die alle entlang einer Geraden (also eindimensional) in einem zweidimensionalen Raum eingebettet sind, denaturieren nach einer Binarisierung auf die Ecken eines Vierecks und werden somit echt zweidimensional, also komplexer als sie es eigentlich waren.

Offensichtlich verlieren wir durch die Binarisierung Strukturinformation der Daten, vor allem, falls wir eine euklidische Metrik für die eigentlichen Berechnungen verwenden. Oftmals jedoch kann man es sich gar nicht aussuchen, dann liegen die



**Abbildung 3.8:** Links: Die Eigenwerte nach PCA auf den relativen Häufigkeiten in  $\mathbf{R}$ , siehe auch Abb. 3.7. Mitte: Eigenwerte nach PCA auf der binarisierten Matrix  $\mathbf{D}$ . Man kann deutlich sehen, dass die Eigenwerte wesentlich langsamer kleiner werden. Die Datenstruktur ist komplexer geworden. Rechts: Eigenwerte der Ähnlichkeitsmatrix  $D_{\text{subsim}}$  nach PCA auf MDS-eingebetteten Punkten. Im Vergleich zur PCA auf  $\mathbf{R}$  erkennt man ein ähnliches Profil, die 90% Eigenwertenergie ist allerdings bereits für 4 Dimensionen erreicht.

Daten bereits in binarisierter Form vor.

## Subdimensional Distance

In vorausgegangenen Arbeiten (wiederum für den Geruchssinn) haben wir eine Metrik eingeführt, die solche binarisierten Daten gut beschreiben kann [56, 57]. Es handelt sich hierbei um eine Art gewichtete Hamming-Distanz bzw. eine gewichtete Variante einer Crossentropie-Abschätzung [15]. Sie hat sich durch ihre Gewichtung besonders bei Beschreibungen von Geruchswahrnehmung bewährt, da hier teilweise einige Beschreiber (wie z.B. *fruity*) wesentlich häufiger zur Beschreibung verwendet werden als andere (z.B. *grapefruit*). Dies macht den direkten Vergleich über eine Hamming-Distanz oder die Crossentropie schwierig. Die Subdimensionale Distanz  $d_s$  ist definiert als:

$$d_s(\mathbf{d}_x, \mathbf{d}_y) = \frac{\sum_{i=1}^n (|\mathbf{d}_x(i) - \mathbf{d}_y(i)| \cdot \mathbf{d}_x(i)) + \frac{\sum_{i=1}^n (|\mathbf{d}_x(i) - \mathbf{d}_y(i)| \cdot \mathbf{d}_y(i))}{\sum_{i=1}^n \mathbf{d}_y(i)}}{\sum_{i=1}^n \mathbf{d}_x(i) + 1} \quad (3.1)$$

wobei  $\mathbf{d}_j = (\mathbf{d}_j(1), \dots, \mathbf{d}_j(n))$  das  $n$ -dimensionale binäre Beschreibungs-Profil für einen Stimulus  $j$  ist.  $\mathbf{d}_x$  sei hierbei im Vergleich zu  $\mathbf{d}_y$  der weniger prominent beschriebene Reiz (beinhaltet also weniger Einträge mit 1).

Es handelt sich hierbei nicht mehr um eine Metrik, da die Dreiecks-Ungleichung nicht notwendigerweise erfüllt ist. Ein solches Ähnlichkeitsmaß nennt man Semimetrik. Wir erzeugen also eine Ähnlichkeitsmatrix  $D_{\text{subsim}} \in \mathbb{R}^{n \times n}$  für unsere  $n = 90$

Farbstimuli mit

$$D_{\text{subsim}} = \begin{pmatrix} d_s(\mathbf{d}_1, \mathbf{d}_1) & \dots & d_s(\mathbf{d}_1, \mathbf{d}_n) \\ \vdots & \ddots & \vdots \\ d_s(\mathbf{d}_n, \mathbf{d}_1) & \dots & d_s(\mathbf{d}_n, \mathbf{d}_n) \end{pmatrix} \quad (3.2)$$

Das Semimetrische bringt mit sich, dass wir nicht mehr einfach die PCA anwenden können, da hierbei implizit die euklidische Metrik verwendet wird. Es gibt aber eine Methode, das sogenannte Multidimensional Scaling (MDS, [85, 100]) mit der wir eine gegebene nicht-metrische Ähnlichkeitsmatrix – wie z.B.  $D_{\text{subsim}}$  – wiederum in einen euklidischen Raum beliebiger Dimension einbetten können.

Dabei werden  $n$  zufällige Punkte  $\mathbf{m}_i \in \mathbb{R}^q, i = 1, \dots, n$  im  $q$ -dimensionalen (euklidischen) Zielraum (mit  $q < p$ ) so verschoben, dass die Abweichung zwischen den Distanzen dieser Punkte und der Ausgangsmatrix minimal wird. Auf diese eingebetteten euklidischen Punkte kann man dann eine PCA anwenden. Eine knappe Darstellung des Multidimensional Scaling findet sich im Anhang A.1.

Es ergibt sich dann eine neue Darstellung unserer Daten durch die eingebettete Punktmatrix  $\mathbf{M} \in \mathbb{R}^{n \times q}$  mit

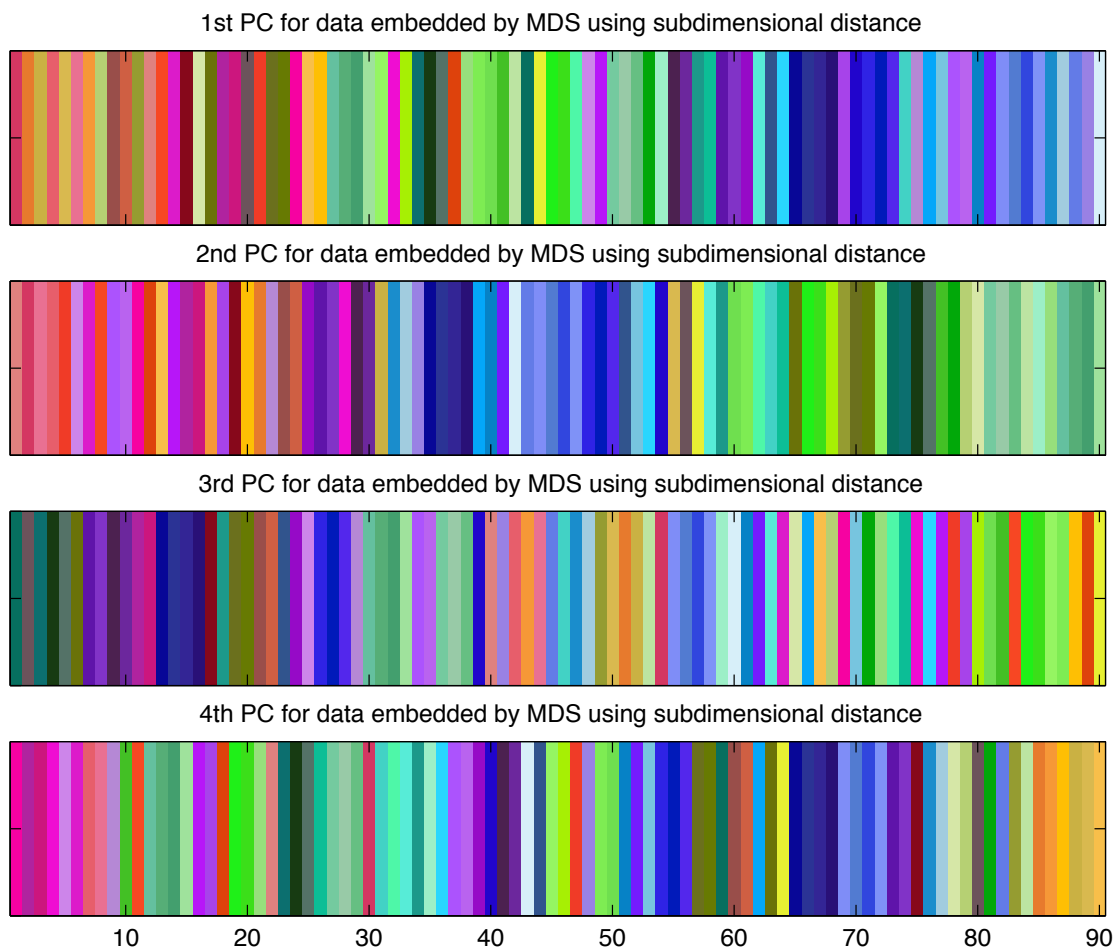
$$\mathbf{M} = \begin{pmatrix} \mathbf{m}_1 \\ \vdots \\ \mathbf{m}_n \end{pmatrix}.$$

Jeder Zeilenvektor  $\mathbf{m}_i$  von  $\mathbf{M}$  repräsentiert also die  $q$ -dimensionale Einbettung des Beschreibungsvektors  $\mathbf{d}_i$  unter Berücksichtigung seiner subdimensionalen Ähnlichkeiten zu allen anderen Beschreibungsvektoren.

Wie in Abbildung 3.8 zu sehen ist, kommen wir mit der Datenbeschreibung mittels der vorgestellten Metrik auf eine Abschätzung von 4 Dimensionen. Dies ist zwar durchaus in derselben Größenordnung, wie wir es auch für die metrischen Daten vor der Binarisierung erhalten, allerdings werden wir in Kapitel 4 Daten betrachten, die bereits binarisiert vorliegen. Daher wollen wir im weiteren die Datenanalyse auf den binarisierten und eingebetteten Farbdaten weiterführen.

## Der Merkmalsraum der Farben

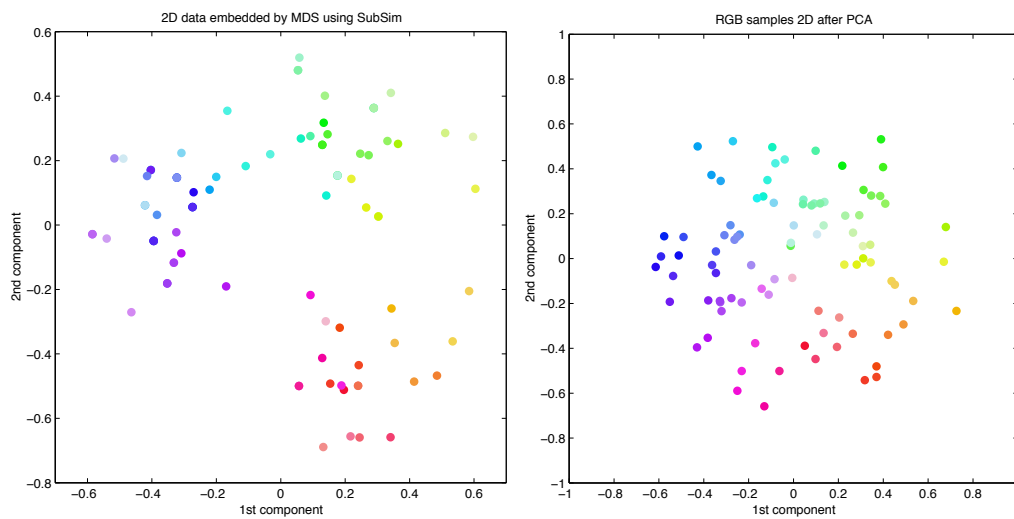
Um die Daten zu beschreiben, betrachten wir im weiteren nun also die binarisierte Beschreibungsmatrix  $\mathbf{D}$ , welche mittels der SD-Semimetrik in die Distanzmatrix  $D_{\text{subsim}}$  überführt und anschließend mit MDS wieder in einen 4-dimensionalen Raum eingebettet wurde. In Abbildung 3.9 sind die vier Hauptkomponenten der resultierenden Punktmatrix  $\mathbf{M}$  dargestellt. In den vier Zeilen sind jeweils alle 90 Stimuli in aufsteigender Reihenfolge entlang der jeweiligen Hauptkomponente ge-



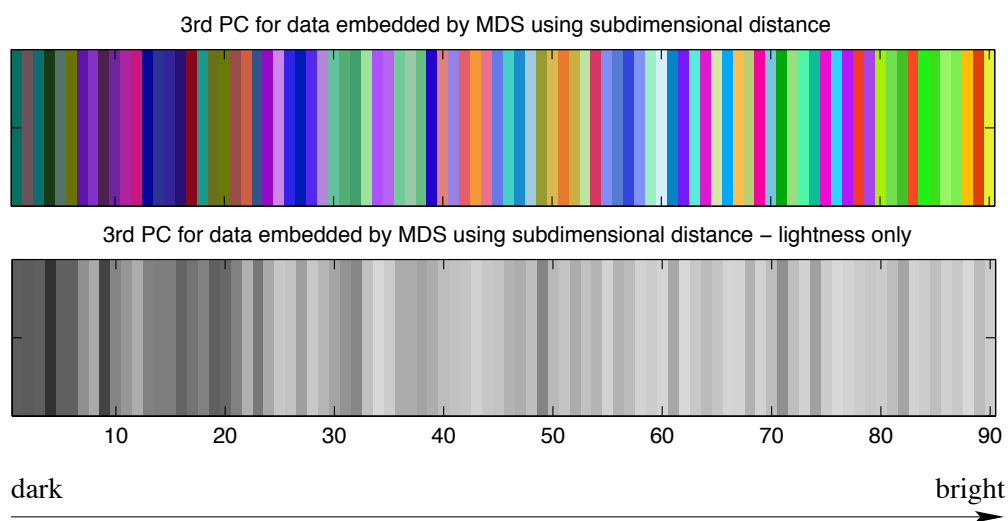
**Abbildung 3.9:** Die Hauptkomponenten der 90 eingebetteten Punkte aus **M**. Zeilenweise sind die Stimuli äquidistant und aufsteigend entlang der Hauptkomponenten (PCs) aufgezeigt. Die ersten beiden PCs zeigen eine Art Dreiteilung: Für die erste PC laufen die Farben von *rot* über *grün* nach *blau*, für die zweite PC von *rot* über *blau* zu *grün*. Trotz einer gewissen Ordnung ist eine spontane Deutung der PC schwierig. Bei der dritten PC scheinen die Farben von gedämpft nach grell zu verlaufen. Für die vierte ist nicht unmittelbar ein Kontinuum zu erkennen.

zeigt. So kann man qualitativ einen guten Überblick darüber erhalten, wie die Punkte zueinander liegen.

Die ersten beiden Hauptkomponenten zeigen eine relativ deutliche Organisation um die Farbtöne *rot*, *grün* und *blau*, jeweils in unterschiedlicher Reihenfolge. Mit dem Vorwissen, dass wir im Wahrnehmungsraum einen zirkulären Schluss haben, also einen Farbkreis, glaubt man auch hier diesen Zusammenhang zu erahnen: Einige der Farben, die durch die erste Achse aufgespannt werden (von *blau-lila* nach *gelb-orange*) scheinen in der zweiten Hauptkomponente zu kollabieren. Ob man aber auch ohne das Vorwissen sofort auf diese Idee gekommen wäre, sei dahinge-



**Abbildung 3.10:** Links: Die ersten beiden Hauptkomponenten von  $M$  gegeneinander geplottet. Die Farbe jedes Punktes  $m_i$  entspricht dem Farbwert des zugehörigen RGB-Stimulus  $i$ . Rechts: Zum Vergleich die ersten beiden Hauptkomponenten für die 90 RGB-Farbstimuli.



**Abbildung 3.11:** Oben: Entspricht Abb. 3.9. Die Farben scheinen von dunkel nach hell zu verlaufen. Unten: Wie in Abb. 3.2 wurden die RGB-Farbstimuli entfärbt, so dass nur noch die Helligkeit bestehen bleibt. Man erkennt den bereits vermuteten Verlauf nun wesentlich deutlicher.

stellt. Tragen wir die ersten beiden Komponenten gegeneinander auf, so erhalten wir tatsächlich eine subjektiv sehr harmonische Farbanordnung. Dies ist in Abbildung 3.10 (links) dargestellt, im Vergleich zu der Ground Truth, den Farbstimuli in Farbkreisdarstellung.

Betrachten wir die dritte Hauptkomponente, so fällt zunächst einmal auf, dass die



---

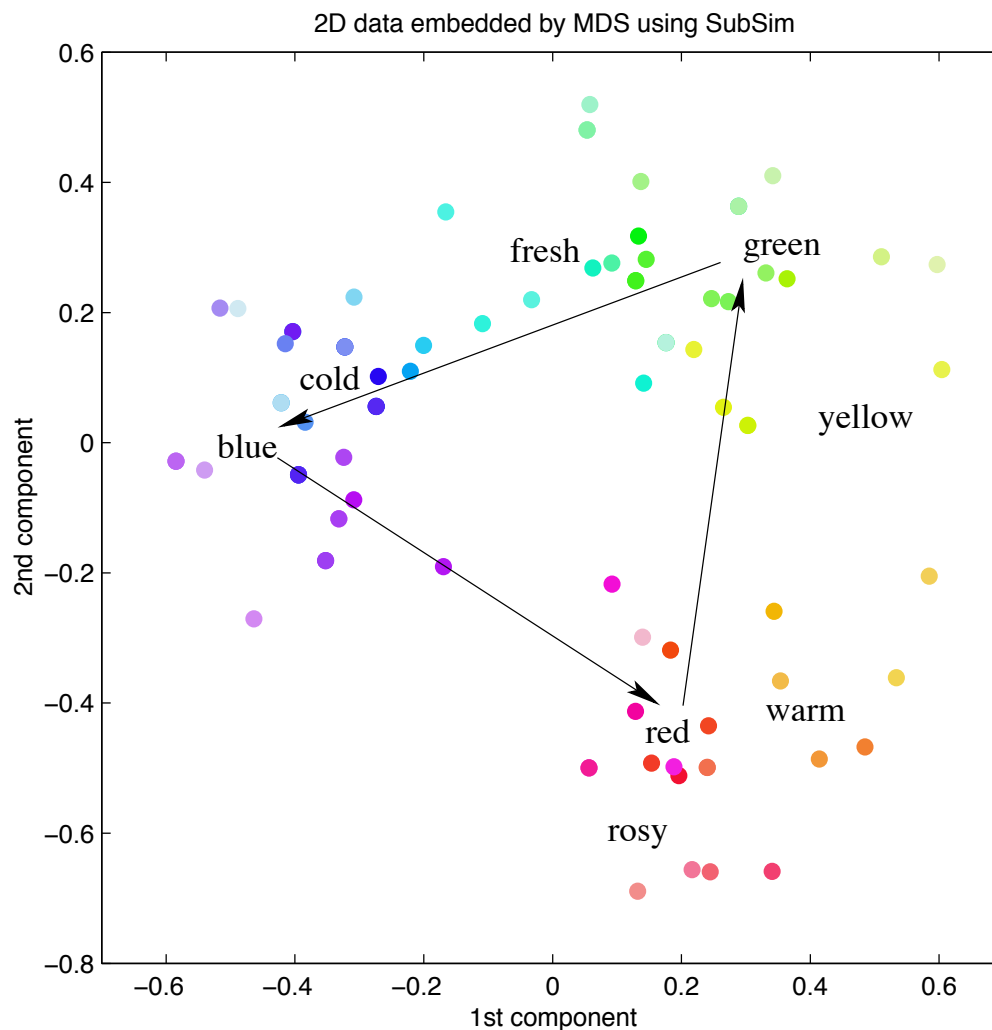
Farben rechts gedeckter wirken, nach links wird der Gesamteindruck der Farben eher grell. Die Achse scheint also von dunklen Farben bis zu recht hellen Farben zu gehen. In Abbildung 3.11 haben wir auch dies vergleichend dargestellt: Wir haben die dritte Hauptkomponente entfärbt, d.h. wir haben die Sättigung entfernt und erhalten somit Grauwerte unterschiedlicher Helligkeiten. Der Übergang ist nicht ganz glatt, die Grundtendenz jedoch – ein Verlauf von dunklen Farben zu hellen Farben – lässt sich darauf durchaus erkennen.

Die vierte Dimension ist nicht so leicht darzustellen. Hieran kann man aber wiederum gut erkennen, wie schwierig die Interpretation eines Merkmalsraumes ist, wenn einem subjektiv nicht gleich ein Kontinuum auffällt. Betrachten wir auch hier die Extremstellen der Achsen so stehen sich da grelle Pink- und Grüntöne auf der einen Seite und eher gedeckte, herbstliche Orange- und Fliedertöne gegenüber. Das Problem jedoch bleibt: Welche Systematik kennen wir noch, um Farbunterschiede systematisch zu beschreiben? Selbst wenn wir also eine Idee entwickeln oder eine Systematik auf den Bildern zu erkennen vermögen, wie können wir den Zusammenhang dann nachweisen?

Gehen wir also noch einmal einen Schritt zurück und betrachten die ersten beiden Achsen. In Abbildung 3.12 haben wir einmal alle typischen subjektiven Interpretationen per Hand eingetragen. Es gibt eine Gruppe mit Rottönen im unteren rechten Teil des Plottes, der sich in zwei Richtungen wandelt. Schräg nach links werden die Töne erst *rosa* dann *lila* und münden in eine Gruppe von Blautönen. Der andere Pfad verläuft über *orange* und *gelb* in eine Gruppe von Grüntönen. Die Grüntöne wiederum verlaufen über *türkis* ebenfalls in die Gruppe von Blautönen. Während sich bei den Rottönen die meisten *warmen* Farben gruppieren, finden sich zwischen den Grün- und Blautönen die *frischen* bzw. die *kalten* Farben.

### Warum glauben wir das so genau zu wissen?

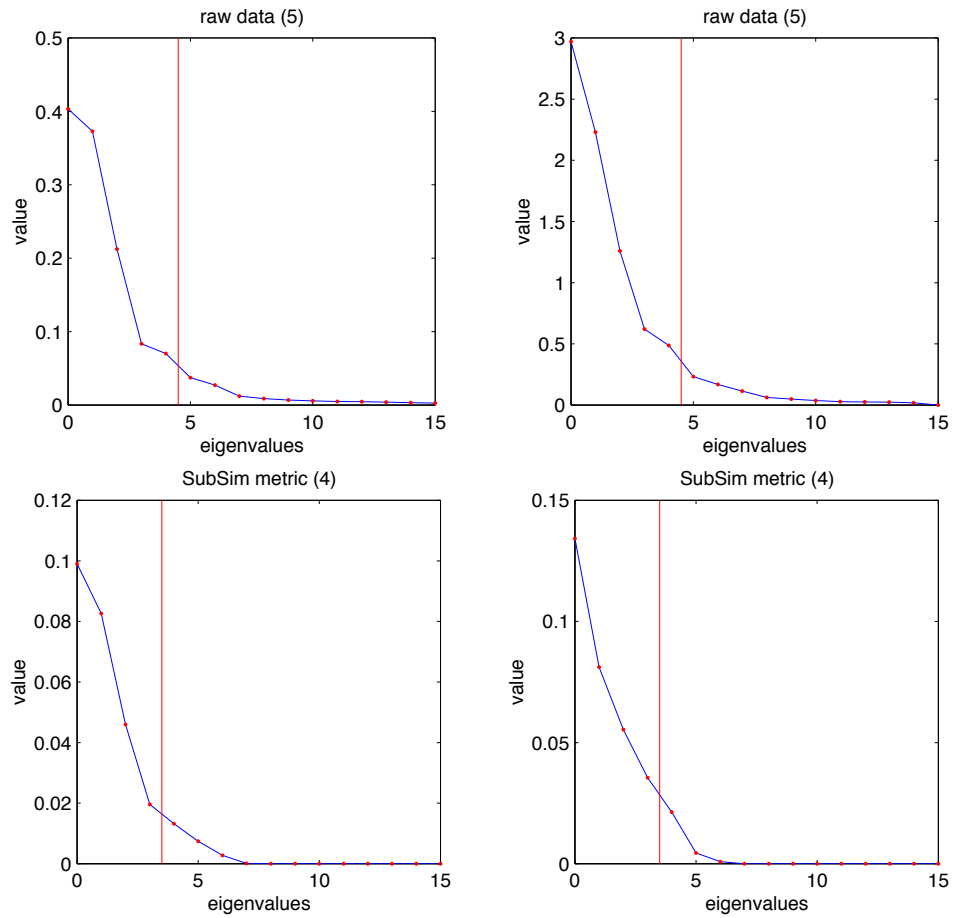
Warum fällt es uns wesentlich leichter, Abbildung 3.10 als sinnvoll zu akzeptieren als die ersten beiden Zeilen von Abbildung 3.9? Interessanterweise verwenden wir die typischen Wahrnehmungsbeschreiber, um unsere Wahrnehmung beim Betrachten der Zusammenhänge auszudrücken. Andererseits ist dies nicht überraschend, denn mit welchen Worten sollen wir unsere Wahrnehmungen denn sonst beschreiben. Spätestens jetzt sollte uns aber ein wesentlicher Aspekt dieses Farbenexperimentes auffallen: Die gesamte Zeit über haben wir die Stimuli direkt vor den Augen, die Stimuli, um die es im Experiment geht und die wir einordnen wollen. Stellen wir uns aber vor, wir hätten dasselbe Experiment mit Gerüchen gemacht. Dann würden die Bilder vermutlich nicht „riechen“ und wir hätten dort bestenfalls die Molekülnamen der Substanzen geplottet. Dann wäre (zumindest den meisten) der Zusammenhang zwischen den Beschreibungen und den Stimuli in Abbildung 3.12 vermutlich nicht mehr so deutlich.



**Abbildung 3.12:** Das Ergebnis aus 3.10 ist per Hand mit Beschriftungen versehen worden, die den typischen spontanen Eindrücken beim Betrachten dieser Darstellung nachempfunden sind. Man erkennt ganz deutlich den kontinuierlichen Übergang der Rottöne in die Blautöne und von dort über die Grün- und Gelbtöne wieder zurück zu den Rottönen. Auch Beschreibungen wie *frisch*, *kalt* und *warm* lassen sich gut platzieren.

Wir interpretieren also eigentlich nicht, wir vergleichen das Ergebnis mit dem, was wir bereits über die Farben wissen. Selbst wenn der anfangs erwähnte griechische Philosoph ebenfalls an diesem Punkt angekommen wäre und er dann sogar noch die ersten beiden PCs gegeneinander plotten würde: Auch er würde bereits „Expertenwissen“ einfließen lassen, bereits beim Sichten der Ergebnisse.

Das, was wir in Abbildung 3.12 bereits skizziert und in Abschnitt 3.1.2 bereits an-diskutiert haben, ist das gesuchte Ergebnis: Eine Darstellung des *Wahrnehmungs-raumes*, mit den Bezeichnern als Punkten und aufgespannt durch die Farbstimuli. Dies soll im folgenden auch praktisch umgesetzt werden.

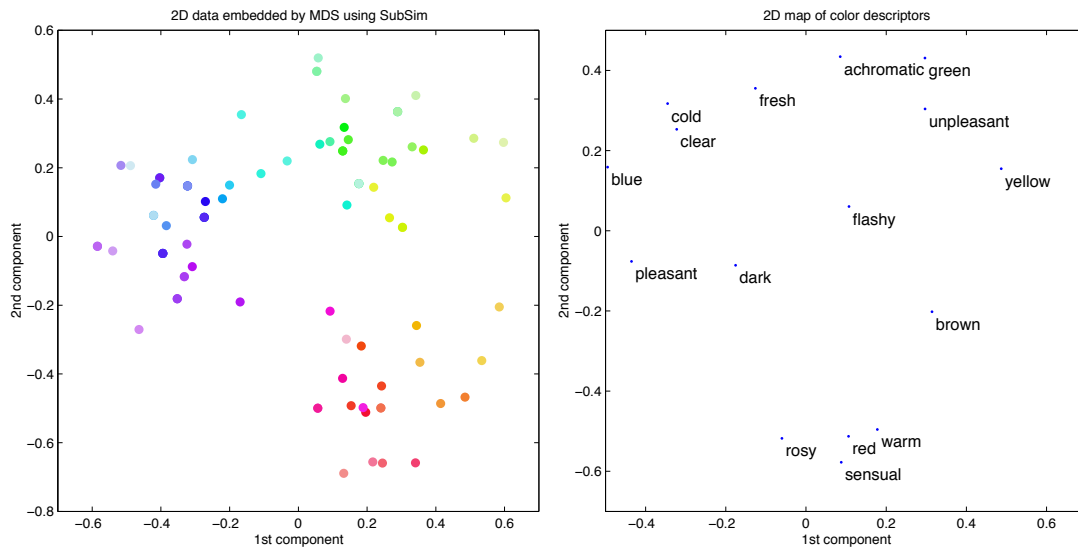


**Abbildung 3.13:** Eigenwertverlauf für  $\mathbf{D}$  und  $\mathbf{D}^T$  nach der PCA. **Oben:** Eigenwerte auf den Rohdaten, links für  $\mathbf{D}$ , rechts für  $\mathbf{D}^T$ . Jeweils für 5 Dimensionen werden 90% der Eigenwertsumme erreicht. **Unten:** Eigenwerte auf eingebetteten SD-Ähnlichkeiten, wieder links für  $\mathbf{D}$ , rechts für  $\mathbf{D}^T$ . 90% der Eigenwertsumme werden für beide nach Dimension 4 erreicht.

## Der Farbwahrnehmungsraum

Um nun den Farbwahrnehmungsraum darzustellen, bedarf es nicht mehr vieler Überlegungen. Wir verwenden wieder die binarisierten Ausgangsdaten  $\mathbf{D} \in \mathbb{R}^{n \times m}$ , jedoch transformieren wir sie nun zu einer Matrix  $\mathbf{D}^T \in \mathbb{R}^{m \times n}$ . Dadurch erhalten wir statt der  $n = 90$  Punkte in einem  $m = 16$  dimensionalen Raum (die Farbpunkte aufgespannt durch die Beschreiber) nun 16 Punkte in einem 90-dimensionalen Raum (16 Beschreiber aufgespannt durch die 90 Farbstimuli).

Die Datenanalyse folgt analog: Eine PCA zur ersten Abschätzung der Datenkomplexität, dann eine Ähnlichkeitsabschätzung mittels der SD-Semimetrik und abschließend eine erneute Einbettung in einen niedrig-dimensionalen euklidischen Raum mittels MDS.



**Abbildung 3.14:** Links: Entspricht Abb. 3.10. Die ersten beiden Hauptkomponenten werden dargestellt, die Farbe des Punktes entspricht dem zugehörigen Stimulus. Rechts: Die ersten beiden Hauptkomponenten der Matrix  $\mathbf{D}^T$ . An die jeweiligen Punktkoordinaten wurden die Deskriptoren geplottet.

$\mathbf{D}\mathbf{D}^T$  als auch  $\mathbf{D}^T\mathbf{D}$  haben jeweils maximal 16 von Null verschiedene Eigenwerte, denn  $\mathbf{D}$  besteht nur aus 16-dimensionalen Vektoren und  $\mathbf{D}^T$  ist zwar 90-dimensional aber nur durch 16 Punkte beschrieben, womit auch hier nicht mehr als 16 Dimensionen aufgespannt werden können. Die erste Zeile in Abb. 3.13 beschreibt die Eigenwerte auf den Rohdaten, links für die Matrix  $\mathbf{D}$  und rechts für die Matrix  $\mathbf{D}^T$ . Man erkennt den ähnlichen Verlauf der Eigenwerte. Die 90%-Abschätzung auf den PCA-Ergebnissen ergibt für beide eine geschätzte Dimension von 5. In der zweiten Reihe sind die Eigenwert-Entwicklungen für die beiden Matrizen nach der SD-Einbettung dargestellt, auch hier ergibt sich für beide ein vergleichbarer Verlauf und dieselbe Abschätzung von 4 Dimensionen.

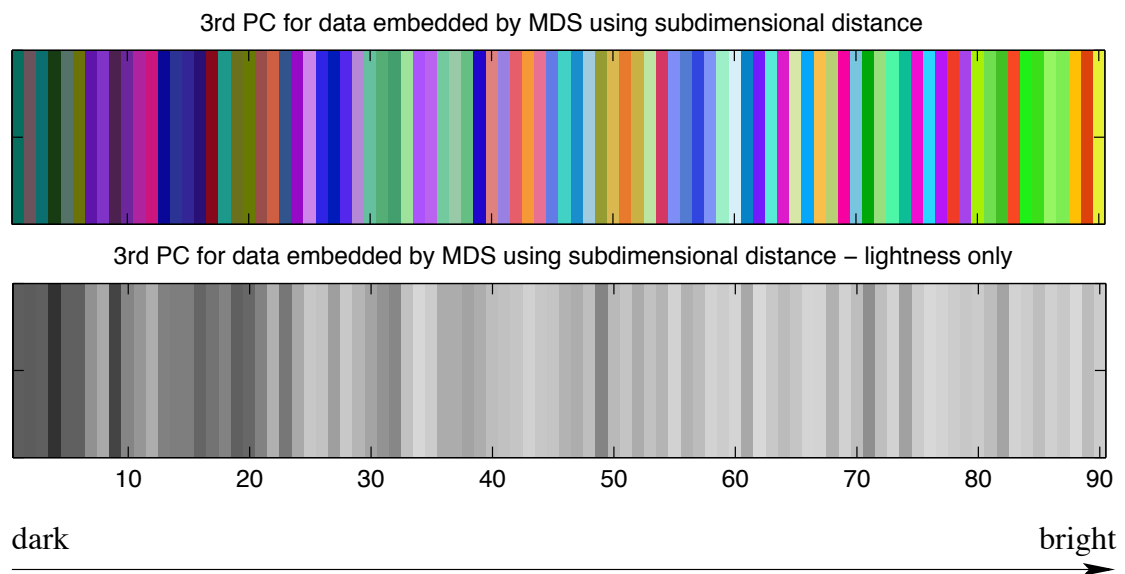
Dass sich die Eigenwert-Verteilungen der PCA so ähnlich sehen, ist wenig verwunderlich, da die Kovarianzmatrizen der beide Matrizen  $\mathbf{D}$  und  $\mathbf{D}^T$  sogar identische Eigenwerte besitzen [55]. Im Rahmen der PCA werden die Daten jedoch vor der Eigenwertberechnung mittelwertbefreit, daher entstehen hier Unterschiede, die Verteilung der Eigenwerte wird dadurch aber nicht grundlegend verändert.

Bleibt die Frage, wie die eingebettete transponierte Matrix aussieht. In Abbildung 3.14 sind die beiden ersten Hauptkomponenten beider Ansätze gegenübergestellt, links die ersten beiden Hauptkomponenten aus  $\mathbf{D}$ , rechts das Ergebnis für  $\mathbf{D}^T$ . Im direkten Vergleich kann man erkennen, dass die rechte Abbildung eine gewisse Ähnlichkeit hat mit den handgesetzten Beschriftungen, die wir in Abb. 3.12 in die Grafik gesetzt haben. Die Darstellung von  $\mathbf{D}^T$  scheint tatsächlich mehr der gesuchten Darstellung des Farbwahrnehmungsraumes zu entsprechen. Die Karte

von **D** dagegen macht nur Sinn, wenn wir bereits in der Lage sind, eine subjektive Interpretation der Ergebnisse vornehmen zu können.

Der Farbkreis ist auch in der neuen Darstellung wiederzuerkennen. *rot*, *grün* und *blau* spannen hier ebenfalls die dominanten Dimensionen auf, *rosig* liegt dicht bei *rot* und *kalt* bei *blau*. Auf dieser objektiven Karte sehen wir nun auch Zusammenhänge, die zuvor nicht so einsichtig waren, z.B. die Lage des Bezeichners *unangenehm* zwischen den Farben *grün* und *gelb*.

Ähnliche Analogien zu den subjektiven Interpretationen von **D** finden wir auch für die dritte Hauptkomponente von **D<sup>T</sup>**. In Abbildung 3.15 findet sich in den ersten Zeilen noch einmal die 3. Hauptkomponente von **D**, darunter die entfärbte Variante und ganz unten sind die 16 Bezeichner in ihrer Ordnung entlang der 3. Hauptkomponente von **D<sup>T</sup>** abgedruckt. Auch diese Achse wird aufgespannt von den Bezeichnern *dunkel* und *braun* bis zu den gegenüberliegenden Bezeichnern *frisch* und *grell*. Dazwischen scheinen nicht alle Bezeichner sinnvoll angeordnet zu sein, es ist in dieser äquidistanten Darstellung aber auch nicht klar zu erkennen, wie dominant die entsprechenden Bezeichner für diese Dimension sind.



3rd PC for alternative presentation with descriptors along the axis

dark brown blue cold yellow warm red rosy clear fresh flashy

**Abbildung 3.15:** Oben: Zum Vergleich die dritte Hauptkomponente aus der Analyse von **D**. Die Farben verlaufen von *dunkel* nach *hell* wie an der desaturierten Darstellung in der zweiten Zeile zu erkennen ist. Unten: Eine Auswahl der 16 Bezeichner sind nach ihrer Ordnung entlang der 3. PC von **D<sup>T</sup>** angeordnet. Auch hier finden sich *dunkel* und *grell* an den Maxima.

### 3.3 Diskussion

In diesem Kapitel haben wir am Beispiel des Farbensehens eine Analysemethode vorgestellt, die es ermöglicht, *objektive* Wahrnehmungskarten zu erstellen anstatt in reinen Merkmalsräumen nachträglich subjektive Interpretationen zur Deutung der Zusammenhänge einführen zu müssen.

Wir haben uns dem Problem auch auf der experimentellen Seite genähert. Zwar hätte schon alleine die Auswahl der Farbbezeichner ein eigenes Experiment gerechtfertigt, ebenso wie wir den Test schuldig geblieben sind, ob die Einträge der transponierten Matrix experimentell gleichartig hätten erzeugt werden können, d.h. ob die Probanden den Bezeichnern ebenso die Farben zuordnen wie umgekehrt. Insgesamt denken wir aber, dass vor allem der positive Ausgang des Experiments unsere Annahmen stützt.

Wir haben zeigen können, dass es alleine auf der Basis der von uns durch Befragung von Personen erhobenen Daten möglich war – ohne fundierte Kenntnisse der sensorischen Zusammenhänge – eine Darstellung des Merkmalsraumes zu erhalten. Diese Darstellung besitzt erstaunliche Ähnlichkeit mit dem wohlbekannten Farbkreis.

Es wurden nicht *exakt* drei Dimensionen rekonstruiert, aber mit den ermittelten vier Dimensionen haben wir uns dem schon erfreulich angenähert. Vor allem aber konnten wir zeigen, dass die subjektiven Ergebnisse aus der Analyse des Merkmalsraumes durch unsere objektive Untersuchung des Wahrnehmungsraumes reproduziert werden können.

Es ist nun nur noch ein kleiner Schritt, um endgültig Wahrnehmungskarten zu erstellen, auf denen auch die Stimuli wieder einen Platz finden. Es wäre zum Beispiel sehr leicht möglich, den einzelnen Bezeichnern die Stimuli zuzuordnen, die der entsprechenden Wahrnehmung am ähnlichsten sind bzw. die Stimuli bei der Darstellung zu berücksichtigen, die am eindeutigsten einen Bezeichner charakterisieren. Man könnte also zum Beispiel auf der Farbwahrnehmungskarte den Bezeichner *rötlich* in der typischen Stimulifarbe darstellen, die *rötlich* zugeordnet wurde.

Wie gut dies funktioniert, ist aber natürlich nicht nur von den Bezeichnern sondern auch von den Stimuli abhängig. Können wir hier für die Farben bereits konkrete Vorschläge machen, stellt uns die Einbettung auf der Wahrnehmungskarte bereits im nächsten Kapitel bei der Beschreibung von Geruchswahrnehmung vor eine ganz andere Herausforderung.

# 4 Der Geruchswahrnehmungsraum

## 4.1 Einleitung

Nachdem wir im vorangegangenen Kapitel am Beispiel des Farbensehens einen allgemeingültigen Analyserahmen vorgestellt haben, wollen wir diese Methoden auch auf Daten des menschlichen Geruchssinns anwenden. Es handelt sich hierbei um eine Wahrnehmung, zu der wesentlich weniger Referenzen existieren, um potentielle Ergebnisse dagegen abzugleichen.

Wie wir es in Kapitel 3 schon angedeutet haben, steht uns für die Geruchswahrnehmung eine ähnliche Datenbank zur Verfügung, in der Informationen darüber stehen, wie Personen den Geruch bestimmter Substanzen mit Hilfe von vorgegebenen Worten charakterisiert haben. Diese Datenbank wurde uns im Rahmen einer Kooperation mit dem California Institute of Technology zur Verfügung gestellt.

Das Ziel dieser Analyse wird es sein, diese Geruchsdatenbank in den kleinstmöglichen euklidischen Raum einzubetten, um aus dieser Einbettung etwas über die Struktur des Geruchswahrnehmungsraumes zu lernen. Aus einer objektiven Darstellung der Zusammenhänge innerhalb dieses Wahrnehmungsraumes erhofft man sich Ideen und Inspirationen, um gezielt praktische Geruchsexperimente zu entwerfen, die letztlich helfen können, um Geruchswahrnehmung – wie beim Farbensehen mittels der Wellenlänge – endlich auch quantifizieren zu können.

Beim Farbensehen hat man bereits eine konkrete Vorstellung von der Komplexität des zugehörigen Wahrnehmungsraumes: Man geht davon aus, dass drei bis vier Dimensionen ausreichend sind, um den Farbwahrnehmungsraum aufzuspannen (siehe Abschnitt 3.1.3 und [53]). Es ist möglich, nahezu alle Farbtöne, die der Mensch unterscheiden kann, durch eine Linearkombination von drei Grundfarben in einem euklidischen Raum zusammenzusetzen. In diesem Raum können auch Wahrnehmungsähnlichkeiten zwischen zwei Farben leicht abgelesen werden.

Beim Riechen hingegen würden fundierte Kenntnisse über einen entsprechenden Wahrnehmungsraum ganz neue Möglichkeiten eröffnen. Derzeit ist eine fundierte

Analyse und eine Abschätzung von Geruchsähnlichkeiten kaum möglich. Nehmen wir zum Beispiel die Geruchsbeschreiber *apple*, *cherry* und *banana*. So banal es klingt, aber es ist noch nicht einmal verstanden, ob die drei Geruchsbezeichner das gleiche Geruchsmerkmal in unterschiedlichen Ausprägungen beschreiben oder ob es sich hierbei schlicht um Beschreibungen von ganz unterschiedlichen Wahrnehmungsmerkmalen handelt.

### 4.1.1 Grundlagen des Riechens

Alles was einen eigenen Geruch besitzt, gibt beständig winzige Mengen von Geruchsstoffen (sog. *Duftstoffe* oder auch *Odorants*) an seine Umwelt ab. Daher ist die Luft erfüllt von einer Mischung aus beliebig vielen solcher Duftstoffe, ganz egal ob es sich dabei um den Geruch einer wundervollen Rose oder um einen vergammelten Fisch handelt. Diese Moleküle sind so winzig, dass sie nahezu nicht nachweisbar sind. Typischerweise handelt es sich um hochreaktive chemische Substanzen. Ein Sensor, der in der Lage ist, einen solchen Stoff zu detektieren, wird auch ein „chemischer Sensor“ genannt. Entsprechend handelt es sich bei der Nase um ein „chemisches“ Sinnesorgan.

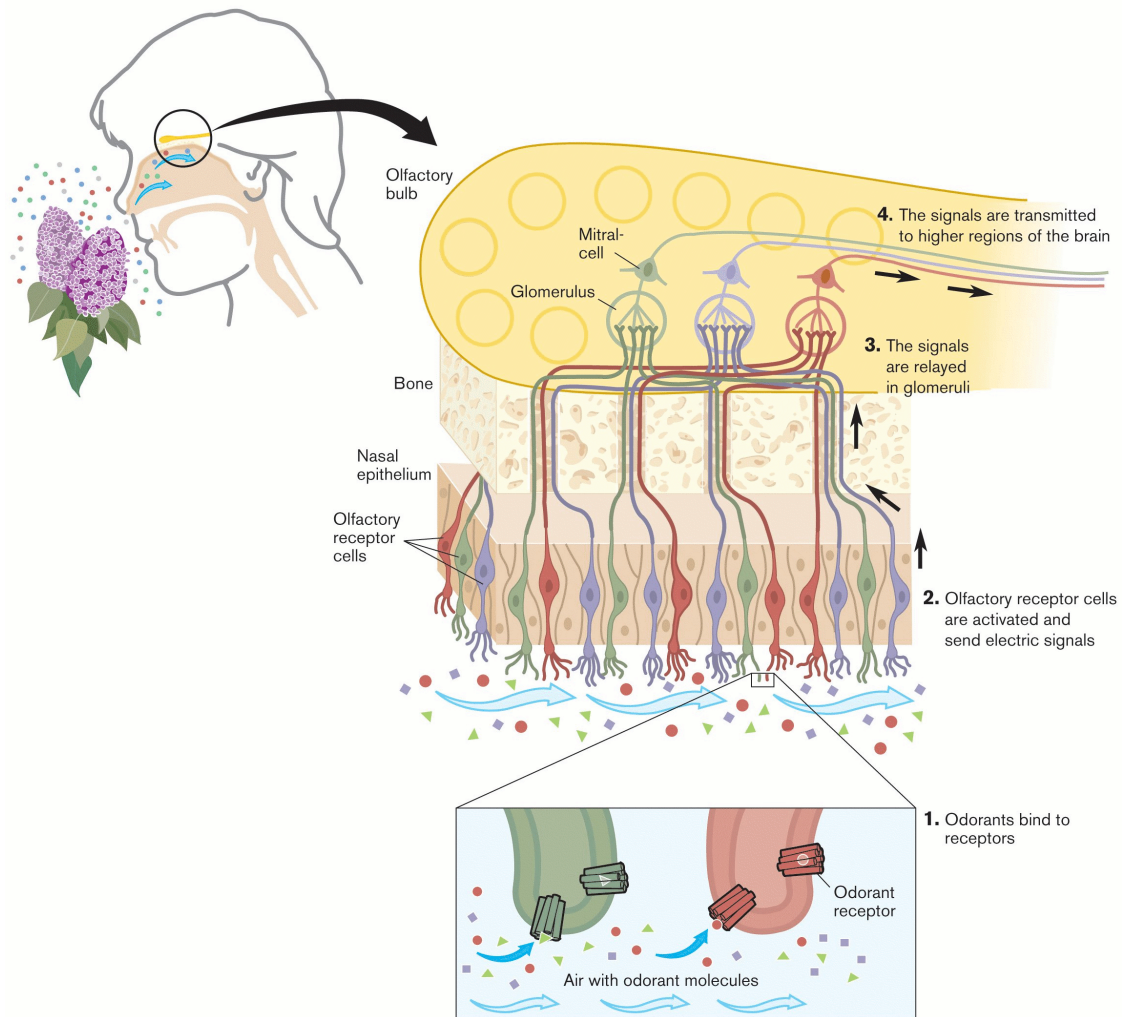
Die meisten Menschen sind sich in ihrem Alltag nicht mehr der aktiven Rolle ihres Geruchssinns bewusst. Dabei handelt es sich beim Riechen für die meisten Säugetiere noch immer um den primären Sinn. Sie benutzen ihren Geruchssinn bei so essentiellen Aufgaben wie der Auswahl ihrer Nahrung, ihrer Sexualpartner oder auch dem Wittern von Gefahren. Duftstoffe sind in der Lage, ganz wesentlich unser Befinden zu beeinflussen und können sowohl Unbehagen, Sympathie wie auch Ablehnung auslösen.

Auch wenn der moderne Mensch sehr viel Energie darauf verwendet, natürliche Gerüche, besonders Körpergerüche, aus seinem Alltag zu verbannen, besitzt dennoch jeder Mensch einen einzigartigen, genetisch bestimmten Eigengeruch. Daher ist es nicht unwahrscheinlich, dass auch für den Menschen einst die Nase eine zentrale Rolle gespielt hat und es vielleicht immer noch tut. Wells und Hepper [97] haben in einer Studie die Kapazität des menschlichen Geruchssinns auf eine harte Probe gestellt: Sie haben Hundebesitzer darauf getestet, ob sie in der Lage seien, ihren eigenen Hund anhand seines Geruches von anderen Hunden zu unterscheiden. Interessanterweise konnten 88.5% der Teilnehmer erfolgreich ihren Hund an seinem Geruch erkennen.

Säugetiere sind in der Lage, eine enorme Vielzahl von Duftstoffen zu Unterscheiden. Auch Menschen können etwa 10.000 verschiedene Duftstoffe differenzieren [4]. Eine Geruchswahrnehmung (ein sog. *Geruch* oder auch *Odor*) wie etwa *blumig*, kann schon ausgelöst werden durch winzigste Konzentrationen eines Duftstoffes. Einige Moleküle können bereits bei einer Konzentration von einem Duftmolekül pro einer



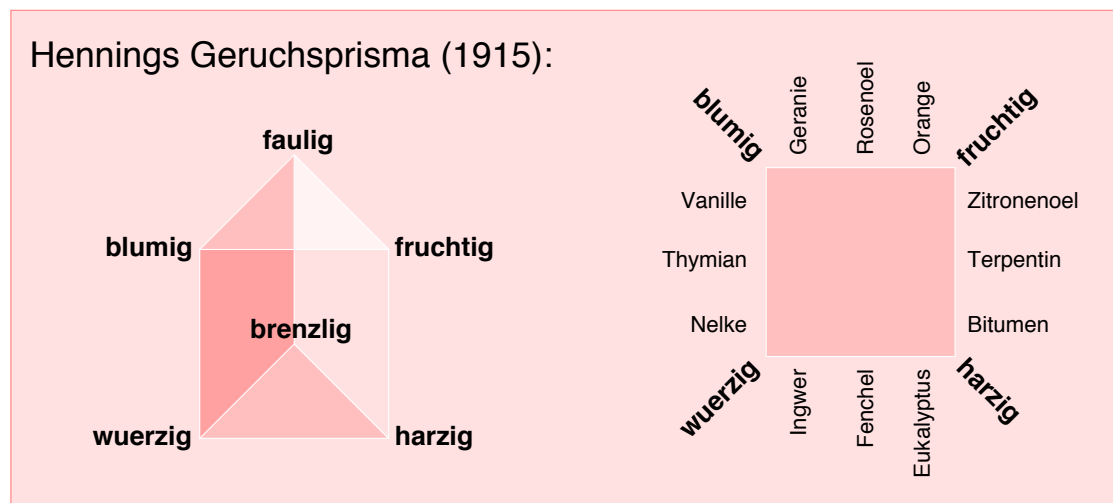
## Odorant Receptors and the Organization of the Olfactory System



**Abbildung 4.1: Schematische Darstellung der menschlichen Nase.** Eingeatmete Duftstoffe (Odorants) binden an Neurone in der Riechschleimhaut (Olfactory Epithelium). Die Riechschleimhaut befindet sich an der oberen Wand der Nasenhöhle. Riechneurone (ORNs) senden ihre Axone durch das Siebbein in den Riechkolben (OB), wo gleichartige ORNs in sog. Glomeruli zusammenlaufen. Von dort propagieren Neurone (Mitralzellen) die Glomeruli-Aktivität in höhere Hirnregionen wie die Riechrinde (Pyriform Cortex). *Picture taken from press release at [www.nobelprize.org](http://www.nobelprize.org).*

Billion Luftteilchen wahrgenommen werden.

Es befinden sich ca. 1.000 verschiedene Geruchsrezeptoren in der menschlichen Nase [4, 12]. Sie unterscheiden sich in einem genetisch kodierten Rezeptor-Protein, welches sich an der Spitze der Neurone in der Nasenschleimhaut ausbildet und für die spezifische Molekülbindung eines Rezeptors verantwortlich ist. Es handelt sich



**Abbildung 4.2:** Henning schlug 1915 ein Geruchsmodell in Form eines dreieckigen Prisma vor. Die Basisgerüche sind in den Ecken platziert: *Blumig*, *würzig*, *fruchtig*, *harzig*, *faulig* und *brenzlig*. Die übrigen Gerüche lassen sich als Mischung der anderen darstellen und haben daher alle Koordinaten innerhalb oder auf der Oberfläche des Prismas. **Rechts:** Das Prisma in einer schematischen Darstellung. **Links:** Beispielhaft sind Gerüche aufgetragen, die (nach Henning) auf den Verbindungslinien der Prismenbasis zwischen vier der Basisgerüchen liegen.

hierbei um eine bemerkenswert hohe Anzahl unterschiedlicher Rezeptoren, da es auch genauso viele Gene geben muss, die diese Rezeptorproteine kodieren. Somit sind etwa 1 – 2% der insgesamt etwa 50.000 bis 100.000 menschlichen Gene dem Geruchssinn zuzuordnen. Dies mag ein weiterer Hinweis auf die hohe Bedeutung von chemischen Sensoren in der Evolution des Menschen sein.

Auch wenn die sensorischen Grundlagen des Riechens erst für die folgenden Kapitel relevant werden, so soll hier doch zu Beginn des Kapitels die grundlegende Funktionsweise des Riechens – soweit bekannt – knapp dargestellt werden.

### Duftstoff-Bindung in der Riechschleimhaut

Die Duftstoffe sind biochemische Liganden für spezifische Rezeptorproteine (Odorant Receptor Proteins (ORPs)). Die sogenannten Riechzellen (Odorant Receptor Neurons (ORNs)) exprimieren diese ORPs an ihrer Zelloberfläche und befinden sich in in dichter Anordnung in der Riechschleimhaut (Olfactory Epithelium (OE)). Das Epithelium befindet sich an der oberen Wand der Nasenhöhle und hat eine Größe von etwa  $6.55\text{cm}^2$  [80]. Dies ist auch in Abbildung 4.1 noch einmal in schematischer Darstellung zu sehen.

---

Man konnte zeigen, dass jedes ORN auf seiner Rezeptoroberfläche genau einen Typ von ORPs exprimiert [62]. Die unterschiedlichen Typen von ORNs verteilen sich auf vier unterschiedliche Regionen des OE. Innerhalb dieser Regionen sind die ORNs im OE zufällig verteilt [13]. In-situ-Hybridisierungen konnten nachweisen, dass alle ORNs, die dasselbe ORP auf ihrer Oberfläche tragen, mit ihren Axonen zu denselben Glomeruli im Riechkolben (Olfactory Bulb, OB) konvergieren [4]. Dies ist in Abbildung 4.1 ebenfalls dargestellt. Ein Glomerulus ist gleichzeitig ein Verschaltungspunkt zu den sogenannten Mitralzellen, welche die eingehenden Signale aus dem Bulbus hinaus in die höheren Verarbeitungsschichten des Gehirns wie z.B. die Riechrinde (Pyriform Cortex) propagieren.

Selbst angesichts der 1.000 verschiedenen Rezeptortypen ist die Fähigkeit, mehr als 10.000 Duftstoffe unterscheiden zu können, noch nicht hinreichend geklärt. Da Experimente nicht nur darauf hindeuten, dass jeder Rezeptortyp eindeutig einem Glomerulus zugeordnet werden kann [4] sondern auch jedes ORN genau einen einzigen ORP-Typ exprimiert [62], geht man davon aus, dass ein kombinatorischer Kode existiert, der substanzspezifische Aktivierungsmuster der Glomeruli als Basis für eine nachgeschaltete Riechwahrnehmung verwendet [13].

Diese Art der Kodierung wird das zentrale Thema in den folgenden Kapiteln 5 und 6 sein. In diesem Kapitel hingegen werden wir uns in Fortführung des letzten Kapitels, nämlich mit der Analyse von Geruchsdatenbanken, auseinandersetzen. Denn obwohl es immer wieder Versuche gab, sich diesem Wahrnehmungsraum systematisch zu nähern, ist nach wie vor wenig über Geruchswahrnehmungen verstanden. Noch immer sind sowohl die Parfumindustrie wie auch Geruchsforscher darauf angewiesen, als Referenz auf große Sammlungen von Geruchsstoffen und ihrer Gerüche wie dem *Dravnieks Atlas of Odor Character Profiles* [19] oder sogar einfach nur auf die persönlichen Erfahrungen von Experten zurückzugreifen.

## 4.1.2 Geruchskarten

Bereits in der Antike wurde nicht nur über das Farbensehen philosophiert sondern auch über die mögliche Funktion des Geruchssinns. Aristoteles (384-322 v.Chr.) versuchte z.B. Geruchswahrnehmungen mit den gleichen Beschreibern zu erklären, die er auch verwendet hatte, um den Geschmackssinn zu erklären. Einzig fügte er eine Geruchsqualität hinzu, die er *stinkend* nannte. Aber auch er musste bereits eingestehen, dass sich der Geschmackssinn wesentlich leichter kategorisieren lässt als der Geruchssinn [53].

Wesentlich später, im 18. und 19. Jahrhundert haben Wissenschaftler versucht, die unterschiedlichen Geruchsqualitäten in verschiedene Klassen einzuordnen, ebenso wie auch Tiere und Pflanzen kategorisiert wurden. Carl von Linné (1707–1778) hat in einer Arbeit von 1752 Gerüche in sieben Klassen eingeteilt. Eine 1895 auf neun

Klassen verfeinerte Version von Hendrik Zwaardemaker (1857–1930) blieb bis in das 20. Jahrhundert als Kategorisierung von Düften weitgehend akzeptiert:

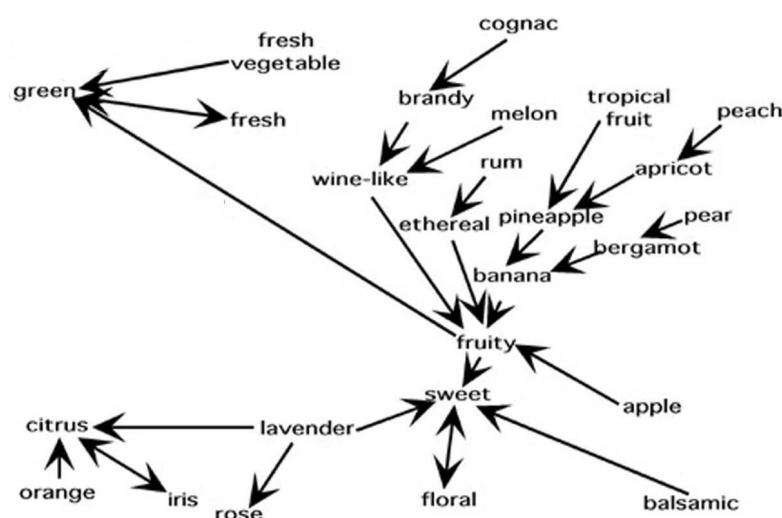
1. aromatische Gerüche (z.B. Anis)
2. ätherische Gerüche (z.B. Apfel)
3. balsamische Gerüche (z.B. Jasmin)
4. Moschusartige Gerüche (z.B. Patchouli)
5. abstoßende Gerüche (z.B. Verwesung)
6. Lauchartige Gerüche (z.B. Zwiebel)
7. Rauchartige Gerüche (z.B. brennendes Holz)
8. Käseartige Gerüche (z.B. Limburger Käse)
9. betäubende Gerüche (z.B. Opium)

Diese frühen Klassifikationen basierten in erster Linie auf persönlicher Erfahrung und kaum auf experimentellen Beobachtungen oder Erhebungen [14]. Blickt man in die Geschichte der Geruchswahrnehmung zurück, so gibt es nur wenige Arbeiten, in denen explizit versucht wird, die Dimensionalität der Geruchswahrnehmung abzuschätzen.

Eine der ersten Arbeiten, in der dieses explizit versucht wird, geht auf Hans Henning (1885–1946) zurück, der 1915 ein sogenanntes „Geruchsprisma“ vorgestellt hat [27]. In Abbildung 4.2 ist sein Prisma aufgezeigt. Er verwendete darin sechs Basisgeruchsqualitäten (*faulig*, *blumig*, *würzig*, *harzig*, *brenzlich* und *fruchtig*), um mit diesen in Form eines dreidimensionalen Prismas den Geruchsraum aufzuspannen. Jeder Duftstoff hatte in seinem Modell einen festdefinierten Platz, z.B. liegt der Duft von Thymian irgendwo zwischen *blumig* und *würzig*. Hennings Theorie konnte sich jedoch nie wirklich durchsetzen.

Erst viele Jahre später (1968) stellte Woskow eine frühe Version des sogenannten *Multidimensional Scaling* vor, mit dem er psychophysikalische Daten in einen euklidischen Kontext bringen konnte. Er stellte die Hypothese auf, dass der Wahrnehmungsraum tatsächlich dreidimensional sei, allerdings auf der Basis von Ähnlichkeitsbeschreibungen von 25 Substanzen gegeneinander. Nach seinen eigenen Angaben waren die Substanzen zufällig ausgewählt und teilweise so ähnlich, dass bereits während des Experimentes nur anhand des Siedepunktes überprüft werden konnte, ob die Flaschen auch richtig beschriftet waren [100].

Darauf aufbauend veröffentlichte 1974 Susan Schiffman [81] eine Arbeit, in der sie gar schließt, der Geruchsraum lasse sich in einem nur zweidimensionalen Raum einbetten. Ihre Abschätzungen basieren auf zwei Datensätzen, die beide initial schon nur einen Raum zulassen der eher niedrigdimensional ist. Der erste bestand aus 50 Duftstoffen, die in ihrer Ähnlichkeit gegen 9 Basisduftstoffe verglichen wurden



**Abbildung 4.3: Ausschnitt aus dem Geruchs-Graph von Chee-Ruiter.** Der gerichtete Graph besteht aus Verbindungen zwischen einem Geruchsbeschreiber A und seines nächsten Nachbarn B. Der komplette Graph ist in Abbildung B.1 gezeigt.

(d.h. dieser Ansatz konnte höchstens einen 9-dimensionalen Raum als Geruchsraum ergeben). Der zweite Datensatz war der oben bereits erwähnte Datensatz von Woskow [100]. Interessanterweise ergab eine erste Faktorenanalyse für beide Datensätze 8 Dimensionen.

Zusätzlich zu diesen psychophysikalischen Karten gibt es einige empirische Ansätze, die sich weit verbreitet haben, vor allem in der Parfumindustrie. In allen Fällen werden zwei- bis dreidimensionale Räume vorgeschlagen.

In den letzten Jahrzehnten setzte ein so enormer Wissenszuwachs ein, dass es nur noch als eine Frage der Zeit schien, bis die ersten fundierten Geruchskarten entstehen würden. Aber noch immer stellt die Quantifizierung von Geruchsqualitäten ein ungelöstes Problem dar.

Christine Chee-Ruiter hatte dann genau die Idee, die wir im vorangegangenen Kapitel andiskutiert haben: Alle existierenden Ansätze haben tatsächlich Karten des *Merkmalsraumes* erstellt und sind dann an der Interpretation gescheitert. Sind die Stimuli selbst – hier also chemische Moleküle – die Punkte in dem untersuchten Raum, so kann die Interpretation nur entlang der bereits vorhandenen Kenntnisse entstehen. Diese sind aber beschränkt durch den Wissenstand über das untersuchte Kontinuum. Mit anderen Worten: Es ist schwer, Dinge zu erkennen, von denen man gar nicht weiß, dass man sie sucht.

Im Jahre 2000 schlug Chee-Ruiter nun vor, Informationen über die Beschreiber aus den Qualitätsprofilen der Moleküle abzuleiten. Sie hat damit implizit begonnen, ei-

ne Karte des Wahrnehmungsraumes zu erstellen, in dem eben genau die Bezeichner die Punkte sind und ihre Anordnung durch die Stimuli bestimmt wird [15].

Ihre Analyse umfasste beinahe 900 Substanzen, beschrieben durch 278 Beschreiber. Die Karten bestanden aus gerichteten Graphen, was eine Interpretation schwierig macht. Tatsächlich ist alleine die planare Darstellung eines Graphen mit 278 Knoten eine schwierige Aufgabe, da es beliebig viele Möglichkeiten gibt, die Knoten anzuordnen. In Abbildung 4.3 ist ein Ausschnitt aus Chee-Ruiter's Graph zu sehen. Jeder Beschreiber ist mit dem ähnlichsten Nachbarn verbunden. Da das von ihr verwendete Maß nicht symmetrisch ist, gilt nicht immer die Umkehrung. Der komplette Graph ist im Anhang in Abbildung B.1 zu sehen.

Wir werden im folgenden eine Erweiterung dieses Ansatzes entlang unseres Analyserahmens aus Kapitel 3 vorstellen. Die Datenbank von Chee-Ruiter soll mittels Multidimensional Scaling [51] in den kleinsten euklidischen Raum eingebettet werden. Anders als bei den Farben – und entgegen der bisherigen Arbeiten – werden wir jedoch feststellen, dass der Geruchsraum komplex und hochdimensional ist. Um die Ergebnisse jedoch lesbar und vergleichbar zu machen, werden wir die eingebetteten Daten mittels einer selbstorganisierenden Karte [77] darstellen, um letztlich doch noch eine zweidimensionale Näherung zu erhalten.

Da die MDS-Analyse einen hochdimensionalen Geruchswahrnehmungsraum erwarten lässt, stellt sich als nächstes natürlich die Frage, *welche* Merkmale es denn explizit sind, die diesen Raum aufspannen. Daher scheint es zunächst ein essentieller Schritt zu sein, die Anzahl zu erwartender Merkmale möglichst präzise abzuschätzen.

### 4.1.3 Geruchsbeschreiber

Es gibt einige Datensätze, in denen Geruchsqualität beschrieben wird. Üblicherweise werden die Beschreibungen von Experten oder einer kleinen Gruppe von Probanden erstellt, die ihren Riecheindruck mittels einer vorgegebenen Liste von Worten kategorisieren. Dies ist vergleichbar mit einer Phantombild-Erstellung bei der Polizei und entspricht im Wesentlichen der Datenerfassung, wie wir sie auch schon für die Farben in Kapitel 3 beschrieben haben.

Wenn durch ein „**X**“ die Assoziation eines Beschreibers zu einer Substanz gegeben ist, ergibt sich eine Datenmatrix der Form:

| Duftstoff                                        | fruchtig | ananas   | süßlich  | apfel    | kokos    | nussig   |
|--------------------------------------------------|----------|----------|----------|----------|----------|----------|
| <b>c</b> <sub>1</sub> : Hexyl Butyrat            | <b>X</b> | <b>X</b> | <b>X</b> |          |          |          |
| <b>c</b> <sub>2</sub> : Methyl-2-Methylbutyrat   | <b>X</b> |          | <b>X</b> | <b>X</b> |          |          |
| <b>c</b> <sub>3</sub> : 6-Amyl- $\alpha$ -Pyrone |          |          | <b>X</b> |          | <b>X</b> | <b>X</b> |

---

In diesem Beispiel riecht Substanz  $\mathbf{c}_3$  also *süßlich* aber nicht *fruchtig*. Dasselbe, ein wenig mathematischer ausgedrückt, ergibt:

$$\mathbf{C} = \begin{pmatrix} \mathbf{c}_1 \\ \mathbf{c}_2 \\ \mathbf{c}_3 \end{pmatrix} = \overbrace{\begin{bmatrix} 1 & 1 & 1 & 0 & 0 & 0 \\ 1 & 0 & 1 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 & 1 & 1 \end{bmatrix}}^{\text{odor descriptors } \mathbf{o}_1, \dots, \mathbf{o}_6}$$

Hier enthält jede Zeile  $i$  das Geruchsprofil (oder auch den Merkmalsvektor) der Substanz  $\mathbf{c}_i$ . Jede Spalte  $j$  hingegen speichert die Information, welche Substanzen  $\mathbf{c}_i, i = 1, \dots, 3$  die Geruchswahrnehmung  $\mathbf{o}_j$  auslösen können oder nicht. Analog zum Vorgehen im vorangehenden Kapitel werden wir durch einfaches Transponieren der Matrix  $\mathbf{C}$  eine neue Matrix  $\mathbf{O}$  erzeugen:

$$\mathbf{C}^T = \mathbf{O} = \begin{pmatrix} \mathbf{o}_1 \\ \mathbf{o}_2 \\ \mathbf{o}_3 \\ \mathbf{o}_4 \\ \mathbf{o}_5 \\ \mathbf{o}_6 \end{pmatrix} = \begin{bmatrix} 1 & 1 & 0 \\ 0 & 0 & 1 \\ 1 & 1 & 1 \\ 0 & 1 & 0 \\ 1 & 0 & 0 \\ 1 & 0 & 0 \end{bmatrix}$$

Somit steht nun in jeder Zeile  $i$  die Information über eine Geruchsqualität  $\mathbf{o}_i$ . Dies entspricht dem Vorschlag von Chee-Ruiter, um aus diesen Daten etwas über die Zusammenhänge im Wahrnehmungsraum herauszufinden.

Wir verwenden im weiteren einen Beschreiber-Datensatz, der weitestgehend auf dem *Aldrich Flavor and Fragrances Catalog* [1] basiert. Diese Sammlung führt zwei Datensätze (Den Arctander [3] und die Fenaroli-Daten [22]) und besteht im wesentlichen aus 851 Substanzen, deren Geruchsqualität durch insgesamt 278 Geruchsbeschreiber charakterisiert wird. Dieser Datensatz ist derselbe, den auch schon Chee-Ruiter zur Erstellung ihrer Geruchs-Graphen (wie im letzten Abschnitt beschrieben) verwendet hat [15].

Wir haben den Beschreibersatz von 278 auf 171 Beschreiber reduziert, indem alle Beschreiber aus der Datenbank gelöscht wurden, die nur genau eine Substanz charakterisieren. Die Ähnlichkeiten der verbleibenden 171 Beschreiber wurden dann mit Hilfe der SD-Ähnlichkeit (siehe Gleichung 3.1 in Kapitel 3) abgeschätzt.

Dies ist die Ähnlichkeitsmatrix, die wir im weiteren verwenden werden, um die Dimensionalität des Geruchswahrnehmungsraumes abzuschätzen. An dieser Stelle möchten wir noch einmal erwähnen, dass wir hier jetzt keine experimentell abgeleiteten Ähnlichkeiten untersuchen, sondern vielmehr alleine auf der Basis von Geruchsbeschreibungen den zugehörigen Wahrnehmungsraum analysieren werden.

## 4.2 Datenanalyse des Geruchswahrnehmungsraumes

Gesucht ist also eine metrische Darstellung der Zusammenhänge zwischen den einzelnen Wahrnehmungsbeschreibern in einem möglichst überschaubaren Raum, d.h. möglichst eine zweidimensionale, euklidische Geruchskarte. Am wichtigsten ist, dass sie auch alle Relationen zwischen den Beschreibern konserviert, d.h. benachbarte Beschreiber im Wahrnehmungsraum sollten auch auf der Karte beieinander liegen und sehr unterschiedliche entsprechend voneinander getrennt.

Grundsätzlich ist in einem solchen Fall die Untersuchung der intrinsischen Dimension des gegebenen Datenraums ein erster Schritt, um die Struktur der Daten einschätzen zu lernen. Wie wir es im Kapitel 3 bereits vorgeschlagen haben, ist die Hauptkomponentenanalyse (PCA) das Mittel der Wahl, um mittels der Eigenwerte die intrinsische Dimension abzuschätzen. Die Eigenwerte geben die Varianzen wieder, und eine Dimension mit Varianz Null deutet auf eine irrelevante Dimension hin [44].

Da wir die binären Geruchsdaten ebenfalls mit der SD-Ähnlichkeit abschätzen, verwenden wir auch hier wieder das Multidimensional Scaling (MDS, [51]) zum Einbetten der Ähnlichkeitsmatrix  $D_{\text{subsim}} \in \mathbb{R}^{171 \times 171}$  in einen euklidischen Raum (siehe auch Absatz 3.2.3 in Kapitel 3).

MDS versucht,  $n$  euklidische Punkte  $\mathbf{m}_i \in \mathbb{R}^d$  so in einem  $d$ -dimensionalen Raum zu verschieben, dass der Einbettungsfehler zwischen der euklidischen Distanzmatrix  $D_{\mathbf{M}}$  dieser Punkte minimal wird verglichen mit der Ähnlichkeitsmatrix  $D_{\text{subsim}}$ .

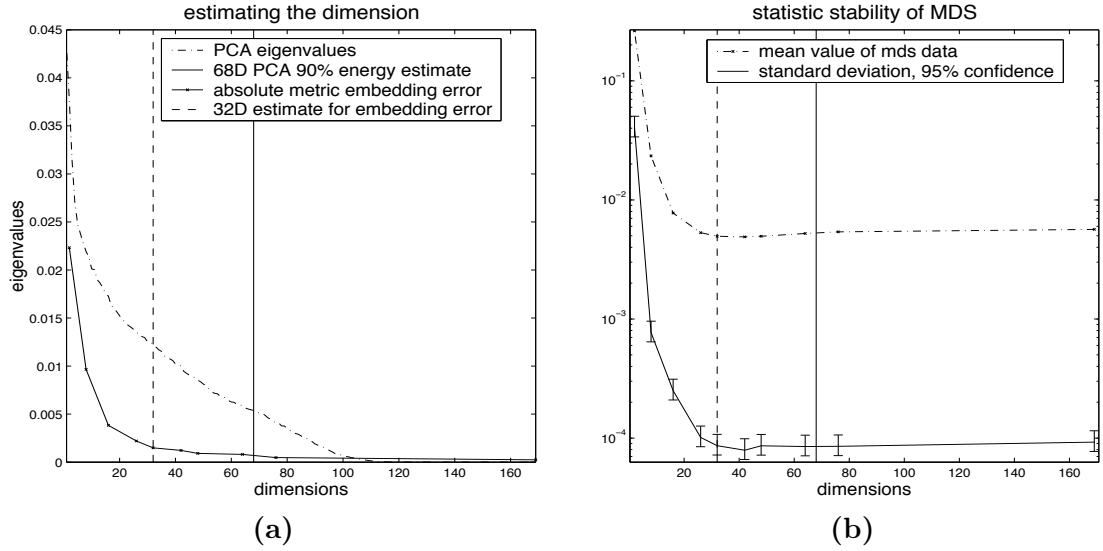
Für  $n$  Beschreiber ergibt sich für  $d = n - 1$  Dimensionen eine triviale Einbettung, denn  $n$  Punkte können höchstens einen  $(n - 1)$ -dimensionalen Raum aufspannen. Somit wäre  $D_{\text{subsim}}$  mit  $D_{\mathbf{M}}$  perfekt zu rekonstruieren, sofern die Ähnlichkeiten mit einem euklidischen Maß geschätzt wären.

Sei nun also  $d_{\text{sim}}(i, j)$  die Ähnlichkeit zwischen zwei Beschreibern  $i$  und  $j$ , und sei  $d_{\text{eucl}}^{(n-1)}(i, j)$  die euklidische Distanz zwischen den beiden Beschreibern nach der Einbettung aller  $n$  Punkte in einen  $(n - 1)$ -dimensionalen euklidischen Raum, dann ergibt sich der *nicht-metrische* Einzeldefekt  $\Delta$  durch:

$$\Delta(i, j) = |d_{\text{sim}}(i, j) - d_{\text{eucl}}^{(n-1)}(i, j)|$$

Dieser Defekt bleibt konstant für alle höheren Dimensionen mit  $d > (n - 1)$ . Es wäre jedoch auch möglich, dass der Defekt auch für kleinere Dimensionen mit  $d < (n - 1)$  konstant bleibt. Dann nimmt man an, dass die Daten auf einem intrinsischen Unterraum liegen. Die zugehörige intrinsische Dimension kann abgeschätzt werden





**Abbildung 4.4:** (a) Die strichpunktierte Kurve zeigt die geordneten Eigenwerte für die triviale Einbettung in 170D für die 171 Beschreiber. Die vertikale Linie zeigt die 90%-Eigensumme (bei 68D). Der metrische Einbettungsfehler  $\varepsilon_{\text{metric}}^D$  (durchgezogene Linie) nimmt signifikant zu für Dimensionen kleiner als 32D (gepunktete Senkrechte). Der Fehler ist in den Plot hineinskaliert. (b) Standardabweichung und Mittelwerte des tatsächlichen Einzelfehlers  $\varepsilon^D(i, j)$ . Je stärker die Werte streuen, desto größer wird auch der Lösungsraum. Hier wird die Schranke von 32D noch deutlicher. Die gemittelten Fehler in beiden Plots beziehen sich auf 100 MDS-Läufe pro getesteter Dimension.

durch die niedrigste Dimension, in die die  $n$  Punkte eingebettet werden können, ohne dass der euklidische Defekt signifikant erhöht wird.

Der tatsächliche Einzelfehler  $\varepsilon^D(i, j)$  für die Einbettung in  $D$  Dimensionen beträgt

$$\varepsilon^D(i, j) = |d_{\text{sim}}(i, j) - d_{\text{eucl}}^D(i, j)|$$

Mit dem nicht-metrischen Einzeldefekt erhalten wir direkt eine Abschätzung des metrischen Einzelfehlers  $\varepsilon_{\text{metric}}^D(i, j)$  mit

$$\varepsilon_{\text{metric}}^D(i, j) = |\varepsilon^D(i, j) - \Delta(i, j)|$$

und folglich auch den metrischen Einbettungsfehler  $\varepsilon_{\text{metric}}^D$  mit

$$\varepsilon_{\text{metric}}^D = \sum_{i, j} |\varepsilon^D(i, j) - \Delta(i, j)| \quad (4.1)$$

Es sei angemerkt, dass die Qualität der Einbettung direkt vom *metrischen* Einbettungsfehler abhängt, der *nicht-metrische* Defekt hingegen wird konstant bleiben.

### 4.2.1 Dimensionen des Riechraumes

Um die Daten initial zu untersuchen, haben wir die Rohdaten – also die  $(171 \times 171)$ -Matrix mit den SD-Ähnlichkeiten – mittels MDS in den  $\mathbb{R}^{170}$  eingebettet. Dies sollte uns einen minimalen metrischen Einbettungsfehler geben, und somit eine gute Annäherung an den nicht-metrischen Defekt. Auf diesen eingebetteten Punkten haben wir dann mittels einer PCA die Eigenwertentwicklung der Daten untersucht.

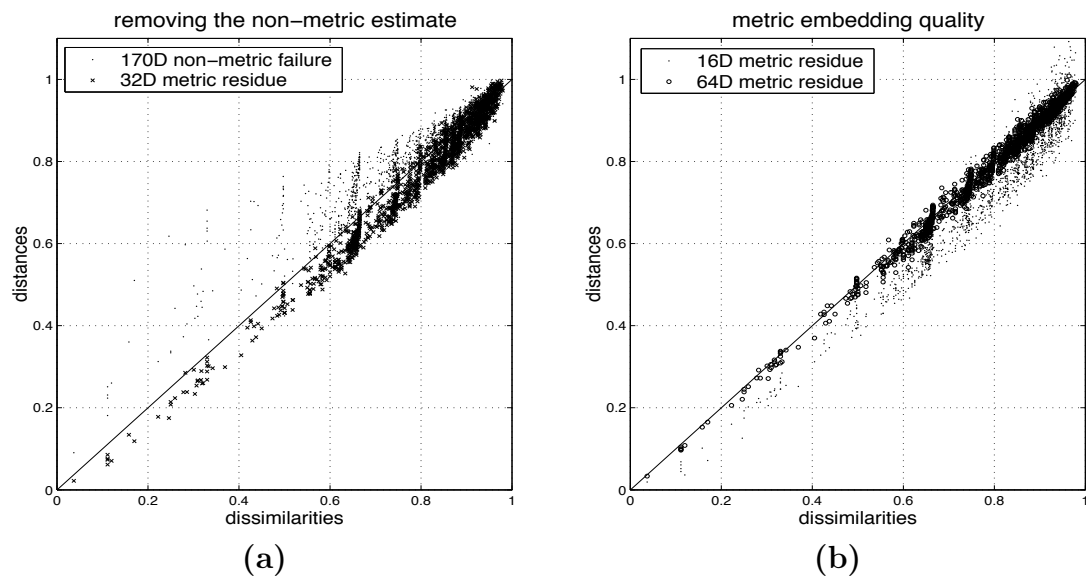
Abbildung 4.4.a zeigt die sortierten Eigenwerte. Wie bereits im letzten Kapitel erläutert, wird typischerweise 90% der Eigenwertenergie als relevanter Datenanteil betrachtet, die verbleibenden 10% als Rauschen. Hier ergibt sich eine Relevanzschwelle bei Dimension 68. Diese Schranke ist in beiden Plots in Abbildung 4.4 als durchgezogene senkrechte Linie zu sehen und gibt eine gute obere Schranke für die intrinsische Dimension der Daten. Eine exaktere Abschätzung erhalten wir jedoch über mehrere Läufe der MDS-Einbettung. Wir haben für jede Dimension 100 MDS-Läufe durchgeführt und dabei drei Meßgrößen

liefert uns der Einbettungsfehler  $\varepsilon_{\text{metric}}^D$ , der als durchgezogene Funktion in Abbildung 4.4.a eingezeichnet ist. Hierfür wird nach und nach die Einbettungsdimension für das MDS verringert und jeweils der Einbettungsfehler betrachtet. Ist die intrinsische Dimension unterschritten, kann nur noch mit einem erheblichen Fehler eingebettet werden, er steigt merklich an im Vergleich zum nicht-metrischen Defekt, den wir in höheren Dimensionen beobachten können.

Die Kurve in Abbildung 4.4.a zeigt einen konstanten Verlauf zwischen den Dimensionen 170 und 32 für mehrere Einbettungen mittels MDS. Danach steigt der Fehler signifikant an, was ein typisches Verhalten wäre für eine 32-dimensionale intrinsische Topologie.

In Abbildung 4.4.b sehen wir die Schranken auch durch eine zweite Meßgröße bestätigt: Hier betrachten wir die Standardabweichung und den Mittelwert des Einbettungsfehlers für mehrere Einbettungen mittels MDS mit zufälligen Initialverteilungen. Unterschreitet die Einbettung die intrinsische Dimension, so steigt auch die Anzahl der Freiheitsgrade und somit auch die Anfälligkeit, in möglichen lokalen Minima zu landen. Oberhalb der intrinschen Dimension ist ein relativ konstantes Verhalten bzgl. der statistischen Stabilität der Einbettung zu erwarten. Auch in diesem Plot sind die beiden Schranken bei Dimension 32 und 68 eingetragen. Es ist deutlich zu sehen, wie beide Größen für Dimensionen unterhalb der 32D-Schranke signifikant ansteigen.

In Abbildung 4.5 schließlich ist der sogenannte *Scatterplot* für unterschiedliche Einbettungsdimensionen dargestellt. In einem solchen Plot sind die Differenzen zwischen den Ausgangsdaten und der Einbettung abzulesen, da beide Werte paarweise



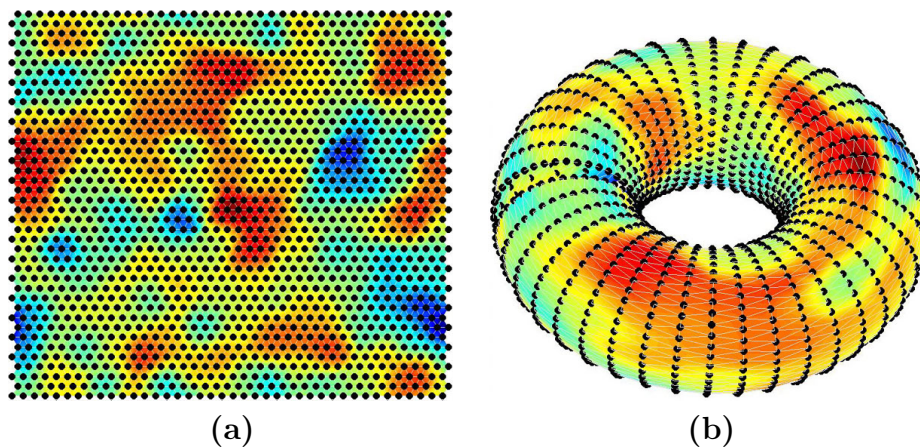
**Abbildung 4.5:** Scatterplots für die Geruchsdaten nach MDS. (a) Die kleinen Punkte zeigen den nicht-metrischen Defekt, d.h. den Einbettungsfehler der 171 Punkte in 170D. Die Kreuze zeigen die Güte der metrischen Einbettung für 32D. (b) Für 16D (kleine Punkte) ist der Einbettungsfehler sichtbar schlechter als für 32D, der Fehler für 64D (Kreise) ist vergleichbar zu 32D.

gegeneinander aufgetragen werden. D.h. im Idealfall erhofft man sich eine scharfe Einheitsgerade als Ergebnis. Dann nämlich sind die Daten mit einem minimalen Einbettungsfehler dargestellt.

In Abbildung 4.5.a sieht man, dass bereits schon die triviale Einbettung in  $\mathbb{R}^{170}$  bei weitem nicht perfekt ist. Wie zuvor beschrieben, haben wir diese nicht-metrische Fehlerkonstante beim Scatterplot für 32 Dimensionen entfernt. Was wir sehen ist nur noch der metrische Fehler bei einer Einbettung in 32D im Vergleich zur Einbettung in 170D. Vergleicht man diesen Plot zu den Scatterplots für 16D (kleine Punkte) und 64D (große Punkte) in Abbildung 4.5.b, so stellt man fest, dass die Einbettung in 64D in der Tat keine große Verbesserung bringt, die Verschlechterung durch die Einbettung in 16D ist aber mehr als deutlich zu erkennen.

#### 4.2.2 Kohonen-Karten zur Projektion von 32D in 2D

Wir haben nun also eine Einbettung der Geruchsähnlichkeiten in einen 32-dimensionalen Raum vorgenommen. Der nächste Schritt ist es, diese Daten in eine Darstellung zu bringen, die unseren Farbkreis-Rekonstruktionen (siehe z.B. Abb 3.14 aus dem vorangegangenen Kapitel) entspricht.

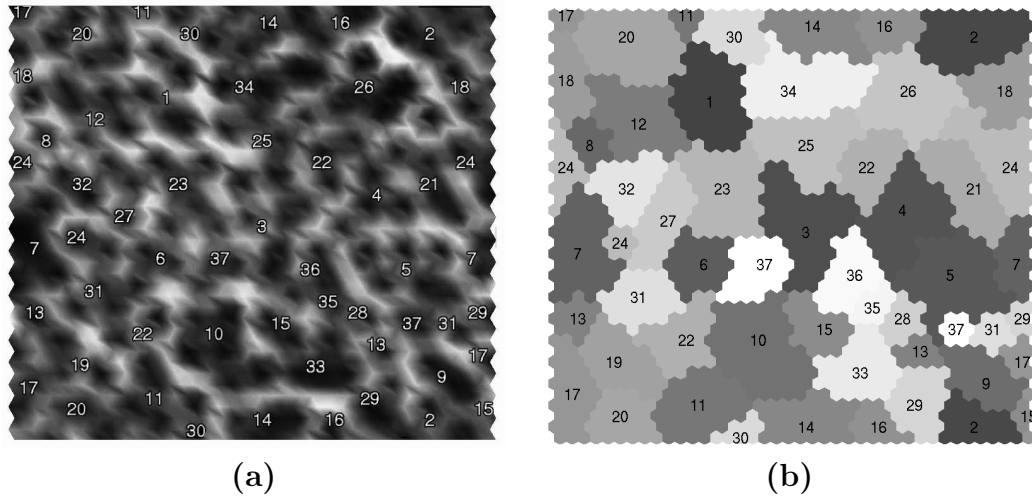


**Abbildung 4.6: Topologie der Kohonen-Karte.** Das 2D-Gitter der Kohonen-Karte kann dreifach strukturiert sein: Als einfaches Blatt (a), als Zylinder oder als Toroid (b). Die tatsächlich Dislokalisierung der Knoten im hochdimensionalen sorgt dort für eine Deformation der Topologie. In beiden Abbildungen sind die Abstände zwischen den Knoten farblich auf Blatt und Toroid projiziert.

Diese 32 Dimensionen scheinen eine gute Annäherung zu sein im Sinne einer möglichst geringen Verzerrung der tatsächlichen Ähnlichkeiten und einer möglichst starken Dimensionsreduktion.

Um diesen 32-dimensionalen Datenraum nun aber tatsächlich in eine für Menschen lesbare Version zu bringen, haben wir eine zweidimensionale Selbstorganisierende Karte (SOM, [77]) verwendet. SOMs (auch Kohonen-Karten genannt) erhalten die Topologie der Ausgangsdaten soweit wie möglich und liefern somit eine möglichst ähnlichkeitserhaltende Darstellung des hochdimensionalen Kontextes. Diese Karte beinhaltet 32-dimensionale Vektoren, welche die gegebenen Datenpunkte bzw. deren Topologie lernen. Diese Vektoren sind jedoch in einer festen zweidimensionalen Ordnung miteinander verbunden und können somit leicht als 2D-Karte angezeigt werden. Eine knappe Darstellung des hier verwendeten Verfahrens findet sich im Anhang A.2.

Es handelt sich bei der hier verwendeten Kohonen-Topologie um einen Torus, d.h. die linke und rechte Seite sind verbunden, ebenso wie oben und unten. In dreidimensionaler Projektion ist die Topologie also mit der eines Donuts zu vergleichen. Diese Variante einer zweidimensionalen Karte ist auch in Abbildung 4.6 dargestellt. Diese Darstellung des Torus verdeutlicht auch die Gefahr von sog. topologischen Defekten, die besonders dann auftreten können, wenn die zu lernenden Daten auf einer höherdimensionalen, nicht toroiden Mannigfaltigkeit liegen. Driften im Datenraum während des Lernens z.B. Ober- und Unterseite des Torus zusammen, so liegen die beiden Seiten zwar auf der 2D-Karte nach wie vor weit auseinander, im Datenraum jedoch beschreiben sie dicht benachbarte Punkte.



**Abbildung 4.7:** Geruchswahrnehmungsraum für 32D-MDS-Einbettung der Aldrich-Daten. (a) Die Distanzen zwischen benachbarten Knoten der SOM sind farblich dargestellt. Große Abstände sind hell markiert. (b) Die SOM-Knoten sind mittels des k-means-Verfahrens in 32D zu Clustern zusammengefasst worden. Die Zentren dieser Cluster und Cluster-Fragmente sind durch eine Nummer gekennzeichnet, z.B. Cluster 15 befindet sich in der rechten unteren Ecke und mittig im unteren Drittel der Karte. Die Nummern entsprechen den Nummern in (a).

Abbildung 4.7.a zeigt die Kohonen-Karte, die sich auf den 32-dimensionalen Geruchsdaten ergibt. Dargestellt ist hier der Abstand (im 32D-Datenraum!) benachbarter Knoten der Kohonen-Karte. Je heller die Verbindung, desto weiter liegen die Knoten auseinander. Ziehen sich die Knoten zu einer Gruppe dicht beieinander liegender Knoten zusammen, so spricht man von Clustern. Solche Cluster sind durch die Zahlen in der Abbildung gekennzeichnet. Die Cluster wurden im Datenraum durch das sogenannte *k*-means-Clusterverfahren [33] berechnet. Die Cluster sind in ihrer tatsächlichen Ausdehnung besser auf der Karte in Abbildung 4.7.b zu sehen.

Hier erkennt man auch die eben bereits erwähnten topologischen Defekte, die auf einer solchen Karte auftreten können: Einige der Cluster (z.B. 15 und 37) findet man an zwei verschiedenen Orten der Karte, obwohl diese im 32-dimensionalen Raum dicht beieinander liegen. Dies kann passieren, wenn sich die 2D-Struktur der Karte der Topologie der Datenverteilung im 32D-Raum nicht perfekt annähern kann. Bei einer Punktmenge mit einer intrinsischen Dimension von 32 ist dies auch nicht zu erwarten.

Für die Darstellung der Karten ebenso wie für das Lernen der Kohonen-Karten wurde die *“SOM-Toolbox for Matlab”* von Vesanto et al. [95] verwendet.



---

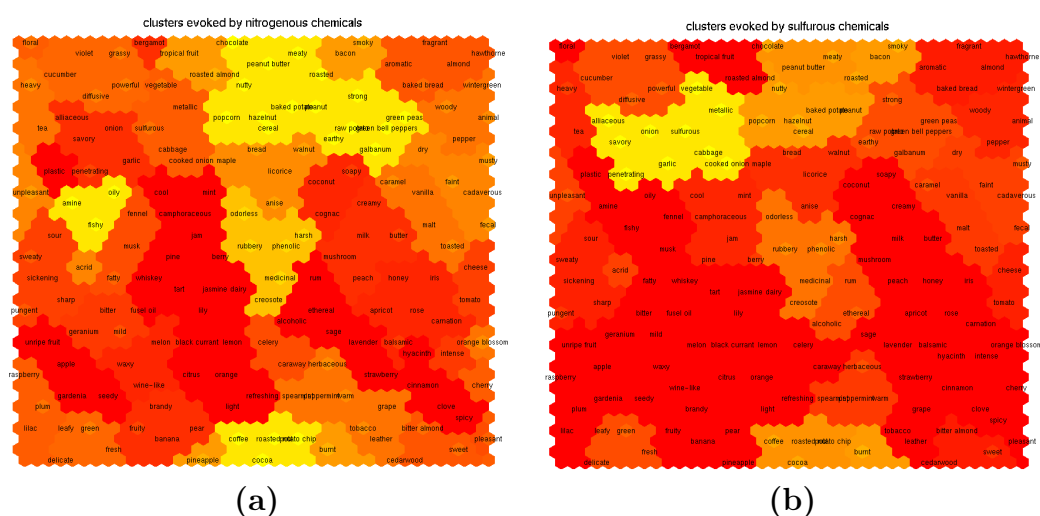
ständig auch weiterhin sämtliche topologischen Defekte, die in Abb. 4.7 sichtbar sind. Somit gehören z.B. die Bezeichner *celery*, *caraway* und *pleasant* allesamt in Cluster 15.

Auf diesen Karten können wir nun versuchen, die Zusammenhänge zwischen den Geruchsqualitäten herauszuarbeiten. Wir hatten zu Beginn dieses Kapitels beispielhaft nach den Zusammenhängen zwischen den Bezeichnern *apple*, *banana* und *cherry* gefragt. Dies wollen wir nun noch einmal auf der Basis der Karten näher beleuchten:

Betrachten wir also Abbildung 4.8 und vergleichend dazu die Clusterkarte in Abbildung 4.7.b, so stellen wir fest, dass *cherry* zu Cluster 17 gehört, *apple* zu Cluster 19 und *banana* zu Cluster 11. Da unsere Kohonen-Karte eine toroide Form hat, sind die Cluster 17 und 19 direkt benachbart, ebenso wie auch Cluster 11 direkt an Cluster 19 angrenzt. Andererseits findet sich jeweils immer mindestens ein anderer Cluster der zwischen den Clustern 11 und 17 liegt. Folglich scheint *cherry* ähnlicher zu *apple* als zu *banana* zu sein. *apple* wiederum liegt qualitativ irgendwo zwischen *cherry* und *banana*.

Solche Zusammenhänge lassen sich jetzt grundsätzlich für alle Bezeichner herausarbeiten. Man sollte sich jedoch eines Problems noch einmal bewusst werden, welches wir in der Einleitung bereits angedeutet haben: Einige der Bezeichner sind kaum miteinander zu vergleichen. Zwar kann man schnell sinnvoll erscheinende Zusammenhänge auf der Karte erkennen, z.B. dass *fecal* und *cadaverous* in einer Klasse mit *unpleasant* sind, bei anderen hingegen ist es nicht so klar, ob die Relationen stimmen oder ob es sich um irgendwelche Artefakte bei der Erzeugung der Karte handelt.

Nehmen wir als weiteres Beispiel die Relation zwischen *roasted almond*, *pineapple* und *tropical fruit*, welche sich gemeinsam in Cluster 30 (mittig, am oberen und unteren Rand von Abb. 4.8) befinden. *pineapple* und *tropical fruit* scheinen zusammenzugehören, aber *roasted almond* passt an diesen Ort höchstens aufgrund der Nähe zu den übrigen Beschreibern für gebackene und gekochte Gerüche, die auch dort liegen. Oder gibt es da doch einen Zusammenhang zwischen gerösteten Mandeln und tropischen Früchten? Diese Frage ist eigentlich nur erschöpfend zu beantworten durch experimentelle Ansätze, die gezielt einige aus den Karten herausgearbeitete Hypothesen testen.



**Abbildung 4.9: Cluster für Stickstoff bzw. Schwefel-tragende Substanzen.** Je heller ein Cluster ist, desto höher ist der Anteil der Beschreiber, die den Substanzen der Klasse mit Stickstoff (a) bzw. Schwefel (b) zugeordnet wird.

#### 4.2.4 Visualisierung von stoffklassenspezifischen Geruchsqualitäten

Abschließend sei noch eine weitere Anwendung vorgestellt, die demonstriert, wie der Wahrnehmungsraum für das Riechen durch die hier vorgestellten Karten langsam handhabbar wird.

Chee-Ruiter [15] hat auf den von ihr vorgestellten Geruchsgraphen festgestellt, dass besonders die Bezeichner von Qualitäten, die typischerweise von Stickstoff- bzw. Schwefel-tragenden Substanzen ausgelöst werden, einen starke Zusammenhangsgraphen auf ihren Karten formen. Dies erscheint durchaus plausibel, bedenkt man, dass ein Molekül, welches sich in so einer speziellen Konfiguration befindet, auch durch ganze Stoffwechselkreisläufe meist seine grundlegende Struktur nicht mehr verändert. Ist die menschliche Geruchswahrnehmung womöglich entlang der Struktur von tierischen und pflanzlichen Stoffwechsel-Kreisläufen orientiert?

Auch diese Frage werden wir im Rahmen dieser Arbeit nicht beantworten können, wir konnten jedoch durchaus Chee-Ruiters Thesen an unseren Karten auf Plausibilität testen. Wir betrachten jeweils alle Substanzen, die ein Stickstoff- bzw. ein Schwefel-Atom tragen. Für beide Klassen haben wir nun die Menge aller von ihnen typischerweise ausgelösten Beschreiber gebildet und auf unseren Karten geschaut, wie sie sich verteilen. Unter der Annahme, diese Auswahl erfolgte zufällig, sollten auch diese Beschreiber relativ zufällig auf der Karte verteilt sein. Eine Ordnung wäre nur zu erwarten, sollten die Gruppen doch in irgendeiner Art und Weise mit der Geruchsqualität korreliert sein.



---

Das Ergebnis dieser Untersuchung ist in Abbildung 4.9 zu sehen. Wir haben dazu wieder die Clusterkarte aus Abb. 4.7 bemüht und den Prozentsatz an N- bzw. S-Bezeichnern pro Klasse in den Farbwerten kodiert. Je höher der Wert ist, desto heller ist der Cluster dargestellt. Ist die Rate ganz niedrig, so ist der Cluster dunkel dargestellt. Sowohl die Stickstoff-Substanzen (**a**) als auch die Moleküle mit Schwefelatomen (**b**) scheinen Beschreiber auszulösen, die ziemlich klar gruppiert auf unseren Karten angeordnet sind.

Unter der Annahme, dass unsere Karten die Relationen zwischen den Bezeichnern korrekt darstellen, finden wir in beiden Fällen eine Bestätigung zu dem, was Chee-Ruiter auch auf ihren Geruchsgraphen gesehen hat.

## 4.3 Diskussion

Wir haben die einzige verfügbare Quelle für die Quantifizierung von Geruchsqualität verwendet, nämlich psychophysikalische Beschreibungen von Gerüchen durch menschliche Probanden.

Mit Hilfe von Multidimensional Scaling im Rahmen der in Kapitel 3 bereits vorgestellten Analyse konnten wir zeigen, dass der Geruchswahrnehmungsraum auf der Basis unserer umfangreichen Datenbank hochdimensional ist. Wir haben sowohl durch Abschätzung des Einbettungsfehlers wie auch über die statistische Varianz der Einbettungen Hinweise auf eine intrinsische Dimension von etwa 32 Dimensionen ausgearbeitet.

Durch die nachträgliche Projektion des Wahrnehmungsraums auf eine zweidimensionale Karte mittels Selbstorganisierender Karten konnten wir auch für die Wahrnehmung von Gerüchen eine Darstellung ableiten, die vergleichbar dem Farbkreis ist. Hier haben wir demonstriert, wie man für einzelne Bezeichner bereits Zusammenhänge ablesen kann.

Die Karte ist nicht so leicht zu lesen wie z.B. der Farbkreis, nachwievor haben wir zu dem hochdimensionalen Kontext des Riechens ja auch kein physikalisches Kontinuum, um Zusammenhänge exakt vorherzusagen. Dennoch liefert eine Karte wie die hier vorgestellte eine gute Basis, um in weiterführenden praktischen Experimenten verschiedene Hypothesen bzgl. der Kodierung von Gerüchen zu testen.

Wir haben gezeigt, wie man für eine beliebige Stoffklasse eine deutliche Visualisierung ihrer Geruchsqualitäten auf den Karten erreichen kann. Ebenso konnten wir auch eine von Chee-Ruiter aufgestellte Hypothese stützen, nach der die Ordnung entlang Stoffwechselkreisläufen – hier das Tragen einer Komponente Stickstoff bzw. Schwefel im Molekül – einen signifikanten Einfluss auf die Geruchswahrnehmung

hat, unabhängig von der potentiellen Organisation entlang der chemischen Struktur des zugehörigen Rezeptorsystems.

Ein wesentlicher Unterschied zu den bisherigen Ansätzen ist der Tatsache geschuldet, dass wir nun nicht mehr versuchen, den Merkmalsraum zu analysieren und so die Wahrnehmungen zu verstehen [81, 100]. Vielmehr verwenden wir umgekehrt die verfügbaren Beschreibungen einzelner Substanzen, um einen Wahrnehmungsraum aufzuspannen. Diese Erkenntnis ist nicht auf den Geruchssinn beschränkt sondern betrifft Wahrnehmungsprozesse jedweder Art.

### Teil III: Der kombinatorische Kode des Riechens

|   |   |   |   |   |   |   |   |   |
|---|---|---|---|---|---|---|---|---|
|   |   |   | 3 |   |   |   |   |   |
|   |   |   |   |   | 6 |   |   | 9 |
| 7 | 6 |   | 5 |   |   | 3 | 8 | 4 |
| 5 |   | 9 | 8 | 4 |   |   |   |   |
|   | 8 | 6 |   |   |   |   |   |   |
|   | 4 | 2 | 6 |   |   |   | 5 |   |
|   |   | 7 |   | 3 | 2 |   |   |   |
| 8 | 2 |   | 7 | 6 | 1 |   | 9 | 3 |
|   |   | 3 |   |   |   |   | 7 | 2 |

*Man muß das Unmögliche versuchen,  
um das Mögliche zu erreichen.*

Hermann Hesse

# 5 Rezeptororganisation im Riechkolben

## 5.1 Einleitung

Das Riechen ist für die meisten Säugetiere nach wie vor einer der wichtigsten Sinne. Auch das menschliche Geruchssystem ist extrem leistungsfähig. Es ist in der Lage, bereits kleinste Unterschiede zwischen zwei Duftstoffen zu erkennen. Selbst chemisch nahezu identische Substanzen wie Hexanol und Heptanol, also zwei Alkohole, die sich nur durch ein Kohlenstoff-Atom in ihrer Kettenlänge unterscheiden, riechen stark unterschiedlich [3]. Die Nase ist in der Lage, selbst solche minimalen Differenzen robust und reproduzierbar zu unterscheiden. Diese Fähigkeit, kleinste Unterschiede trennbar zu machen, nennt man auch *Odorant Fine Discrimination*.

Auf der anderen Seite ist die Nase aber auch in der Lage, zu chemisch ganz unterschiedlichen Substanzen (wie z.B. Blausäure und Benzaldehyd) eine identische Geruchswahrnehmung zu erzeugen (in diesem Fall z.B. Mandel) [79]. Menschen beschreiben diese Wahrnehmungen meist durch Worte, sog. Geruchsbeschreiber. Wenn wir von der Wahrnehmung von Gerüchen sprechen, wird eigentlich von diesen Worten bzw. dem was sie beschreiben gesprochen und nicht von der Sensation eines chemischen Moleküls [3, 18]. Wie wir in den beiden vorangegangenen Abschnitten über die Wahrnehmung von Farben in Kapitel 3 und Gerüchen in Kapitel 4 schon ausführlich beschrieben haben, geben diese Beschreiber eine gute Näherung für die Assoziationskette, die durch einen bestimmten Stimulus ausgelöst wird. Diese Assoziation von biochemisch unterschiedlichen Substanzen zu einer gemeinsamen Geruchsqualität nennt man auch *Odor Clustering*.

Somit scheint das Geruchssystem zwei gegenläufige Probleme zu lösen:

1. Sensitivität auf kleinste molekulare Unterschiede (Odorant Fine Discrimination, OFD)
2. Zusammenfassung unterschiedlicher Substanzen in einem Kontext (Odor Clustering, OC)

### 5.1.1 Sensorische Verschaltung im Riechkolben

Strukturell beginnt dieser Wahrnehmungsprozeß an der Oberseite der Nasenhöhle, genauer am sogenannten Riechepithel. Diese Schleimhaut beherbergt diejenigen Sinneszellen, mit deren Hilfe wir Moleküle riechen können. Durch das Einatmen werden Duftstoffe in die Nasenhöhle gesogen und lösen sich in der Riechschleimhaut. Spezielle Atemtechniken wie z.B. „Schnüffeln“ erhöhen die Wahrscheinlichkeit, dass ein Molekül in der Luftzirkulation der Nasenhöhle auf die Schleimhaut trifft. In der Schleimhaut gelöst kann es dann auf einen sensitiven Rezeptor treffen, der das entsprechende Molekül bindet und ein Aktionspotential an das Gehirn schickt.

Es ist in aufwändigen Experimenten gelungen, die Existenz von ca. 1000 verschiedenartigen Geruchsrezeptoren der Ratte nachzuweisen, d.h. 1000 Rezeptoren mit einem eigenen, substanzspezifischen Bindungsverhalten [12, 72]. Auch konnte man zeigen, dass alle Rezeptoren eines Typs zwar mehr oder weniger unvorhersagbar in der Riechschleimhaut verteilt liegen, die Rezeptoren eines Typs aber zu gleichen Verschaltungspunkten, den sogenannten Glomeruli, konvergieren. Diese Glomeruli liegen auf der nächsten Verschaltungsebene, dem Bulbus Olfactorius oder auch Riechkolben. Da die Zellen ihr Soma in der Schleimhaut haben und der Riechkolben jenseits des Siebbeins liegt, ist die Zuordnung zu genauen Lokalisierungen der einzelnen Rezeptortypen auf der Riechkolben-Seite nur lückenhaft dokumentiert.

Jeder Rezeptortyp wird typischerweise durch mehrere Substanzen stimuliert und jede Substanz passt meist zu mehr als einem Rezeptor. Es gibt weiterhin starke experimentelle Hinweise dafür, dass jede Substanz einen eindeutigen „Geruchsfingerabdruck“ im Sinne von Glomeruli-Aktivität besitzt [2, 21, 42, 62].

Diese Hypothese über reproduzierbare und individuelle Aktivitätsmuster für Einzelsubstanzen und ihre Übertragbarkeit auf unterschiedliche Individuen einer Spezies wurde in etlichen weiteren Studien (wie z.B. mittels [ $^{14}\text{C}$ ]2-Deoxyglucose (2DG) und Optical Imaging) bestätigt [34–43, 88, 93]. Diese gemessenen Aktivitäten korrelieren mit der Stimulation derjenigen Geruchsneuronen (Olfactory Sensory Neurons (OSNs)), die zu den betreffenden Glomeruli projizieren [68, 76].

Noch immer weitgehend unverstanden ist aber das Regelwerk, dem dieser „kombinatorische“ Kode unterliegt. Zwar konnte gezeigt werden, dass 2DG-Uptake-Bilder (siehe auch Abschnitt 5.1.2) räumlich spezifische Aktivitäten für unterschiedliche Stimuli zeigen [16, 28, 29, 54], dennoch fehlt weiterhin eine klare Systematik, die die zugrundeliegende Kodierung entschlüsselt. Es gibt keine klare Zuordnung der einzelnen Regionen des Bulbus zu ihrer Funktion im Sinne der Geruchswahrnehmung. Es ist aber ja auch gerade eine Eigenart einer Kombinatorik, eine ganze Vielzahl von Klassen und Qualitäten eben nicht spezifisch durch die Verwendung eines einzelnen Buchstabens zu kodieren, sondern aus den einzelnen Glomeruli-Buchstaben

---

eindeutige und klar zu erkennende Wörter zu generieren.

Basierend auf räumlichen Ähnlichkeiten der Muster ist es gelungen, Ähnlichkeiten in der Geruchswahrnehmung vorherzusagen [102], was unterstreicht, dass die Information über die Geruchsqualität einer Substanz bereits im räumlichen Muster der Glomerulus-Aktivierung enthalten ist. Aus informationstechnischer Sicht ist die Information über die Geruchsqualität natürlich bereits im eingehenden Signal, also auch schon auf Glomerulus-Ebene, enthalten, schließlich ist dies eindeutig der Sinn, mit dem Gerüche wahrgenommen werden. Wird der Nervenstrang zum Bulbus unterbrochen oder gar durch ein Schädelbasis-Trauma zerstört, ist kein Riechen mehr möglich.

Gottfried et al. [24] und Kadohisa und Wilson [45] haben in kürzlich erschienenen Arbeiten auch erste Verschaltungssystematiken auf der nächst höheren Wahrnehmungsebene nachweisen können: Die Riechrinde (Pyriform Cortex, siehe auch Kapitel 4) lässt sich funktionell in zwei Module unterteilen, dem vorderen und dem hinteren. Der vordere Teil des Pyriform Cortex (anterior PC) scheint unterschiedliche chemische Eigenschaften differenzieren zu können, während der hintere Teil (der posterior PC) auf ähnlich riechende Stimuli spezifisch reagiert. Dies ist insofern bemerkenswert, als beide Teile den gleichen neuronalen Input erhalten und entsprechend allein durch die individuelle Gewichtung der eingehenden Glomerulus-Muster die unterschiedlichen Funktionen bewirken müssen. Es gab bereits auch eher theoretische Abhandlungen, die diese Möglichkeit eingeräumt haben [47, 86].

Im folgenden werden wir versuchen, die verfügbaren Daten – 2DG-Uptake-Bilder von Rattenriechkolben (siehe Abschnitt 5.1.2) und Geruchsqualitätsbeschreibungen von Menschen (siehe Kapitel 4) – zusammenzuführen und etwas über Ordnung und Struktur jenseits der Ergebnisse aus Kapitel 4 ableiten zu können. Wir werden untersuchen, wie die räumlichen Inputmuster gewichtet verschaltet sein müssen, um die einzelnen Wahrnehmungseindrücke (OFD und OC) zu ermöglichen. Zu diesem Zweck haben wir 2DG-Uptake-Bilder von 211 Reinstoffen untersucht und mit Hilfe von linearen Klassifizierern – SVMs mit linearem Kernel [63] – die Relevanz einzelner Regionen berechnet und in einen Kontext zu den chemischen Eigenschaften und Geruchseindrücken der Substanzen gebracht. In den folgenden Abschnitten werden wir erst die Daten wie auch die Analysemethoden vorstellen und anschließend die Ergebnisse gründlich diskutieren.

### 5.1.2 Deoxyglucose-Mapping

Im Rahmen eines gemeinsamen Projektes mit der Gruppe um Michael Leon und Brett Johnson von der University of Irvine, California, USA, wurde uns ein umfangreiches Archiv von Uptake-Bildern des Rattenriechkolbens zur Verfügung gestellt. Die meisten der in dieser Arbeit verwendeten Daten sind unterdessen bereits auch

online verfügbar unter <http://leonservers.bio.uci.edu>. Mit ca. 385 Aktivitätsbildern von nahezu 90% aller im Bulbus befindlichen Glomeruli handelt es sich bei diesen Daten um eine der vollständigsten und umfangreichsten Sammlungen von Glomeruliaktivität.

Bei Deoxyglucose-Mapping nutzt man die erhöhte Metabolismusrate aktivierter Neurone, um deren Aktivität abzuschätzen. Aktive Neurone nehmen verstärkt Glukose auf, um verbrauchte Energiereserven wieder aufzufüllen. Man injiziert den Tieren kurz vor dem Riechexperiment mit  $^{14}\text{C}$  radioaktiv markierte Deoxyglucose, welche von den Neuronen zwar aufgenommen und phosphoryliert, aber nicht verstoffwechselt werden kann. In den aktiven Neuronen reichert sich so markiertes Deoxyglucosephosphat an. Diese Anreicherung von  $^{14}\text{C}$  kann detektiert werden und gibt somit Auskunft über die neuronale Aktivität [48, 87].

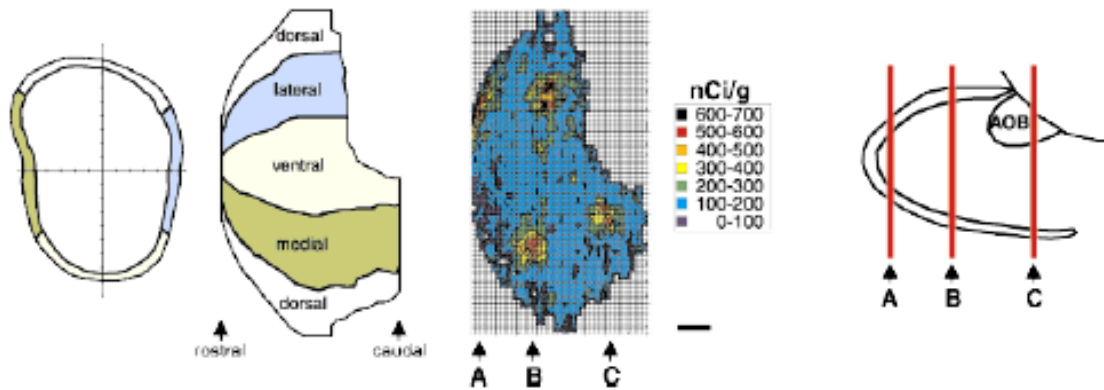
Nach der Injektion der Deoxyglucose in die Blutbahn wird ein Teil der Tiere isoliert, damit sie möglichst gar nichts riechen. Zumindest jedoch sollen sie genau dasselbe riechen wie die anderen Ratten, abgesehen von dem Duftstoff, dem die zweite Gruppe nun über einen längeren Zeitraum, bis zu 45 Minuten, ausgesetzt wird [36, 41].

Anschließend werden die Ratten geköpft und der Kopf tiefgefroren, um den Stoffwechsel zu stoppen. Der freipräparierte Riechkolben wird anschließend in hauchdünne Scheiben geschnitten. Jeder zweite Schnitt wird histologisch bearbeitet, um die Glomeruli sichtbar zu machen, die anderen Schnitte werden auf Fotopapier gelegt, um die aufgenommene Radioaktivität der einzelnen Abschnitte abzuschätzen. Um eine Referenz der Umgebungsluft ohne Stimulus zu haben, wird dasselbe auch mit der Kontrollgruppe gemacht.

Die aufgenommenen Radioaktivitäten werden nach der Anatomie des Rattenbulbus (siehe auch Kapitel 4) in eine standardisierte Datenmatrix eingetragen, bestehend aus einer festen Anzahl von Spalten entlang von anatomischen Markern auf der anterior-posterior-Achse des Bulbus. Da sich die Glomeruli auf der Oberfläche des Riechkolbens befinden, kann jeder koronale Schnitt in der ventral-zentrierten Sicht auf der dorsalen Seite „aufgeschnitten“ und in eine lineare 2D-Darstellung gebracht werden, so dass jedem Schnitt eine Spalte der Datenmatrix entspricht. In Abbildung 5.1 ist dieser Zusammenhang zwischen anatomischer Ausrichtung und Lage des Bulbus und den hier im weiteren verwendeten Datenmatrizen noch einmal dargestellt. Deutlich zu erkennen ist hier der konische Verlauf, dem folgend sich in caudaler Richtung vermehrt Glomeruli (und somit auch mehr Einträge in den Spalten) finden. Nur am Austrittspunkt des olfaktorischen Nerves klafft ein „Loch“ im abgerollten Konus. Dieses ist deutlich zu erkennen, da sich genau dort physiologisch nämlich auch keine Glomeruli mehr finden.

Jeder zweite Schnitt wird histologisch gefärbt, um die Positionen der Glomeruli markieren zu können. Für diese Regionen wird dann auf dem ungefärbten vor-





**Abbildung 5.1:** Jede Spalte der 2DG-Uptake-Bilder entspricht einem entfalteten koronalen Schnitt des Bulbus (zu sehen ganz links), im weiteren jeweils in einer ventral-zentrierten Perspektive. Im zweiten Bild ist hervorgehoben, wo sich in dieser Ansicht die dorsalen, lateralen, ventralen, medialen, rostralen und die caudalen Zonen des Riechkolbens befinden. Die aufgetragene Uptake-Information berechnet sich als relative Abweichung vom Standard (in  $nCi/g$ , siehe auch Skala, drittes Bild). Im rechten Bild ist noch einmal der rostral-caudale Zusammenhang zwischen Spalten und anatomischer Lokalisation dargestellt. Im Bereich des AOB finden sich keine Glomeruli mehr. Abbildung übernommen aus [41].

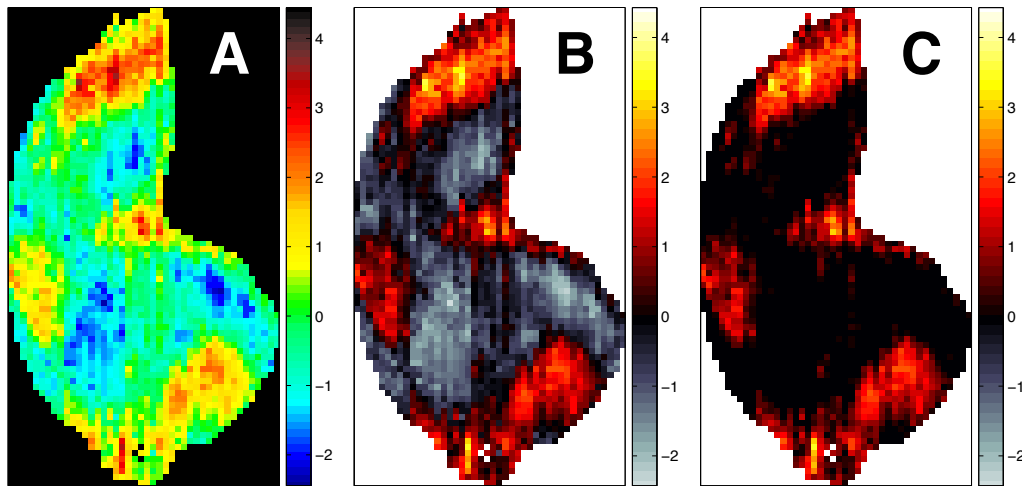
hergegangenen Schnitt der Uptake gemittelt. Ein solcher Wert repräsentiert also ein Glomerulus und entspricht einem Pixel auf den Uptake-Bildern wie in Abbildung 5.1 gezeigt. Um die jeweils aus unterschiedlichen Individuen und aus unterschiedlichen Testreihen stammenden Bilder vergleichbar zu machen, wird der radioaktive Uptake für ein Bild  $\mathbf{x} \in \mathbb{R}^{80 \times 44}$  typischerweise durch den sogenannten *z-score* beschrieben. Dieser Wert drückt die Abweichung der gemessenen Strahlung  $r_i$  an einem Pixel  $x_i$  von der mittleren Strahlung  $\bar{r}$  über alle Pixel dieses Bildes aus:

$$x_i = \frac{r_i - \bar{r}}{\sigma_r}$$

wobei  $\sigma_r$  die Standardabweichung über alle Pixel ist.

Die Daten, die uns von Johnson et al. zur Verfügung gestellt wurden, lagen alle bereits als *z-score*-normierte Bilder vor, daher werden im folgenden die *z-score* Bilder auch als Uptake-Bilder bezeichnet. Da unsere Farbkodierung von der in den Publikationen von Johnson et al. verwendeten abweicht, liefert Abbildung 5.2 A-C noch einmal eine Gegenüberstellung der originalen *z-score*-Farbkodierung von Johnson et al. (A) und der Darstellung desselben Bildes mit unserer Kodierung, die besonders die positiven *z-Scores* betont (B). Im Laufe dieses Kapitels werden wir unsere Sicht auf die positiven *z-scores* reduzieren. Dies ist in (C) dargestellt.

Es liegt in der Natur dieses Messverfahrens, dass die Daten ein gewisses Grundrauschen enthalten, welches zentral durch die spontane Aktivität der Neurone aus-



**Abbildung 5.2:** Farbkodierung der Uptake-Aktivität einer Einzelsubstanz am Beispiel von *Menthone*. **A:** Die originale Darstellung der uptake-Bilder über die angegebene z-Score-Farbtabelle nach Johnson & Leon **B:** Neue Farbkodierung, um den positiven gegen den negativen Uptake hervorzuheben. **C:** Nur positive z-Scores.

gelöst wird. Diese spontane Aktivität wird auch Nullaktivität genannt. Da jedoch die meisten Neurone besonders bei der Stimulation durch Einzelsubstanzen eher inaktiv bleiben werden, erwarten wir eine Datenverteilung, die aus zwei Gauß-Verteilungen zusammengesetzt ist: Einer Gauß-Verteilung um die spontan aktiven Neurone und einer zweiten Gauß-Verteilung um die tatsächlich durch die Stimulation feuernden Neurone. Über das zeitliche Fenster der Aufnahme betrachtet, wird ein „echt“ stimuliertes Neuron eine höhere mittlere Aktivität zeigen als ein nur ab und zu spontan feuerndes Neuron.

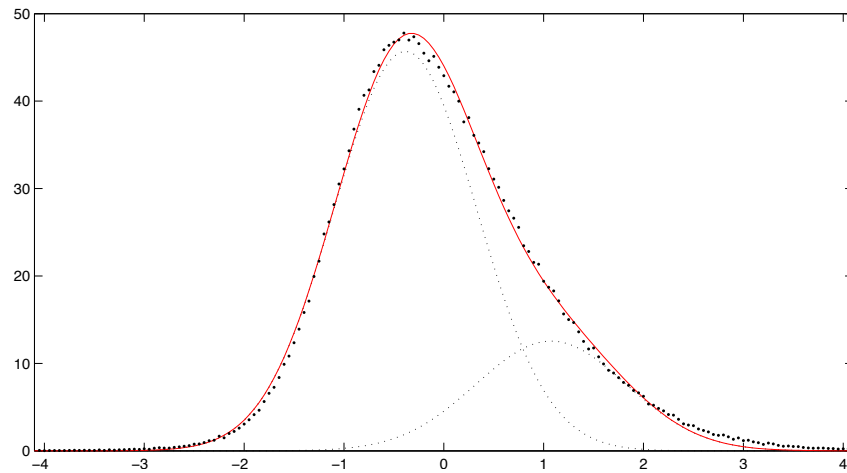
In Abbildung 5.3 zeigen die dickeren schwarzen Punkte eine Histogrammfunktion  $h(x)$  über die in allen 211 Datenmatrizen gemessenen z-Scores. Für  $h(x)$  wurden mit einer Schrittweite von 0.05 alle skalaren z-Scores  $x \in \mathbf{X}$  im entsprechenden Intervall aufsummiert.

Der Peak der Verteilung liegt nicht exakt bei Null sondern – deutlich nach links verschoben – im negativen Bereich. Dies entspricht beiden Annahmen, sowohl dass die meisten Glomeruli nur spontane Aktivität zeigen wie auch dass sich noch eine zweite (kleinere) Verteilung von aktiven Neuronen mit eher positiven z-Scores rechtsseitig von dieser Verteilung findet.

Unter der Annahme, dass sich die Histogrammfunktion linksseitig des Maximums  $x_{\max}$  mit

$$x_{\max} = \arg \max h(x)$$

nur aus spontaner Neuronenaktivität zusammensetzt, haben wir – um diesen Zu-



**Abbildung 5.3:** Die fetten schwarzen Punkte zeigen das Histogramm der originalen z-Scores über alle 211 Bilder. Deutlich sieht man die Verschiebung der Wertevertelung in den negativen Skalenbereich. Die beiden gepunkteten Linien zeigen die an das Datenhistogramm angepassten Gauß-Verteilungen. Die durchgezogene Linie zeigt den Verlauf der Summe der beiden Gauß-Verteilungen.

sammenhang noch weiter zu verdeutlichen – die linke Flanke von  $h(x)$  gespiegelt, um so die Histogrammfunktion  $h^0(x)$  der nicht aktiven Neurone abzuschätzen.  $h^0(x)$  wird also erzeugt mit

$$h^0(x) = \begin{cases} h(x) & \text{falls } x \leq x_{\max} \\ h(x_{\max} - |x - x_{\max}|) & \text{sonst.} \end{cases}$$

Entsprechend ergibt sich auch eine Abschätzung für die Histogrammfunktion  $h^1(x)$  der aktiven Neurone als

$$h^1(x) = h(x) - h^0(x)$$

Anschließend haben wir an beide Funktionen iterativ eine Gauß-Kurve für die beiden Verteilungen an  $h^0(x)$  und  $h^1(x)$  angenähert. In beiden Fällen wurde eine Funktion

$$N(x) = A \exp^{-0.5((x-\mu)/\sigma)^2} / \sqrt{2\pi\sigma^2}$$

mit der Amplitude  $A$ , dem Mittelwert  $\mu$  und der Standardabweichung  $\sigma$  der Gaußfunktion an das Histogramm angepasst. Hierzu haben wir eine least-squares Routine von J. R. Blake<sup>1</sup> verwendet. Die resultierenden Gauß-Kurven sind in Abbildung 5.3 gepunktet dargestellt, die Summe der beiden Kurven ist zum Vergleich mit dem Histogramm als durchgezogene Kurve ebenfalls eingezeichnet.

Um nun nur mit wirklich aktiven Neuronen zu rechnen, müssten wir eigentlich

<sup>1</sup>FITGAUSS

<http://www.mathworks.com/matlabcentral/fileexchange/loadFile.do?objectId=7489>

den Maximum-Likelihood der beiden Kurven zugrunde legen (welcher sich auf dem Schnittpunkt der beiden Gauß-Kurven befinden würde) und alle Datenwerte linksseitig davon als „wahrscheinlich“ spontane Aktivität verwerfen. Da diese beiden Kurven sich jedoch sehr stark überlagern und wir somit trotz starkem Rauschen beinahe die Hälfte aller „echten“ Daten verwerfen würden, haben wir uns entschieden, im weiteren alle positiven z-Scores zu verwenden und nur die in der Mehrheit negativen z-Scores der spontan aktiven Neurone zu verwerfen. Immerhin ist zu erwarten, dass die meisten der Falschwerte sich betragsmäßig nahe bei Null befinden werden.

### 5.1.3 Klassen und Klassenbeschreibungen

Zuallererst besitzt jedes Molekül bestimmte chemische Eigenschaften wie Kettenlänge, funktionale Gruppen und Lösungseigenschaften. Aus biochemischer Sicht wird die Beschreibung von Ähnlichkeiten grundsätzlich komplexer, da auch noch die Molekülfaltung und andere physikalische Faktoren eine Rolle bzgl. der Bindung an den Rezeptorneuronen (ORPs) im Riechepithel (OE) spielen.

Johnson et al. versuchen in vielen ihrer Arbeiten die chemische Organisation der Rezeptoren und die dadurch induzierte funktionale Ordnung im Bulbus zu entschlüsseln [35–37]. Daher steht uns für die Substanzen, von denen uns Uptake-Bilder zur Verfügung stehen, ebenfalls ausführliche Informationen über die chemischen Eigenschaften der Moleküle zur Verfügung, die wir als Klasseninformation nutzen können. Insgesamt sind die Substanzen durch 48 unterschiedliche Moleküleigenschaften beschrieben, die alle im Anhang in Tabelle B.2 aufgelistet sind.

Natürlich haben wir als weitere Quelle für Klasseninformationen die Geruchsbeschreibungen, die wir bereits im Kapitel 4.1.3 eingeführt haben. Die 278 Beschreiber, die hier für die Beschreibung von etwa 851 Substanzen verwendet wurden, finden sich in Tabelle B.1, ebenfalls im Anhang.

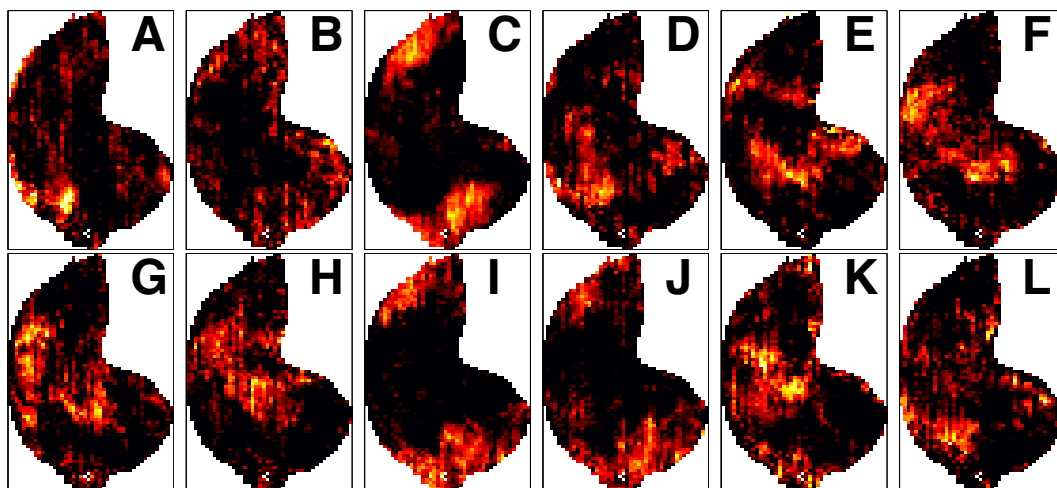
Leider haben wir nicht für jede Substanz auch eine Geruchsbeschreibung und umgekehrt auch nicht für jede Geruchsbeschreibung ein Uptake-Bild. In der Schnittmenge beider Beschreibungsklassen bleiben 211 Bilder, für die wir FEMA<sup>2</sup>-Nummern und somit eine Zuordnung zu unserer Geruchsdatenbank hatten. Diese 211 Substanzen bestehen aus 143 unterschiedlichen Duftstoffen. Einige liegen mehrfach vor, da Versuchsreihen in anderen Zusammensetzungen wiederholt wurden oder eine Substanz speziell auf die Auswirkung unterschiedlicher Konzentrationen untersucht wurde. Im Schnitt wird jede dieser 211 Substanzen durch etwa 3.5 von 139 verwendeten Geruchsbeschreibern charakterisiert.

Jedes der 211 Uptake-Bilder  $\mathbf{x}_i$ ,  $i = 1, \dots, 211$  hat  $(80 \times 44) = 3520$  Pixel, wobei –

---

<sup>2</sup>Flavor and Extract Manufacturers Association

durch die konische Form des Bulbus bedingt – nur 2121 Vordergrundpixel tatsächlich z-Scores enthalten. Um im weiteren besser mit den Daten rechnen zu können, wurde jedes Bild  $\mathbf{x}_i$  spaltenweise in eine vektorisierte Form  $\mathbf{x}_i \in \mathbb{R}^{2121}$  gebracht und als Zeilenvektor in eine Datenmatrix  $\mathbf{X} \in \mathbb{R}^{211 \times 2121}$  überführt.



**Abbildung 5.4:** Für 12 der 37 Bilder der Klasse *Alkohole ohne Phenol* sind die positiven z-Werte aufgezeigt. Jeder Duftstoff erzeugt ein eindeutiges Aktivitätsprofil im Bulbus. Ähnliche Moleküle erzeugen auch ähnliche Profile, es ist jedoch deutlich zu erkennen, wie unterschiedlich hier die Profile der einzelnen Alkohole sein können. Gezeigt sind A: 1-Propanol, B: 1-Butanol, C: 1-Pentanol, D: Geraniol, E: 1-Heptanol, F: 1-Oktanol, G: 1-Nonanol, H: 1-Dekanol, I: 2-Heptanol, J: 2-Methyl-3-Buten-2-ol, K: Phytol und L: 1-Undecanol.

Wir werden alle unsere Berechnungen in der vektorisierten Form durchführen, die meisten Ergebnisse allerdings in der Matrixform darstellen. Die Umrechnung ist vertretbar, da die Bilder auch spaltenweise aus den Koronalschnitten des Bulbus gewonnen werden. In den folgenden Abschnitten dieses Kapitels werden die Begriffe „Bild“ und „Vektor“ in Bezug auf die Uptake-Daten synonym verwendet, da hier jeweils auf die Zeileneinträge der Matrix  $\mathbf{X}$  verwiesen ist.

Diese 211 Uptake-Bilder gehören insgesamt zu 36 der oben erwähnten chemischen Eigenschaften und werden durch insgesamt 44 Geruchsqualitäten beschrieben. Alle Elemente aus  $\mathbf{X}$ , die eine solche Eigenschaft vereinen, nennt man auch eine *Klasse*.

Formal definieren wir für jede der 36 chemischen Eigenschaften einen sog. *Labelvektor*  $L_\mu^c \in \{0, 1\}^p$ ,  $\mu = 1, \dots, 36$  für alle  $p = 211$  Uptake-Vektoren mit  $n = 2121$  Pixeln aus  $\mathbf{X} \in \mathbb{R}^{p \times n}$  mit

$$L_\mu^c(j) = \begin{cases} 1 & \text{falls Bild } x_j \text{ zu einer Substanz mit chem. Eigenschaft } \mu \text{ gehört.} \\ 0 & \text{sonst.} \end{cases}$$

Ebenso definieren wir für jede der 44 Geruchsqualität einen entsprechenden Labelvektor  $L_\nu^o \in \{0, 1\}^p$ ,  $\nu = 1, \dots, 44$  mit

$$L_\nu^o(j) = \begin{cases} 1 & \text{falls Bild } x_j \text{ von einer Substanz mit Geruchsqualität } \nu \text{ stammt.} \\ 0 & \text{sonst.} \end{cases}$$

Es sei schließlich noch angemerkt, dass in diesen Kapiteln in der Regel die englischen Namen der Gerüche und häufig auch für die chemischen Klassen verwendet wurden. Die meisten Bezeichner sind den deutschen Bezeichnungen sehr ähnlich, es befindet sich im Anhang dieser Arbeit aber zusätzlich noch eine Übersetzungstabelle (siehe Anhang B.1).

In Abbildung 5.4 sind Bilder einer der chemischen Klassen, nämlich die Klasse der Alkohole, gezeigt. Es sind z-Score-Bilder von 12 der 37 in **X** verzeichneten Alkohole abgebildet. Man kann deutlich erkennen, dass jede Substanz ein charakteristisches Profil besitzt, dass sich die Bilder aber auch innerhalb dieser einen Klasse teilweise sehr stark unterscheiden. *Nonanol* (G) und *Dekanol* (H) etwa zeigen beinahe identische z-Scores, während *2-Heptanol* (I) ein komplett entgegengesetztes Muster präsentiert.

### Ein-Klassen Klassifikation mit Mittelwerten

Es ist intuitiv eine gute Wahl, eine Gruppe von Geruchsstoffen, z.B. die Gruppe aller Alkohole, durch den Mittelwert aller zugehörigen Bilder zu repräsentieren.

Sei also z.B.  $I_{alc}$  die Indexmenge aller Bilder der Klasse *alcohols* gegeben durch

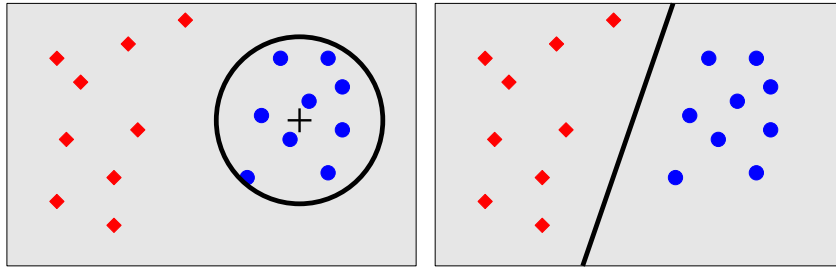
$$I_{alc} = \{i \mid L_{alc}(i) = 1 \quad \forall i = 1, \dots, p\}$$

so würde sich das zugehörige Mittelwert-Template  $\mathbf{c}_{alc}$  berechnen lassen durch

$$\mathbf{c}_{alc} = \frac{1}{|I_{alc}|} \sum_{i \in I_{alc}} \mathbf{x}_i$$

wobei  $|I_{alc}|$  die Anzahl der Elemente in der Indexmenge  $I_{alc}$  bzw. die Elemente in der Klasse der Alkohole *alc* beschreibt.

Betrachtet man Klassen aus dieser Perspektive, sind die Pixel (also Glomeruli), die im Mittel eine hohe Aktivierung zeigen, geeignete Kandidaten, um an dieser Stelle klassenspezifische Rezeptoren zu vermuten. Geometrisch betrachtet erwarten wir also, dass sich z.B. alle Alkohole gaußverteilt um ihr Mittelwert-Bild  $\mathbf{c}_{alc}$  verteilen, und vor allem erwarten wir implizit, dass alle Nicht-Mitglieder geeignet weit davon entfernt liegen.



**Abbildung 5.5:** Links: Der Mittelwert wird als Template der Klassenelemente genommen, die Entscheidungsschwelle ergibt sich aus dem Abstand der Trainingsbilder zum Template. Rechts: Der DMO-Algorithmus wählt aus der Menge aller linearen Lösungen diejenige Hyperebene, welche den minimalen Abstand von Trainingsdaten zur Ebene maximiert.

Für einen beliebigen Radius  $R$  kann man nun also auch unbekannte Bilder  $\mathbf{x}$  mit dem Mittelwert-Template  $\mathbf{c}_{alc}$  vergleichen und mittels  $R$  entscheiden, ob sie noch zur Klasse gehören oder aufgrund ihres hohen Abstandes nicht mehr. Wir formulieren das Ein-Klassen-Problem für ein Testbild  $\mathbf{x}_{test} \in \mathbb{R}^n$  und seine gesuchte Klassenzuordnung  $Y_{test}^{kugel} \in \{0, 1\}$  wie folgt

$$Y_{test}^{kugel} = \begin{cases} 1 & \text{falls } |\mathbf{x}_{test} - \mathbf{c}_{alc}| < R \\ 0 & \text{sonst.} \end{cases}$$

Dies entspricht einer (Hyper-)Kugel mit Radius  $R$  und Zentrum  $\mathbf{c}_{alc}$ , wobei alle Punkte innerhalb der Kugel zur Klasse gehören. Diese Art von Klassifikation ist in Abbildung (5.5) links illustriert.

Es handelt sich in diesem Fall aber nicht um ein echtes Ein-Klassen-Problem, da wir auch Informationen über die Nicht-Klasse haben, nämlich alle Bilder von Substanzen, die nicht zu der Klasse gehören. Beim Mittelwert werden so auch solche Regionen zu den signifikanten Klassenregionen gezählt, wenn dort z.B. einfach für *alle* Riechstoffe Aktivität anliegt.

## Lineare Zwei-Klassen-Klassifikation

Auch für das Zwei-Klassen-Problem gibt es eine ganze Reihe von Ansätzen, mit denen man versuchen kann, eine sinnvolle Lösung für das gegebene Problem zu finden. Wir suchen nun also nicht nur eine gute Repräsentanz für die eine Klasse, sondern auch eine möglichst universelle Unterscheidung zu den restlichen Daten, welche nicht zu der Klasse gehören. Wir suchen also genau die Merkmale, die die beiden Gruppen von Datenpunkten möglichst gut trennen.

In diesem speziellen Fall haben wir zwar echte Zwei-Klassen-Information vorliegen, jedoch handelt es sich um relativ wenig Punkte für die vorliegende Datendimension. Wir haben nur 211 Uptake-Bilder in einem 2121-dimensionalen Vektorraum. Da im Bulbus eine gewisse anatomische Symmetrie herrscht, halbiert sich diese Dimensionalität zwar vermutlich, aber auch dann liegen die 211 Punkte noch immer in einem etwa 1000-dimensionalen Raum.

In einer solchen Konfiguration werden wir vermutlich selbst mit einem sehr einfachen linearen Klassifizierer (einer linearen Hyperebene) eine Lösung finden, ganz gleich, wie die vorliegenden Punkte den Klassen zugeordnet sind.

Um die Hyperebene zu finden, die die gegebenen Datenkonfigurationen bestmöglich trennt, haben wir das sogenannte DoubleMinOver-Verfahren (DMO) verwendet, ein iterativer Lernalgorithmus, welcher für linear separierbare Probleme – und ein solches haben wir hier ja – gegen die Lösung einer Support-Vektor-Maschine konvergiert [64]. Dies ist die Position der Hyperebene mit dem maximalen Margin, also dem größtmöglichen Abstand zu beiden Klassen. Aus Sicht des Maschinenlernens stellt diese Lösung eine optimale Lösung im Sinne der Abstrahierbarkeit dar (d.h. optimal für die erwartete Güte der Klassifikation von unbekannten Daten, siehe z.B. [94]).

Details zum hier verwendeten Algorithmus findet sich im Anhang A.3.

Wir formulieren das lineare Zwei-Klassen-Problem entsprechend zum Ein-Klassen-Problem für ein gegebenes Testbild  $\mathbf{x}_{test}$  und eine gesuchte Zuordnung  $Y_{test}^{lin} \in \{0, 1\}$  als

$$Y_{test}^{lin} = \begin{cases} 1 & \text{falls } \mathbf{w}^T \mathbf{x} \geq \theta \\ 0 & \text{sonst.} \end{cases} \quad (5.1)$$

$\mathbf{w}$  ist der Normalenvektor senkrecht zu der Klassifikationsebene und  $\theta$  definiert den Schwellwert, mit dem die Entscheidungsebene in die Mitte zwischen beide Klassen geschoben wird. Geometrisch betrachtet, zeigt  $\mathbf{w}$  in Richtung von Klasse 1, d.h. Elemente gehören zu Klasse 1 genau dann wenn Gleichung 5.1 erfüllt ist.

Mit anderen Worten, der durch den Gewichtsvektor  $\mathbf{w}_\mu$  beschriebene lineare Klassifikator für eine Merkmalsklasse  $\mu$  hat die Linearkombination von Pixeln gelernt, die am besten geeignet ist, um zu entscheiden, ob ein gegebenes Uptake-Bild  $\mathbf{x}$  zu der entsprechenden Klasse  $\mu$  gehört oder nicht.

Das bedeutet aber auch (da wir hier nur positive Aktivitätsbilder betrachten): Die Gewichtsvektoren  $\mathbf{w}$  haben hohe positive Einträge für die klassenspezifischen Pixel (Glomeruli) und stark negative für die Pixel (Glomeruli), die das Trennen der Nicht-Klasse besonders unterstützen. Diese Entscheidungsbilder wurden auch schon erfolgreich für die Bestimmung des Geschlechtes anhand von Gesichtsbildern verwendet [98]: Verlängert man den Gewichtsvektor im Merkmalsraum, erhält man nach und nach eine Art „Karikatur“ des Klassenmerkmals.



---

Somit trägt  $\mathbf{w}_\mu$  für die Klasse  $\mu$  charakteristische Informationen über die Klasse, da die Komponenten von  $\mathbf{w}_\mu$ , die die höchsten Werte haben, auch die prädiktivsten Komponenten für diese Klasse sind. Da wir in den meisten Darstellungen im folgenden nur an diesen positiven Werten interessiert sind, definieren wir, basierend auf dem Gewichtsvektor  $\mathbf{w}_\mu$ , ein *Decision Image* bzw. einen *Decision Vektor*  $\mathbf{d}_\mu \in \mathbb{R}^{21 \times 21}$  mit

$$d_\mu(i) = \begin{cases} w_\mu(i) & \text{falls } w_\mu(i) \geq 0 \\ 0 & \text{sonst.} \end{cases}$$

## 5.2 Uptake-Pattern des Riechkolbens

### 5.2.1 Chemische Eigenschaften und ihre Repräsentanz im Bulbus

Im letzten Abschnitt wurden bereits zwei Ansätze vorgestellt, die die Charakteristiken der Klassen  $i$  auf den Uptake-Bildern  $\mathbf{x}_\mu$  beschreiben können. Wir wollen in diesem Abschnitt das Mittelwertbild  $\mathbf{c}_i$  und das Decision Image  $d_i$  auf den realen Daten miteinander vergleichen. Die beiden Templates können sich zwar sehr ähnlich sein, es können sich jedoch durch die Hinzunahme der Nicht-Klasse ins Training durchaus ziemliche Unterschiede entwickeln.

Ein grundsätzliches Problem für die Mittelwertbilder ist die auch klassenintern auftretende weite Streuung von Aktivität über den gesamten Bulbus. Dies soll im folgenden noch einmal an der Gruppe der Alkohole verdeutlicht werden: In Abbildung 5.4 sind 12 der 37 verfügbaren Uptake-Bilder  $\mathbf{x}_{alc}$  der Klasse *alcohols* dargestellt. Während einige durchaus ähnliche Muster zeigen (z.B. die Alkohole mit ähnlicher Kettenlänge, Abb. 5.4 links, zweite Reihe: *1-Heptanol* (E) bis *1-Decanol* (H)), liegen für andere Alkohole Muster vor, die nahezu komplementäre Aktivität zeigen (Abb. 5.4 erste Spalte: *1-Propanol*, *1-Heptanol* und *2-Heptanol*).

Wie eben bereits erwähnt wurde, umfassen die aktiven Bereiche der Alkohole weite Teile des Bulbus. Diese Regionen sind aber nicht notwendigerweise alle prädiktiv für Alkohole und nicht alle anderen Regionen sind notwendigerweise unwichtig für die Entscheidung ob es sich um einen Alkohol handelt oder nicht.

Bei der Analyse von homologen Reihen wie z.B. für die primären Alkohole konnte gezeigt werden, dass sich der Schwerpunkt der substanzspezifischen Aktivität gemäß der Kettenlänge des Moleküls entlang der lateral-rostralen Achse verschiebt [39]. Dieses Erkenntnis unterstreicht noch einmal die Tatsache, dass ohnehin kaum zu erwarten ist, einen klassenspezifischen Rezeptor zu finden, dessen Position direkt aus einem Mittelwert-Template abzulesen ist.

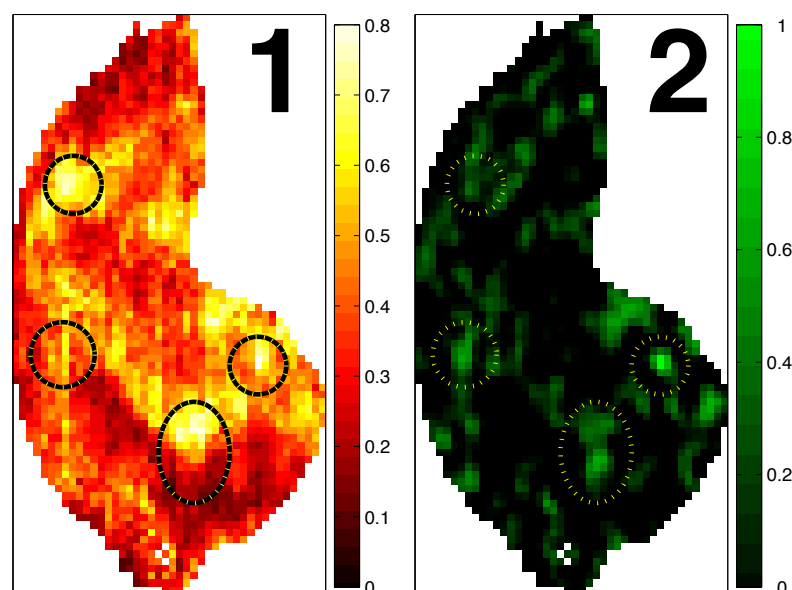
In Abbildung 5.6 sind nun noch einmal das Mittelwertbild  $\mathbf{c}_{alc}$  (links) und das zugehörige Decision Image  $\mathbf{d}_{alc}$  (rechts) für die Klasse der Alkohole zu sehen. Die vier markierten Regionen unterstreichen hier noch einmal die potentiellen Unterschiede zwischen den beiden Klassenbildern. Der caudal-ventrale Kreis zeigt eine Region, die in beiden Mustern eine hohe Signifikanz aufzeigt. Der rostral-ventrale Kreis zeigt hingegen eine Region, die keine auffällige Durchschnittsaktivität für die ganze Klasse zeigt, die aber auf dem Entscheidungsbild deutlich in Erscheinung tritt. Im lateralen Bereich gezeigt ist ein lokales Maximum im Durchschnittsbild, welches im Decision Image keine ausgezeichnete Rolle spielt. Ebenso verliert in der medial positionierten Ellipse die Region mit hohen Mittelwerten an Bedeutung, eine in dorsaler Richtung angrenzende Region wiederum ist signifikant für das Decision Image.

An diesem Beispiel sei auch nochmals erwähnt, dass das Decision Image sich auf eine wesentlich kleinere Entscheidungsfläche zusammenzieht. Abbildung 5.6.1 vermittelt mit dem Mittelwertbild  $\mathbf{c}_{alc}$  für die Klasse der Alkohole – wie jedes der Durchschnittsbilder – den Eindruck, ein Drittel der Bulbusfläche sei von klassenspezifischen Glomeruli durchzogen. Dies ist allerdings sehr unwahrscheinlich. Wie eingangs erwähnt, laufen alle ORNs eines Rezeptor-Typs zu einem Punkt im Bulbus zusammen, d.h. selbst wenn es für dieses Merkmal mehrere rezeptive Neurone geben sollte, würde man diese konzentriert an einigen wenigen Punkten erwarten. Das Decision Image  $\mathbf{d}_{alc}$  in Abbildung 5.6.2 zeigt eine solche Konzentration auf einige wenige Punkte.

Laufen in einer Region Rezeptoren zusammen, die nur auf Alkohole reagieren, so ist diese Region natürlich auch prädestiniert, eine wichtige Rolle im Decision Image zu spielen, denn für Nicht-Alkohole werden wir dort nie Aktivität finden. Somit ist Aktivität dort ein gutes Entscheidungskriterium für den Gewichtsvektor  $\mathbf{w}_\mu$  und somit auch für  $\mathbf{d}_\mu$ . Die Peak-Regionen im Decision Image (DI) sind folglich also gleichzeitig auch potentielle Kandidaten für die Konvergenzpunkte klassenspezifischer ORNs.

In Abbildung 5.7 sehen wir für 5 weitere Klassen Durchschnittsaktivität und DI entgegengestellt. Für die Aromaten mit Sauerstoff (erste Spalte) sowie für die langkettigen aliphatischen Aldehyde ziehen sich die Entscheidungsregionen des DI zwar stark zusammen, sie entsprechen jedoch weitestgehend Regionen, die auch eine hohe durchschnittliche Aktivität aufweisen. In der zweiten Spalte ist die Klasse der kurzkettigen aliphatischen Aldehyde aufgezeigt. Man erkennt deutlich, wie die durchschnittliche Aktivität im Vergleich zu den langkettigen aus Spalte 3 in caudale Richtung verschoben ist. Diese Verschiebung zeigt sich genauso deutlich auch auf den DIs.

Das DI der polyzyklischen Moleküle (Abb. 5.7, vierte Spalte) zeigt erneut viel Korrespondenz mit der Durchschnittsaktivität, allerdings gibt es auch hier zwei prominente Regionen, die davon unabhängig scheinen: Eine an caudal-lateraler Position,



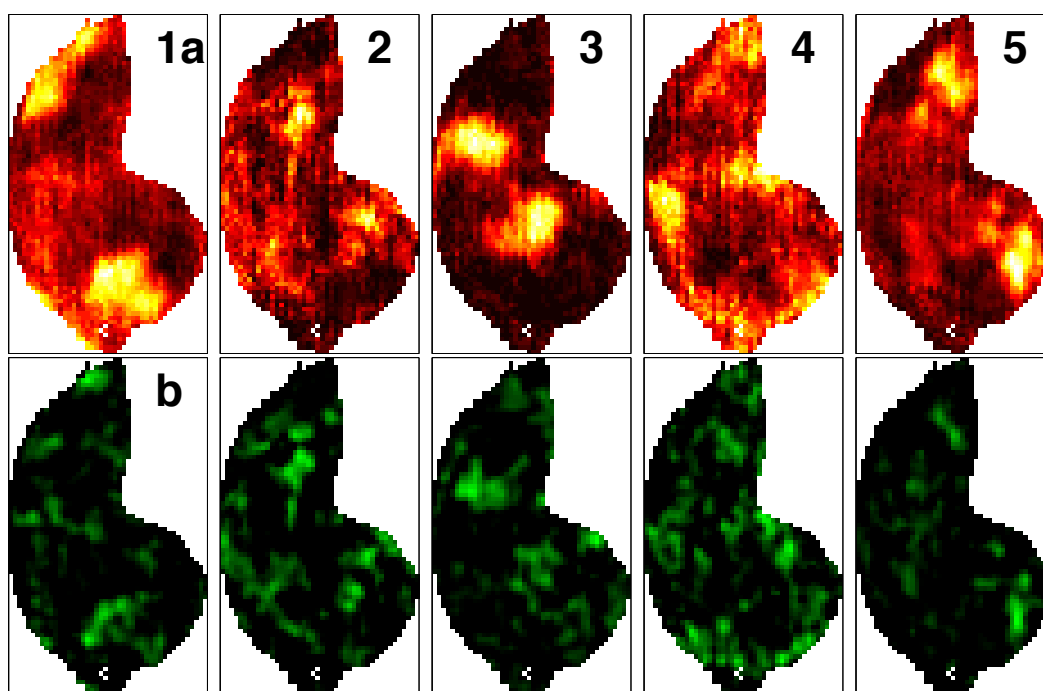
**Abbildung 5.6:** (1) Mittelwertbild  $c_{alc}$  für die Klasse der Alkohole. Je heller die Regionen auf dem Bild sind, desto mehr Alkohole induzieren dort neuralgische Aktivität. Eingekreist sind vier Regionen, die auch auf Bild (2) gekennzeichnet sind. Hier ist das Decision Images  $d_{alc}$  für die Alkohole gezeigt. Je heller ein Pixel, desto prädiktiver ist die Aktivität im zugehörigen Glomerulus für die Klassenzugehörigkeit. Nicht jede Region mit hohem Mittelwert ist stark prädiktiv und nicht jede prädiktive Region ist einer hohen mittleren Aktivität zuzuordnen.

eine medial. Schließlich bleibt noch die Gruppe der kurzkettigen aliphatischen Ester in der fünften Spalte. Die kurzkettigen Ester zeigen ein sehr kondensiertes Decision Image mit nur wenigen scharfen prädiktiven Regionen für diese Klasse. Dies könnte darauf hinweisen, dass diese Substanzklasse in der Tat eine sehr spezifische Rezeptorantwort erzeugt, was auch durch die entsprechenden Peaks im zugehörigen Durchschnittsbild unterstrichen wird.

## 5.2.2 Das Riechalphabet und natürliche Riechwörter

Die Geruchswahrnehmung, also die sensorische Wahrnehmung von Geruchsstoffen durch das nasale System, ist ein komplexer Prozess. Es sind tausende von Rezeptoren beteiligt, die es erst in Kombination möglich machen, zehntausende, teilweise bis auf ein Atom identische Geruchsstoffe zu unterscheiden.

Betrachten wir jeden monomolekularen Geruchsstoff als ein spezifisches Uptake-Bild  $\mathbf{x}_i$ , so wird er beschrieben durch eine Mischung aus den unterschiedlichsten Rezeptoren bzw. Aktivität in den zugehörigen Glomeruli (siehe z.B. Abb. 5.4). Die Variation an Rezeptoren, die durch eine Substanz stimuliert werden, ist maßgeblich



**Abbildung 5.7:** Gegenüberstellung von Mittelwert  $c_\mu$  und Decision Image (DI)  $d_\mu$  für einige chemische Eigenschaften. Die erste Reihe zeigt die durchschnittliche Aktivität  $c_\mu$  der Klassen (a), die zweite Reihe die Decision Images  $d_\mu$  (b). Das DI ist typischerweise flächenkleiner und zeigt auf, welche Mittelwertpeaks wirklich prädiktiv für die Klasse sind. Es gibt auch signifikante Regionen außerhalb der hohen Mittelwerte. Chemische Klassen von links nach rechts: (1) Aromate mit O-Substituent, (2) kurzkettige ( $\leq 8C$ ) und (3) langkettige ( $\geq 9C$ ) aliphatische Aldehyde, (4) Polyzyklische Substanzen und (5) kurzkettige aliphatische Ester.

durch seine biochemischen Eigenschaften bestimmt. Und – hier kommt dann die Kombinatorik ins Spiel – jeder der Rezeptoren wird typischerweise durch mehr als eine Substanz stimuliert. Zusammengefasst heißt das also: Jeder Geruchsstoff ist eindeutig beschrieben durch eine spezifische Mischung von Glomeruli-Aktivität.

Um diese Kombinatorik zu verstehen und vielleicht sogar zu entschlüsseln, muss dieser Kode irgendwie wieder *zerlegt* werden in seine fundamentalen Bestandteile. Informationstheoretisch betrachtet ist jede Substanz also repräsentiert durch ein Kodewort – in diesem Fall ein spezifisches Muster an aktiven Glomeruli – und wir suchen nun nach dem *Alphabet*, aus dem dieses Kodewort zusammengesetzt wird. Haben wir dieses Alphabet verstanden, könnte man im nächsten Schritt sogar versuchen, die *Grammatik* zu verstehen, um selbst Wörter einer bestimmten Bedeutung zu erzeugen.

Wie verhalten sich nun aber die Decision Images der chemischen Klassen zu dieser allgemeinen Betrachtung der Kombinatorik?

---

Betrachten wir z.B. das DI  $\mathbf{d}_{ester}$  für die Klasse aller Ester. Dieses Muster hat hohe Werte für alle Regionen des Bulbus, die besonders geeignet sind, um Ester von Nicht-Estern zu unterscheiden. Dies entspricht aber weder einem einzelnen Wort noch einem der gesuchten Buchstaben. Die lokalen Maxima des DI stellen eher eine Art Kandidatenmenge derjenigen Buchstaben bzw. Regionen dar, mit deren Aktivierung in einer geeigneten Kombination typische Ester-Wörter erzeugt werden können.

Gesucht sind aber zusammenhängende Regionen oder Module, welche eine klassenspezifische Bedeutung haben. Also Punkte im Bulbus, die man stimulieren müsste, um z.B. eine fruchtig-Wahrnehmung zu induzieren<sup>3</sup>.

Formal sei eine Menge  $\mathcal{A}_\mu$  definiert als

$$\mathcal{A}_\mu = \{\alpha_1, \dots, \alpha_\xi\}$$

die Menge von  $\xi$  Modulen  $\alpha_j$  auf dem Decision Image  $\mathbf{d}_\mu$  der Klasse  $\mu$ . Wir vernachlässigen an dieser Stelle erst einmal komplett die Ermittlung dieser Module  $\alpha_j$  auf den Decision Images. Dies wird allerdings später in Kapitel 6 noch einmal thematisiert.

Ferner nehmen wir an, es gibt eine Obermenge  $\mathcal{A}$ , die die Menge *aller* solcher Module  $\alpha$  beinhaltet, die im Riechkode verwendet werden, also das komplette *Riechalphabet*. Es gilt somit  $\mathcal{A}_\mu \subset \mathcal{A}$ .

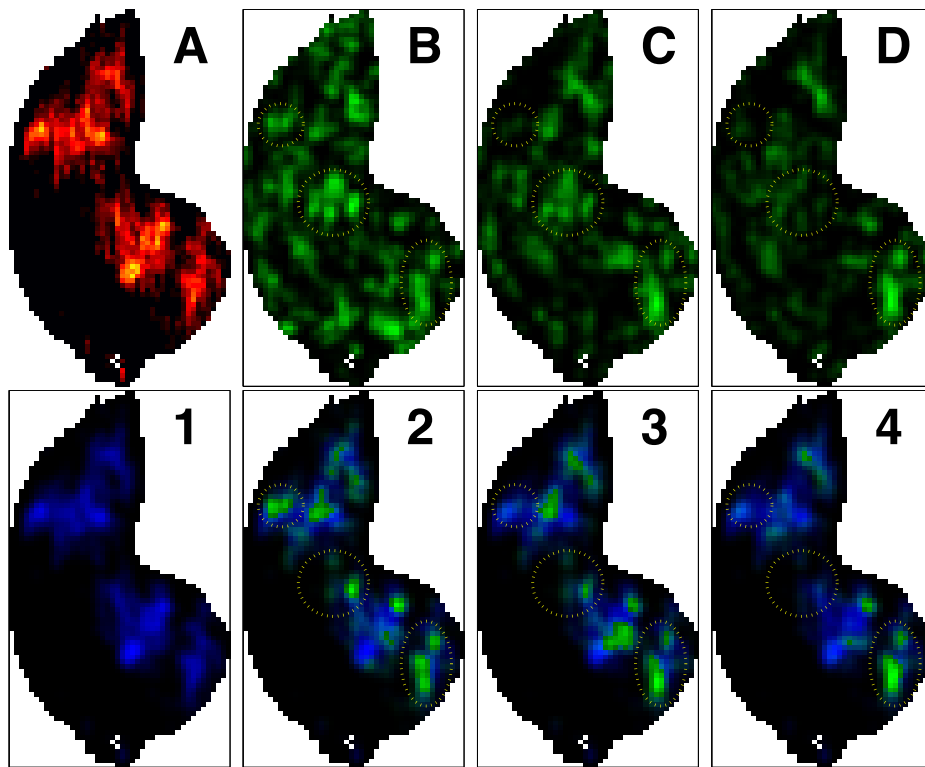
Es handelt sich bei der Menge  $\mathcal{A}_\mu$  um genau die Menge von *Buchstaben*  $\alpha_j$ , aus der die für die Klasse  $\mu$  spezifischen *Riechwörter* zusammengesetzt werden. Ein solches Wort der Klasse  $\mu$  besteht demnach aus einer Kombination von Buchstaben  $\alpha_j \in \mathcal{A}_\mu$ , was genau unserer Grundannahme eines kombinatorischen Codes entspricht. Es reicht demnach vermutlich nicht aus, nur ein einziges Modul (z.B.  $\alpha_1$ ) zu reizen, um eine spezifische Wahrnehmung zu erzeugen, denn  $\alpha_1$  ist vermutlich auch Teil eines anderen Klassen-Alphabets  $\mathcal{A}_j$ .

Um eine solche Kandidatenmenge genauer spezifizieren zu können, bräuchten wir „natürliche“ *Riechwörter*, an denen wir sehen können, welche Buchstaben  $\alpha_j \in \mathcal{A}_\mu$  für eine Substanz aus Klasse  $\mu$  typischerweise kombiniert werden. Solche Wörter stehen uns erfreulicherweise bereits zur Verfügung, denn alle vorhandenen Uptake-Bilder  $\mathbf{x}_{i \in I_\mu}$  der Klasse  $\mu$  entsprechen jedes für sich einem *Riechwort*.

Betrachten wir also ein Uptake-Bild  $\mathbf{x}_i$  mit  $i \in I_\mu$ , also ein Bild aus Klasse  $\mu$ . Dann entspricht die Aktivität von  $\mathbf{x}_i$  im Grunde einer *Maske* für das Decision Image  $\mathbf{d}_\mu$ . Denn hat  $\mathbf{x}_i$  z.B. genau unter den Modulen  $\alpha_1, \alpha_5, \alpha_{12} \in \mathcal{A}_\mu$  Aktivität, so ist das von  $\mathbf{x}_i$  ausgelöste Wort bzgl. Klasse  $\mu$  ganz offensichtlich  $\alpha_1\alpha_5\alpha_{12}$ .

---

<sup>3</sup>Ester haben nahezu ausnahmslos eine fruchtige Note



**Abbildung 5.8:** Decision Images für Amyl Acetate: Oben links ist das Uptake-Bild  $\mathbf{x}_a$  für Amyl Acetate dargestellt (A). Rechts davon sind drei Decision Images  $\mathbf{d}_\mu$  für die Amyl Acetate zugeordneten chemischen Eigenschaften dargestellt: (v.l.n.r.) Ester (B), Aliphatischer Ester (C) und Kurzketziger Alip. Ester (D). In der zweiten Zeile links ist erneut  $\mathbf{x}_a$  gezeigt, allerdings nur im Blaukanal (1). Auf den Bildern links daneben sind zusätzlich im Grünkanal die Spuren der drei DIs  $\mathbf{s}_{\mu,a}$  für die Klassen aus Zeile 1 zu sehen (2–4).

Für diese *Maskierung* definieren wir  $\mathbf{s}_{\mu,i} \in \mathbb{R}^n$  als die komponentenweise Multiplikation eines Decision Image  $\mathbf{d}_\mu$  und einem Uptake-Bild  $\mathbf{x}_i$  als

$$\mathbf{s}_{\mu,i} = \begin{pmatrix} d_\mu(1) \cdot x_i(1) \\ \vdots \\ d_\mu(n) \cdot x_i(n) \end{pmatrix}$$

Das *Riechwort*  $\mathbf{s}_{\mu,i}$  der Klasse  $\mu$  – denn um nichts anderes handelt es sich – werden wir im weiteren Verlauf dieser Arbeit die *Spur* der Klasse  $\mu$  auf dem Bild  $\mathbf{x}_i$  nennen.

Es sei an dieser Stelle bereits angemerkt, dass wir zwar leicht zwei Spuren  $\mathbf{s}_{\mu,i}$  und  $\mathbf{s}_{\nu,i}$ , also bzgl. desselben Uptake-Bildes  $\mathbf{x}_i$  vergleichen können, die Spuren  $\mathbf{s}_{\mu,i}$  und  $\mathbf{s}_{\mu,j}$  eines Decision Images  $\mathbf{d}_\mu$  auf zwei unterschiedlichen Bildern  $\mathbf{x}_i$  und  $\mathbf{x}_j$  hingegen

---

sind nur sehr schwer zu vergleichen. Die Bilder zeigen ihre Aktivität ja potentiell auf komplett disjunkten Regionen.

Es ist also möglich, für jede der Klassen  $\mu$  das potentielle Klassenalphabet  $\mathcal{A}_\mu$ , das durch das Decision Image  $\mathbf{d}_\mu$  repräsentiert ist, noch weiter zu differenzieren. Dazu wird mit Hilfe der Spuren  $\mathbf{s}_{\mu,\xi}$  von  $\mathbf{d}_\mu$  auf den Klassenbildern  $\mathbf{x}_{\xi \in I_\mu}$  das DI in verschiedene klassenspezifische Riechwörter zerlegt.

Diese Betrachtung beschränkt sich dann grundsätzlich auch nicht nur auf ein einzelnes DI, wir können für ein gegebenes Bild  $\mathbf{x}_i$  auch versuchen, seine Aktivitäten komplett in alle seine substanzspezifischen Klassenspuren zu zerlegen, d.h. wir betrachten alle Spuren  $\{\mathbf{s}_{\zeta,i} \mid i \in I_\zeta\}$ .

Im Falle von z.B. *Amyl Acetate* betrachten wir also nicht nur die Spur für die *kurzkettigen aliphatischen Ester*-Eigenschaft, sondern ebenfalls die Spuren der Eigenschaften *aliphatischer Ester* und *Ester*. Dies ist in Abbildung 5.8 gezeigt. Ausgehend von dem Uptake-Bild  $\mathbf{x}_a$  der Substanz *Amyl Acetate* haben wir versucht, etwas über die Bedeutung der durch den Stoff stimulierten Regionen mit Hilfe der zugehörigen DIs abzuleiten. Zu sehen sind neben  $\mathbf{x}_a$  in der ersten Reihe die DIs für die drei zentralen chemischen Eigenschaften von *Amyl Acetate*: *Ester* (B), *aliphatischer Ester* (C) und *kurzkettige aliphatische Ester* (D).

In der zweiten Reihe der Abb. 5.8, Bild 1, ist erneut  $\mathbf{x}_a$  abgebildet. Dies entspricht Bild A, allerdings sind die z-Scores dieses mal nur im Blaukanal gezeigt. Auf diesen Uptake zeichnen wir nun in den Bildern 2–4 die Spuren  $\mathbf{s}_{\mu,a}$  der DIs aus den Bildern B–D in den Grünkanal. Dort wo keine Aktivität in  $\mathbf{x}_a$  ist, wird also auch im resultierenden Bild eine 0 stehen, die höchsten Werte in  $\mathbf{s}_{\mu,a}$  entstehen an den Stellen, an denen hohe Aktivität in  $\mathbf{x}_a$  mit hohen Werten in  $\mathbf{d}_\mu$  für die drei Klassen  $\mu$  korreliert ist.

Aus dieser Zerlegung der durch *Amyl Acetate* induzierten Aktivität können wir einiges über *Amyl Acetate* wie auch über die Klasse der Ester lernen. Betrachten wir die Spuren aus Abbildung 5.8 noch einmal genauer. Die drei Klassen hängen ja insofern voneinander ab, dass jeder kurzkettiger aliphatischer Ester ebenfalls ein aliphatischer Ester ist und jeder aliphatische Ester ist auch ein Ester.

In Abbildung 5.8 sind drei markante Regionen durch Kreise auf den Bildern B–D und 2–4 gekennzeichnet. Im untersten Kreis ist deutlich zu erkennen, wie alle drei DIs dort einen Spuranteil auf  $\mathbf{x}_a$  haben. Dies ist nicht erstaunlich, denn ist eine Region signifikant und spezifisch für kurzkettige aliphatische Ester, so ist diese Region natürlich auch ein guter Indikator für die beiden übergeordneten Klasseneigenschaften. Denn wenn dort Rezeptoren konvergieren, die spezifisch nur kurzkettige aliphatische Ester binden, so ist Aktivität dort natürlich auch ein hinreichendes Kriterium für die Präsenz von aliphatischen und einfachen Estern. Bei dieser Region handelt es sich also potentiell um einen Buchstaben  $\alpha$  des Riechalphabets,

welcher sehr spezifisch für die ganze Familie der Ester kodiert.

Die Region im ventralen Bereich, also in der Mitte der Abbildung, hingegen zeigt für die kurzkettigen Ester keine hohe Signifikanz an, sehr wohl aber für die beiden anderen Ester-Eigenschaften. Dies lässt vermuten, dass hier zwar spezifisch aliphatische Ester gebunden werden, die zugehörigen Rezeptoren im OE allerdings auch langkettige aliphatische Ester binden. Die dritte Region zeigt entsprechend nur für Ester Signifikanz auf dem Spur-Bild. Hier münden also anscheinend Ester-sensitive Rezeptoren, die jedoch keinerlei Spezifität bzgl. *aliphatischer* Ester besitzen. Interessanterweise lassen experimentelle Beobachtungen eine eher ventrale Lokalisierung der Glomeruli sensitiver ORNs für langkettige aliphatische Ester im Vergleich zu den kurzkettigen vermuten [36, 42]. Dies ist damit konsistent mit den hier gemachten Beobachtungen.

Experimentell ist es unmöglich, systematisch die Sensitivität von tausenden Substanzen auf aufwendig aus dem OE extrahierten Neuronen zu bestimmen. Mit Hilfe von Decision Images ist es nun aber offensichtlich möglich, ziemlich scharfe Hypothesenmengen zu definieren, aus denen ein Atlas über Konvergenzpunkte sämtlicher dokumentierter Geruchsrezeptoren erstellt werden könnte.

Auch wenn dies durch den hier vorgestellten Ansatz theoretisch ermöglicht wird, handelt es sich dennoch um ein sehr umfangreiches Projekt, welches vermutlich weitere Jahre intensiver Forschung nach sich ziehen wird. Im Rahmen dieser Arbeit sind diese Rezeptorlokalisierungen aber eher ein willkommenes Nebenprodukt, stellen die Lokalisierung von Rezeptoren und der Vergleich mit vorhandenen Erkenntnissen doch für uns eine Möglichkeit (wenn nicht sogar *die einzige* Möglichkeit) dar, den hier entwickelten Analyserahmen bzgl. der Geruchsqualität auf Plausibilität zu testen.

### 5.2.3 Geruchsqualitäten und ihre Repräsentanz im Bulbus

Die Uptake-Datenmatrix  $\mathbf{X}$  stellen Riechkolben-Aktivitäten von Ratten dar. Wenn wir jetzt versuchen, die von Menschen beschriebenen Geruchsqualitäten mit den Bildern von Rattenriechkolben zu erklären, ist eine gewisse Skepsis angebracht, aber auch viel Neugierde. Viele neurophysiologische Beobachtungen bei Ratten, z.B. des visuellen Systems, ließen auch viele Schlüsse auf die Funktion des entsprechenden menschlichen Systems zu. Genetisch betrachtet besitzen Ratten zwar ein wesentlich breiteres Repertoire an Rezeptoren, die meisten dieser Rezeptoren findet man aber noch im menschlichen Genom, jedoch nur noch als Pseudogene [23, 61, 101].

Wie bereits erwähnt, erwarten wir, dass die Information, aus der auch der Mensch die wahrgenommene Geruchsqualität ableitet, in dem hier gemessenen Signal enthalten ist. Daher haben wir versucht, unsere umfangreiche Geruchsdatenbank (siehe



---

auch Kapitel 4) zu verwenden, um etwas über die Zusammensetzung der Glomeruli-Muster zu lernen. Das Procedere ist dabei grundsätzlich identisch zu der Analyse der chemischen Eigenschaften: Alle Bilder von Substanzen, die von Menschen typischerweise mit Hilfe eines bestimmten Geruchsbezeichners beschrieben werden, gehören zu einer Klasse, alle anderen Substanzen zu der entsprechenden Nicht-Klasse. Basierend auf einer solchen Klasseneinteilung ist es nun entsprechend möglich, Decision Images zu berechnen, die die Uptake-Bilder der Klasse von den übrigen Bildern möglichst zuverlässig trennen. Derart erhalten wir dann auch Decision Images, die uns z.B. für die Geruchsqualität *bananig* die prädiktivsten Regionen beschreiben.

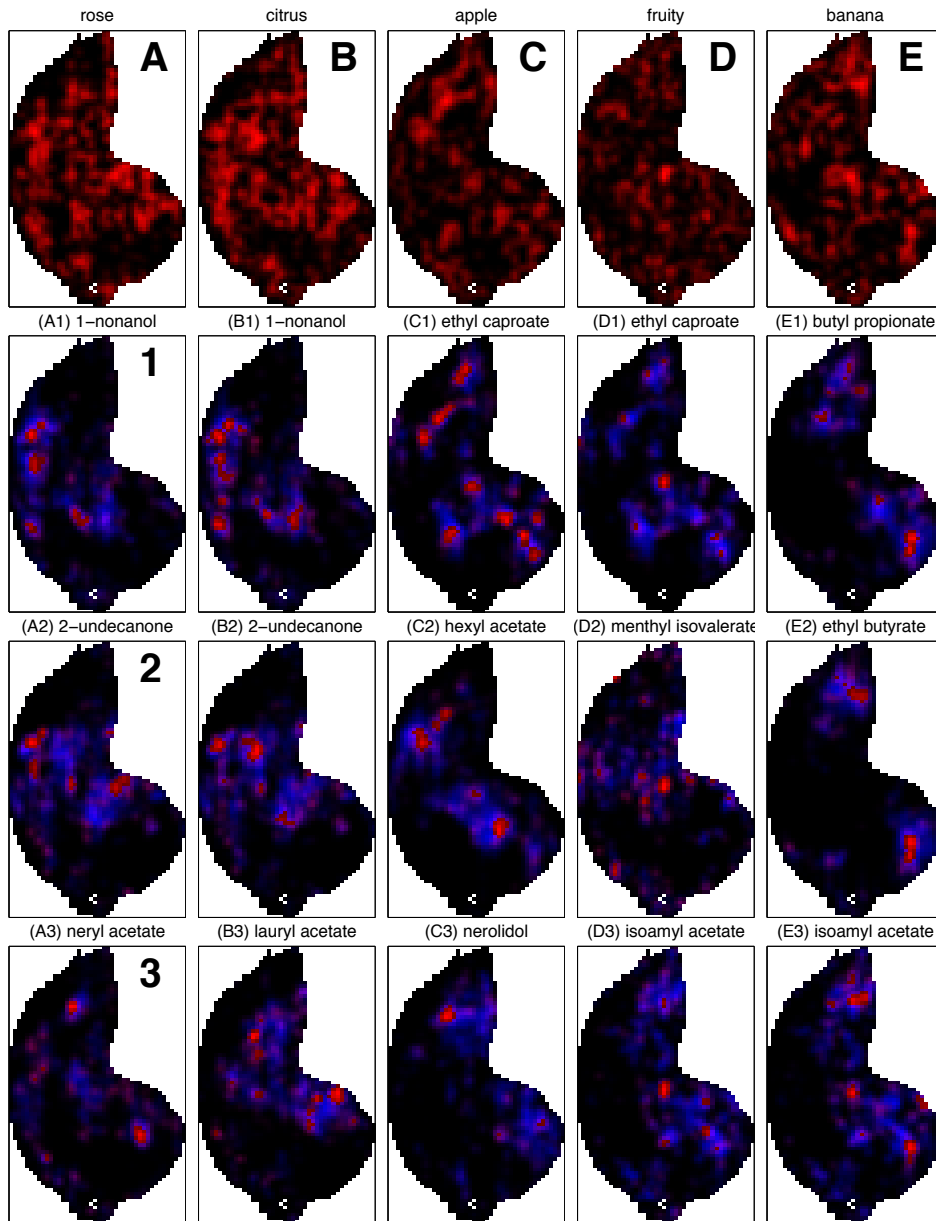
Analog zu Abbildung 5.8 zeigt Abbildung 5.9 in der ersten Zeile die Decision Images für die fünf ausgewählten Geruchsqualitäten: *rose*, *citrus*, *apple*, *fruity* und *banana*. Auch hier sind nur die positiven Werte des Decision Images dargestellt, da wir hier nur daran interessiert sind, zu begreifen, welche Regionen Indikatoren für eine Klasse sind. Für jede dieser Qualitäten  $\nu$  ist in den darunterliegenden Spalten im Blaukanal jeweils das positive Uptake-Bild  $\mathbf{x}_i$  für eine von drei klassenzugehörigen Substanzen gezeigt, überlagert von der Spur des Decision Images  $\mathbf{s}_{\nu, x_i}^O$  auf dem jeweiligen Uptake im Rotkanal.

Für einige Geruchsqualitäten (z.B. *rose* und *citrus*, *apple* bzw. *banana* und *fruity*) sind hierbei gleiche Substanzen gewählt, um die Spuren der Klassen auf dem gleichen substanzspezifischen Unterraum vergleichen zu können. Hier versuchen wir, ebenso wie bei den chemischen Eigenschaften, eine Deutung des Uptakes  $\mathbf{x}_i$  mittels der Spuren  $\mathbf{s}_{\nu, x_i}^O$ . Dieses Mal beziehen sich die Spuren jedoch auf die Geruchsqualität  $\nu$  der jeweiligen Substanz.

### Die Geruchsqualitäten *rose* vs. *citrus*

In Abbildung 5.9 sind für die beiden Qualitäten *rose* (A) und *citrus* (B) zwei Substanzen aufgeführt, die jeweils beide Geruchsassoziationen auslösen können. Solche Beispiele sind besonders interessant, da wir für dasselbe Aktivierungsmuster die kodierenden Anteile für die beiden Qualitäten vergleichen können.

Die beiden betrachteten Qualitäten erscheinen erst einmal grundsätzlich verschieden, wir stellen jedoch fest, dass viele rosige Substanzen auch mit einer *citrus*-Note behaftet scheinen, zumindest empfinden viele Personen sowohl *rose* als auch *citrus* als zutreffende Bezeichner, wenn sie z.B. an den Substanzen *1-Nonanol* (A1+B1) oder *2-Undecanone* (A2+B2) riechen. Beide Substanzen haben übereinstimmende Regionen in der Spur beider Qualitäten, interessanterweise auch an denselben Stellen, nämlich an der rostral-lateralen Spitze des Bulbus. Auch andere Ähnlichkeiten sind augenscheinlich, jedoch scheinen die Spuren zwischen *rose* und *citrus* auch nie identisch zu sein. Ein wenig in caudale Richtung verschoben z.B. ist für *2-Undecanone* ein Peak in der Spur von *citrus* zu erkennen an einer Stelle, die für *rose* keine Entscheidungsrelevanz hat. Eine entsprechende Region gibt es am



**Abbildung 5.9:** Decision Images für Geruchsqualitäten. Die erste Reihe, (A-E), zeigt die DIs  $d_{\nu}^O$  für die Geruchsqualitäten: *rose* (A), *citrus* (B), *apple* (C), *fruity* (D) und *banana* (E). Unter jeder der Qualitäten finden sich drei Beispiele für klassenzugehörige Substanzen  $x_j$  und der Spur der jeweiligen Decision Images  $d_{\nu, x_i}^O$ . Die Aktivität  $x_j$  ist dabei im Blaukanal angezeigt und im Rotkanal die zugehörige Spur des DIs der Spalte. Die gezeigten Beispiele (jeweils von oben nach unten) für *rose* sind *1-Nonanol*, *2-Undecanone*, und *Neryl Acetate*. Für *citrus* sind es *1-Nonanol*, *2-Undecanone* und *Lauryl Acetate*. *Ethyl Caproate*, *Hexyl Acetate* und *Nerolidol* für *apple*. *Ethyl Caproate*, *Menthyl Isovalerate* und *Isoamyl Acetate* für *fruity*. *Butyl Propionate*, *Ethyl Butyrate* und *Isoamyl Acetate* für *banana*.

---

caudal-ventralen Rand der Glomeruli-Karte für *rose*.

Die beiden letzten Beispiele für die beiden Qualitäten sind wiederum interessant, da sowohl *Neryl Acetate* (A3) wie auch *Lauryl Acetate* (B3) sehr unterschiedliche Aktivierungsprofile zeigen im Vergleich zu ihren anderen Klassenvertretern. *Neryl Acetate* hat in seiner *rose*-Spur zwei Peaks (lateral und medial) in Regionen, die durch die anderen beiden Substanzen überhaupt nicht stimuliert werden. Umgekehrt zeigt *Neryl Acetate* in den rostralen Kernregionen der anderen *rose*-Spuren wenig bis gar keine Aktivität.

### Die Geruchsqualität *fruity*

Innerhalb einiger Klassen zeigt die gezeichnete Spur eine erstaunliche Konsistenz für die dargestellten Substanzen, sogar für Duftstoffe, die deutlich unterschiedliche Uptake-Bilder aufweisen. Die fruchtigen (*fruity*) Substanzen etwa weisen alle in der zentralen ventralen Position hohe Werte in der Spur des Decision Images für die *fruity*-Qualität auf. Betrachtet man *Ethyl Caproate* (D1), so fällt einem auf, dass die Spur von *apple* auf derselben Substanz (C1) ebenfalls diesen ventralen Peak aufweist. Gleiches gilt für *Isoamyl Acetate* (D3), eine Substanz, die sowohl zur Klasse *fruity* wie auch zur Klasse *banana* (E3) gehört. Die Unterschiede zwischen den Spuren D1 und C1 bzw. zwischen D3 und E3 finden sich eher an einem umfangreicheren bzw. spezifischeren Muster für *apple* und *banana* im Vergleich zu *fruity*.

Sind also die Geruchsqualitäten *apple* und *banana* präzisere Beschreibungen einer *fruity*-Qualität? Das ist schwer zu sagen, denn es gibt auch Vertreter der beiden Fruchtklassen, die diese Charakteristik nicht zeigen (siehe z.B. Nerolidol (C3) oder Ethyl Butyrate (E2)). Umgekehrt jedoch fällt auf, dass zwar alle *apple*- und *banana*-Vertreter typischerweise auch als *fruity* beschrieben werden, jedoch ebenso konsistent keine der *citrus*-Substanzen (B1–B3) als *fruity* beschrieben wird und genau diesen *fruity*-Peak niemals aufweisen.

Es scheint also zumindest so, dass *citrus* erst einmal nichts mit dem Geruchsbeschreiber *fruity* zu tun hat, *apple* und *banana* aber auch sensorisch mit *fruity* korreliert sind. Evtl. sind die Rezeptoren, die Teil des typischen *banana* oder *apple*-Aktivitätsmusters sind, aus *fruity*-spezifischen Rezeptoren entstanden. Vielleicht deutet sich hier aber auch an, dass eine Ratte *apple* und *banana* als sehr ähnlich wahrnimmt, *citrus* aber wesentlich anders.

### Die Geruchsqualitäten *apple* vs. *banana*

Für beide Qualitäten gilt erst einmal ähnliches wie für *rose* und *citrus*: Wir finden für alle Beispiele gewisse Konsistenzen innerhalb der Klasse, jedoch ist es schwer – alleine durch die stark unterschiedlichen Aktivierungsbereiche – ein einheitliches Muster herauszuarbeiten. Vergleicht man die Decision Images der beiden Quali-

täten (C+E), so fällt einem auf, dass die signifikanten Regionen zwar räumlich beieinander, aber gleichzeitig auch nahezu disjunkt zu liegen scheinen.

Die Geruchsqualität *banana* zeigt noch am deutlichsten ein einheitliches Aktivierungsmuster für alle drei Beispiele. Die caudal-mediale Region, die durch alle Beispiele sehr prominent ist, ist uns bereits zuvor schon einmal über den Weg gelaufen: In Abbildung 5.8 haben wir sie als Konvergenzpunkt für Rezeptoren identifiziert, die vermutlich sensitiv für kurzkettige aliphatische Ester sind. Tatsächlich gehören die meisten der *banana*-Substanzen zu dieser Gruppe. Insofern wundert es nicht, dass einige der qualitätsspezifischen Regionen hier auch mit den chemischen Eigenschaften korrelieren.

An diesen wenigen Beispielen sei auch demonstriert, wie schwierig es ist, hier einen kombinatorischen Kode herauszuarbeiten. Es ist schwierig, überhaupt für eine Klasse von Substanzen eine allgemeingültige Systematik abzuleiten. Besonders gravierend ist hier zum einen, dass die Klassen eigentlich nur sinnvoll an gleichen Substanzbildern verglichen werden können, da wir sonst implizit Aussagen über zwei komplett unterschiedliche Subräume treffen. Weiterhin gibt es innerhalb der Klassen oftmals Substanzen, die bzgl. ihres Uptakes eine nahezu leere Schnittmenge besitzen, d.h. die zwei Substanzen teilen sich zwar eine Geruchsqualität, teilen sich aber faktisch kaum einen Pixel an positiver Uptake-Aktivierung (siehe z.B. in Abb. 5.9 die beiden *citrus*-Vertreter 1-Nonanol (B1) und Lauryl Acetate (B3) oder auch die beiden *apple*-Vertreter Hexyl Acetate (C2) und Nerolidol (C3)). Somit ist dann auch die Menge an möglichen klassenübergreifenden Kodes nahezu leer, da jeder Kode nur für jeweils eine dieser Substanzen Gültigkeit haben könnte.

Wir werden uns im nächsten Kapitel noch einmal explizit mit unterschiedlichen, auf dieser Datenlage überhaupt realistischen und vor allem nachweisbaren Kodierungsvarianten auseinandersetzen. Zuvor soll jedoch noch etwas zu der statistischen Signifikanz der hier eingeführten und bereits verwendeten Decision Images ausgeführt werden. Schließlich vollführen wir in den Betrachtungen einen weiten Spagat zwischen Aktivitätsprofilen vom Ratten-Bulbus und Geruchsbeschreibungen von Menschen. Haben die Beschreiber für die Ratten überhaupt eine Bedeutung?

### 5.2.4 Statische Signifikanz der Klassifikation mittels Decision Images

Im Allgemeinen werden die gegebenen Daten bei einer  $k$ -fachen Kreuzvalidierung in  $k$  Gruppen eingeteilt, um dann unter zufälligem Auslassen einer der Gruppen den Klassifikator zu trainieren und das Ergebnis an der weggelassenen Gruppe zu testen. Durch die mittlere Klassifikationsleistung des Klassifizierers auf den Testdaten wird ungefähr angegeben, wie gut der Klassifizierer mit unbekannten Daten umgehen

---

kann. Kreuzvalidierung wird besonders bei kleinen Datenmengen eingesetzt, wenn das Abspalten einer Testmenge zu einem nicht unerheblichen Datenverlust führt.

Wir werden im folgenden die Kreuzvalidierung nicht nur zu der Abschätzung der Abstraktionsfähigkeit verwenden. Für uns ist es weiterhin interessant, ob die von uns vorgeschlagenen Prototypen für die jeweiligen Klassen eine statistische Signifikanz besitzen oder ob sie sich von zufälligen Partitionierungen der Daten nicht unterscheiden lassen. Zu diesem Zweck haben wir auf den Uptake-Bildern eine 10-fach Kreuzvalidierung angewendet. Dies wurde nicht nur für die gegebenen Klasseneinteilungen, sondern auch für zufällige Labelungen derselben Datenmenge durchgeführt.

Wir haben in diesem konkreten Fall nicht nur sehr wenige Punkte – gemessen an der Dimensionalität des Datenraumes – es liegen in der Regel auch immer asymmetrische Klassenverhältnisse vor, d.h. typischerweise befinden sich weit weniger Elemente in der konkreten Klasse als in der verbleibenden Nicht-Klasse. Z.B. gibt es 37 Bilder von Alkoholen, in der Nicht-Klasse befinden sich dann jedoch entsprechend 174 Bilder<sup>4</sup>.

Um den Fehler trotzdem zu balancieren, verwenden wir im weiteren eine geschichtete oder auch *Stratified Cross-Validation*. Bei dieser Variante der Kreuzvalidierung berechnet sich der Gesamtfehler aus dem gewichteten Fehler aus FAR und FRR [10]. Eine Darstellung des hier verwendeten Algorithmus findet sich im Anhang A.4.

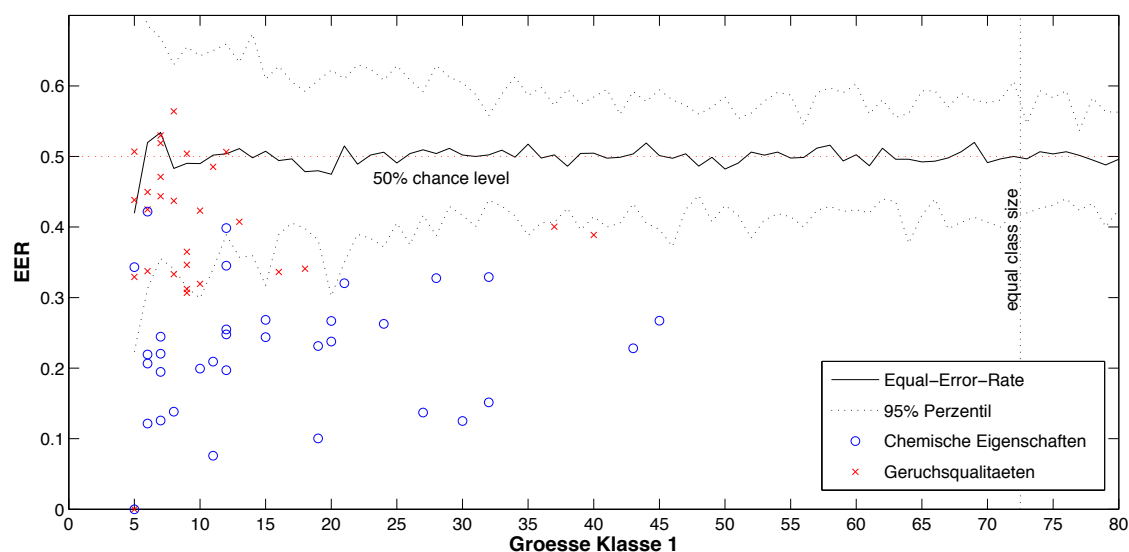
Mit Hilfe des DMO-Algorithmus haben wir für die Klassen der chemischen Eigenschaften als auch der Geruchsqualitäten klassenspezifische Gewichtsvektoren berechnet. D.h. unabhängig voneinander wurden für die 211 Uptake-Bilder in  $\mathbf{X}$  für alle vertretenen 36 chemischen Merkmale die zugehörigen Gewichtsvektoren  $\mathbf{w}_j^C$  berechnet, ebenso wie die Gewichtsvektoren  $\mathbf{w}_k^O$  für die 44 enthaltenen Geruchsqualitäten (siehe Abschnitt 5.1.3 für eine komplette Liste aller chemischen Eigenschaften und der verfügbaren Geruchsqualitäten).

Da wir in diesem Fall wesentlich weniger Punkte (211 Bilder) zur Verfügung haben als Dimensionen (2121 Pixel), gibt es für jede beliebige Labelung  $L$  der Daten (also für jede beliebige Einteilung der gegebenen Bilder in zwei Klassen) eine perfekte Lösung, d.h. eine Hyperebene, die die Trainingsdaten fehlerfrei trennen kann. Betrachten wir die chemischen Klassen wie auch die Geruchsklassen als Hypothesenmengen, so stellt sich die Frage, wie willkürlich diese Einteilung bezogen auf die Bilder in  $\mathbf{X}$  ist. Beschreiben die Gewichtsvektoren  $\mathbf{w}_j^C$  und  $\mathbf{w}_k^O$  auch real in den Daten existierende Cluster oder ist die Einteilung von einer zufälligen Einteilung der Daten nicht zu unterscheiden?

Man sollte an dieser Stelle erwähnen, dass es bei so wenigen Datenpunkten und

---

<sup>4</sup> $211 - 37 = 174$



**Abbildung 5.10:** Kreuzvalidierung auf 211 Uptake-Bildern. Die durchgezogene Linie zeigt die mittlere Fehlerrate ( $\sim 50\%$ ) für jeweils 50 Kreuzvalidierungsschritte mit zufällig gelabelten Daten, die gepunkteten Linien das zugehörige 5%- bzw. das 95%-Perzentil. Kreuze kennzeichnen die Equal Error Rate (EER) für die Geruchsqualitäten und Kreise die EER der chemischen Eigenschaften.

speziell bei nicht balancierten Klassen nicht einfach ist, „echte“ Cluster von zufälligen zu unterscheiden. Mit anderen Worten, ist eine Konfiguration statistisch nicht signifikant, heißt das nicht automatisch, dass sie auch ohne reale Bedeutung ist. Im Gegenteil, es kann sogar gezeigt werden, dass für bestimmte Konstellationen keine metrik-basierte Klassifikation mehr sinnvoll funktionieren kann [49].

Wir hoffen aber natürlich dennoch, dass es uns gelingt, die statistische Signifikanz der Klassen-Entscheidungsbilder – also der Decision Images – zu zeigen: Welche sind nachweisbar nicht zufällig? Um dies zu testen, haben wir uns für jede Klassengröße die Streuung bei der *Stratified Cross-Validation* für jeweils 50 zufällige Partitionierungen angeschaut. Durch Vergleich haben wir dann für die vermeintlich echten Klassen (chemische Eigenschaften und Geruchsqualitäten) deren „Nicht-Zufälligkeit“ abgeschätzt.

Abbildung 5.10 zeigt, wie sich die Gewichtsvektoren für die chemischen Eigenschaften (Kreise) und für die Geruchsqualitäten (Kreuze) relativ zu dem Zufallskorridor (strichpunktierte Linien) verhalten. Man erkennt, dass die meisten der chemischen Klassen – wie zu erwarten war – deutlich „nicht-zufällig“ sind. Hier liegt ganz offensichtlich eine direkte Korrelation zwischen den Bilddaten und den Klassenbeschreibungen vor. Die Geruchsqualitäten hingegen scheinen sich eher innerhalb des Zufallskorridors zu bewegen, hier ist kaum eine Signifikanz zu erkennen.

Es ist nicht überraschend, dass die chemischen Klassen hier so eindeutig signifikante

---

Fehlerraten ergeben, denn es gibt eine ganze Reihe von experimentellen Arbeiten, die alle erwarten lassen, dass die räumliche Anordnung der Glomeruli tatsächlich entlang der chemischen Eigenschaften organisiert und somit diese Klassen natürlich auch in unseren Daten eine Entsprechung in räumlichen Clustern haben [34–43, 88, 93].

Für die Geruchsqualitäten hingegen verwundert es schon, dass nur so wenige Klassen außerhalb der zufälligen Labelung liegen. Hat die Geruchsqualität, beschrieben durch Menschen, etwa gar nichts mit den Rattenaktivitäten zu tun? Leider kann im Rahmen dieser Arbeit diese Frage nicht endgültig beantwortet werden. Es wäre spektakulär gewesen, hier eine Signifikanz nachweisen zu können, andererseits wäre es auch nicht zu erwarten gewesen, denn es gibt keine Rezeptorfamilie mit Rezeptoren, die geruchsspezifisch Substanzen binden. Und die Rezeptororganisation gibt auch die räumliche Struktur im Riechkolben vor.

Angenommen, die Geruchsqualitäten haben hier doch eine Bedeutung, scheinen die zugehörigen Merkmale zumindest schwieriger aus den gegebenen Daten extrahierbar zu sein als dies der Fall ist für die chemischen Eigenschaften. Einen letzten Versuch, hier doch noch einen Schritt weiter zu kommen, werde ich im nächsten und letzten Kapitel dieser Arbeit beschreiben. Dabei sollte auch der Zusammenhang zwischen OFD und OC etwas transparenter werden.

## 5.3 Diskussion

Das Lernen von Decision Images stellt eine interessante Möglichkeit dar, um die vorhandenen Glomeruli-Aktivitäten wieder zu entmischen und in den Kontext einer gegebenen Klasse zu stellen. Wir können für den weitgehend verstandenen Zusammenhang zwischen Sensoren und ihrer Organisation entlang der chemischen Eigenschaften signifikante Ergebnisse erhalten und erstellen Vorhersagen z.B. von Geruchsneuronen-Konvergenzpunkten.

Wir stellen einen Analyserahmen vor, der es zum ersten Mal möglich macht, etwas über die Organisation des kombinatorischen Kodes der Gerüche herauszuarbeiten. Mittels der Decision Images und den zugehörigen Spuren auf den Uptake-Bildern können wir einem Großteil dieser Muster eine funktionelle Zuordnung geben. Dies eröffnet komplett neue Möglichkeiten bei der Kartografierung des Riechkolbens.

Man kann auch darüber spekulieren, ob das Lernen der Decision Images nicht sogar entfernt etwas mit dem Lernen von räumlichen Aktivierungsmustern einer Ratte zu tun haben könnte. Jede Substanz, der eine Ratte ausgesetzt wird, löst ein spezifisches Eingabemuster an Aktivitäten im Bulbus aus. Und basierend auf diesem Signal muss das Tier dann auf höherer Ebene entscheiden, wie es auf den Reiz rea-

gieren soll. So ein Szenario ergibt sich zum Beispiel im Rahmen von Habituationsexperimenten mit Ratten, in denen die Wahrnehmungsähnlichkeit zwischen einzelnen Substanzen abgeschätzt werden soll [16, 54]. Diese Aufgabenstellung ist in ihrer Formulierung erstaunlich ähnlich dem Problem, welches wir in diesem Kapitel mittels der Decision Images versucht haben, möglichst effektiv zu lösen.

Das Gehirn wird vermutlich einen ähnlichen Prozess durchlaufen, wenn es durch sukzessiv gemachte Erfahrung lernt, den Geruch verschiedener Apfelsorten jeweils auch mit einem „typischen“ Apfelgeruch zu verbinden. Intuitiv werden diese Gewichtungen dann natürlich unterschiedlich sein, betrachten wir erneut die Gegenläufigkeit der Prozesse der Odorant Fine Discrimination (OFD) und des Odor Clustering (OC).

Die resultierenden Unterschiede zwischen OFD und OC bezogen auf die verwendete Verschaltung der Bulbus-Information sollte unbedingt Thema weiterführender Untersuchungen sein; in Kapitel 6 werden wir grundlegende Kombinatoriken bezogen auf ihre Plausibilität für die gegebenen Daten untersuchen.



# 6 Kodierungsmodelle für Aktivierungsmuster im Bulbus

## 6.1 Einleitung

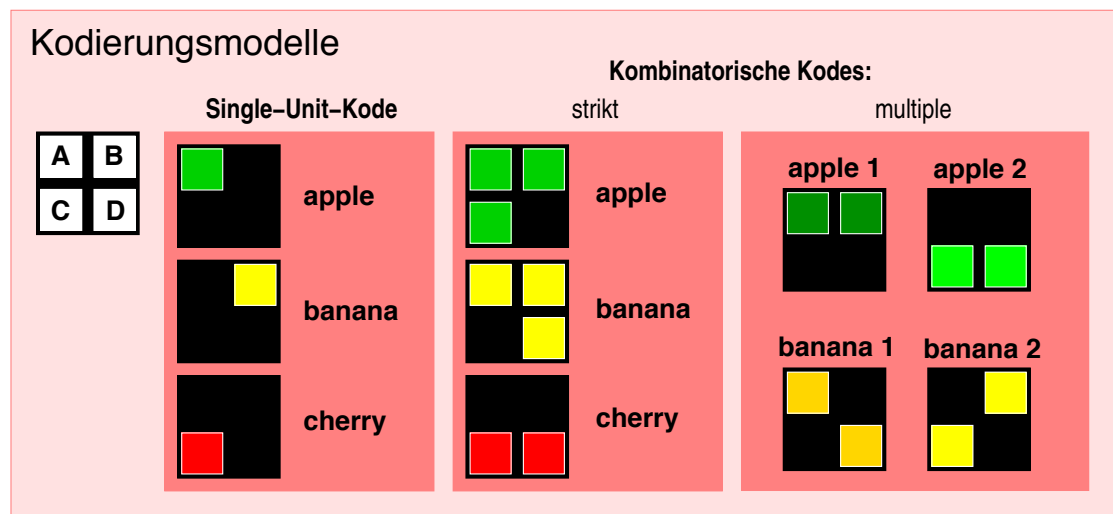
Betrachtet man das Riechsystem und speziell die Organisation des Riechkolbens, so handelt es sich dabei um ein hochkomplexes, leistungsfähiges chemosensorisches Wahrnehmungsorgan. Eine winzige chemische Substanz ist in der Lage, eine ganze Signalkaskade auszulösen, und bereits kleinste Spuren eines Duftstoffes können durch die Nase wahrgenommen werden.

In Kapitel 4 haben wir gezeigt, wie vielseitig auch die Beschreibungen ausfallen, die Personen verwenden, um ihre Geruchswahrnehmungen zu charakterisieren. Immerhin handelt es sich beim Riechen um einen Prozess, in dem mehrere hundert Riechrezeptoren involviert sind. Jeder Duftstoff ist in der Lage, mehrere Rezeptorneurone zu reizen, und jeder Rezeptor hat typischerweise Bindungsaffinitäten zu mehreren Substanzen. Wir haben festgestellt, dass der zugehörige Wahrnehmungsraum in der Tat eine Komplexität von etwa 32 Dimensionen besitzt (siehe dazu Kapitel 4 bzw. [58]).

Wir haben in Kapitel 5 dann versucht, das Riechkolben-Aktivitätsprofil einer Einzelsubstanz in seine funktionalen Module zu zerlegen. Hierbei sind wir einen Schritt weiter als bei der reinen Erstellung von flachen Karten wie in Kapitel 4 oder auch in [57], da wir mit den *Decision Images* (DI, siehe Abschnitt 5.1.3) in der Lage sind, zumindest für ein einzelnes Aktivierungsprofil den Zusammenhang zwischen bestimmten Beschreibern herauszuarbeiten. Weiterhin helfen die DIs die Konvergenzpunkte bestimmter Rezeptortypen weiter einzugrenzen.

Dennoch können wir damit noch nicht vollständig den kombinatorischen Kode der Geruchswahrnehmung entschlüsseln. In diesem Kapitel wollen wir abschließend zumindest noch einmal testen, ob nicht doch eines der ganz einfachen Kodierungsmodelle in der Lage ist, die Wahrnehmung von bestimmten Geruchsqualitäten oder einigen der chemischen Eigenschaften gut zu beschreiben.

Es gilt dabei auch hier wie in den anderen Kapiteln: Da die Datenlage sehr dünn und



**Abbildung 6.1:** Die intuitivste Variante ordnet jedes der vier Module A, B, C und D jeweils einer Klasse zu (Single-Unit-Kode). Beim kombinatorischen Kode setzt sich eine Klasse jeweils aus einer Kombination von aktiven Modulen zusammen. Bei einer strikten Kombinatorik gibt es ein eindeutiges Muster (z.B. A, B und C für *apple*). Bei einer multiplen Kombinatorik (hier zweifach) gibt es mehrere Kombinationen, die zu einer Klasse zusammengefasst sind (z.B. A+B und C+D ergeben beide *apple*).

das System so komplex ist, versuchen wir hier eine Hypothesenmenge von möglichst einfachen Zusammenhängen auf den existierenden Daten zu testen, mit der Idee, die Geruchsforschung in zukünftigen Experimenten in eine möglichst vielversprechende Richtung zu lenken.

### 6.1.1 Kodierungsmodelle

Das einfachste Modell wäre eine Kodierung entlang einzelner Module des Bulbus. Dies entspricht einer eineindeutigen Korrelation zwischen einzelnen Duftstoffen und der Aktivität in einem bestimmten Sektor des Riechkolbens. Ob es sich hierbei dann um die Aktivierung eines einzelnen Rezeptors oder einer Rezeptorfamilie handelt, ist hierbei erst einmal sekundär. Schematisch ist dieses Modell in Abbildung 6.1 als *Single-Unit-Kode* dargestellt. Teilen wir einen Modellbereich in vier Module A, B, C und D auf, so entspricht in diesem einfachen Kode ein aktives Modul A beispielsweise einer Substanz, die apfelig riecht. Eine weitere Substanz, die eher bananig riecht, stimuliert in Modul A keine Zelle, sondern ausschließlich in Modul B. Ein drittes Modul C könnte dann z.B. den kirschtigen Substanzen zugeordnet werden usw..

Es sei an dieser Stelle angemerkt, dass sich diese Überlegungen zur Kodierung der Geruchsqualitäten nicht auf selbige beschränken. Ganz analog und ohne Ein-

---

schränkung lässt sich dieses Modell auch auf Single-Unit-Kodierung bzgl. chemischer Eigenschaften übertragen: Ein Modul zeigt genau dann Aktivität, wenn das Riechsystem durch eine Gruppe von z.B. Alkoholen oder Estern stimuliert wird.

Ist die Aufnahmequalität der Uptake-Bilder ausreichend, so ist bei diesem Kodierungsmodell sofort einsichtig, dass wir eine solche Region leicht identifizieren könnten. Man braucht nur die Bereiche des Bulbus betrachten, die Aktivität nur für genau eine der Substanzklassen zeigen.

Nun ist diese Art der Kodierung nicht zu erwarten. Selbst wenn wir z.B. bezüglich *fruity* in Abschnitt 5.2.3 bereits sehen konnten, dass es spezifische Regionen gibt, die sowohl strukturelle wie auch funktionelle Bedeutung haben, so gibt es dennoch keine verallgemeinerbare Systematik, die auf ein solches Paradigma hinweisen würde.

Johnson et al. [20, 29, 38] haben bereits in mehreren Arbeiten versucht, eine solche Modularisierung der Bereiche vorzunehmen, allerdings mit nur wenig Erfolg. Für jedes Modul und jede Modulkategorie finden sich immer wieder Ausnahmen. Zwar gibt es immer mehr Arbeiten, in denen Karten bzgl. der chemischen Eigenschaften erstellt werden [69, 71, 104], aber bei der Interpretation und den Schlüssen auf den eigentlichen Kode bzgl. der Geruchswahrnehmung stoßen wir erneut auf das bereits in Abschnitt 3.1.2 ausführlich dargestellte Dilemma: Kennen wir weder den Wahrnehmungsraum noch ein für ihn geltendes physikalisches Kontinuum, so können wir aus einer Darstellung des Merkmalsraums so gut wie nichts über das Wahrnehmungssystem lernen.

Weiterhin muss man feststellen, dass die Modularisierung des Bulbus in klar abgegrenzte Module dem eigentlichen Gedanken eines kombinatorischen Kodes widerspricht. Stellen wir uns jedes Modul als einen Buchstaben des Riechalphabets vor (siehe auch Abschnitt 5.2.2), so kann und wird vermutlich auch jeder Buchstabe in mehreren Wörtern eine möglicherweise andere Bedeutung haben.

Ein etwas komplexeres Modell, welches allerdings den Gedanken des kombinatorischen Kodes aufgreift und im Grunde weitestgehend unseren (theoretischen) Vorstellungen entspricht, ist ebenfalls in Abbildung 6.1 dargestellt. Es handelt sich hierbei um den einfachen oder auch *strikten kombinatorischen Kode*. Strikt, weil wir annehmen, dass eine Klasse durch *genau einen* kombinatorischen Kode repräsentiert wird. Die Qualität *apple* würde beispielsweise genau dann wahrgenommen werden, wenn es eine spezifische Modulaktivierung in der Kombination A, B und C gibt. A und B sind hier zwar auch Teil des Geruchswortes für *banana*, aber nur in der Kombination A, B und D wird diese Wahrnehmung tatsächlich ausgelöst.

Können wir dieses Kodierungsmodell ebenfalls aus gutartigen Daten ablesen? Implizit haben wir mit der Einführung der Decision Images in Kapitel 5 bereits angenommen, dass wir eben solche Kodes erkennen können. Wir gehen davon aus, dass

es eine klassenspezifische (Linear-)Kombination von aktiven Regionen gibt, anhand derer wir Klassenelemente von nicht-Klassenelementen unterscheiden können.

Es besteht allerdings auch die Möglichkeit, dass es sich um einen *multiplen kombinatorischen Kode* handelt. Dies ist ganz rechts in der Abbildung 6.1 dargestellt. Nehmen wir an, unsere Klassenbeschreibung *apple* ist nicht ganz präzise auf die real existierenden Muster abzubilden, vielmehr gibt es zwei verschiedenartige *apple*-Kombinationen. Ein Teil der Substanzen löst eine eher *apple1*-artige Wahrnehmung aus, während die übrigen Klassenelemente eher durch ein *apple2*-Muster zu beschreiben sind.

Wie würde sich eine solche Kodierung der Trainingsdaten auf das Lernen der Decision Images auswirken? Wir haben postuliert, dass wir aus den Mustern der Decision Images Kandidaten für die am Kode beteiligten Module herausrechnen können. Ob dies auch für die eben vorgestellten Modelle funktionieren kann, haben wir unter kontrollierten Bedingungen auf synthetischen Daten getestet.

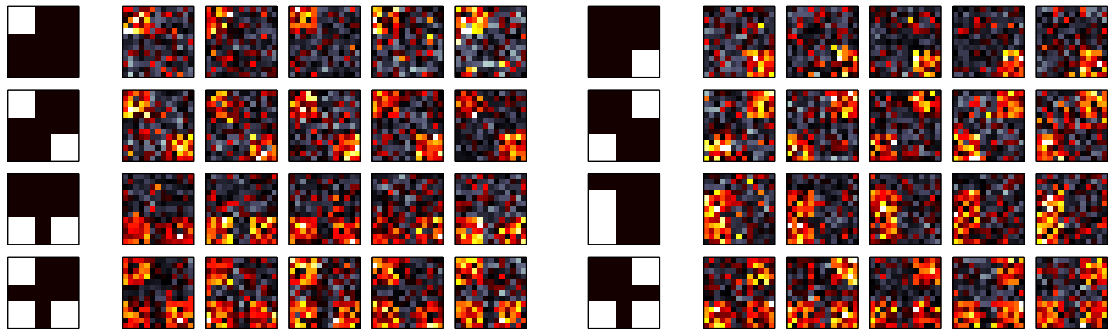
### 6.1.2 Synthetische Modellierung

In Abbildung 6.1 haben wir bereits ein simples Beispielmuster vorgestellt. Wir werden nun einige ganz ähnlich strukturierte Beispiellassen entwerfen. Damit die Ergebnisse aber vergleichbar mit dem realen System werden, werden wir die einzelnen Kodewörter gemäß des in Abschnitt 5.1.2 vorgestellten Rauschmodells verrauschen. Anschließend berechnen wir auch für diese synthetischen Muster  $\mathbf{x}_i$  DIs  $\mathbf{d}_\mu$  für die jeweiligen Klassen  $\mu$ .

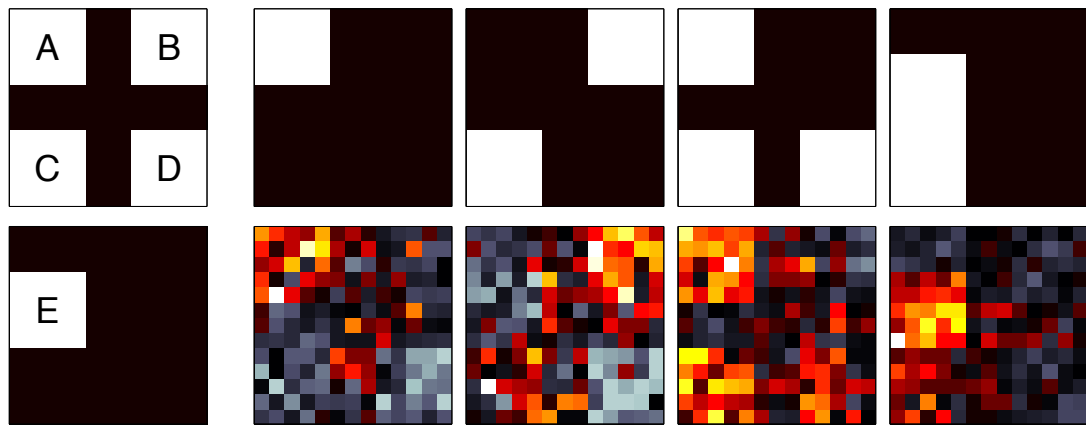
Abbildung 6.2 zeigt den virtuellen Kode, bei dem wir insgesamt 8 Riechwörter aus den in Abb. 6.3 skizzierten fünf Regionen *A-E* zusammengesetzt haben. Im Vergleich zu Beispiel 6.1 wurde eine fünfte Region *E* eingeführt, die mit *A* überlappt. Für jede dieser Kombinationen haben wir fünf Trainingsbilder erzeugt und nach dem Rauschmodell aus Abschnitt 5.1.2 in Anlehnung an unsere Realdaten verrauscht. Es sei hier angemerkt, dass sich auch die Größe der virtuellen Module der vermuteten Größe eines Glomerulus auf den Uptake-Daten entspricht.

Es wird schon durch visuelle Inspektion klar, dass eine Single-Unit-Kodierung auf den Daten herauszuarbeiten wäre. Aufgrund des starken Rauschens sind natürlich die Ausgangskonturen kaum komplett zu rekonstruieren, aber es ergibt sich eine eindeutige Zuordnung von Regionen auf den Decision Images  $\mathbf{d}_\mu$  und dem zugehörigen Klassenalphabet  $\mathcal{A}_\mu$ .

Abbildung 6.3 zeigt vier klassenspezifische Decision Images  $\mathbf{d}_\mu$ , gelernt auf den 40 Beispielen  $\mathbf{x}_i$ , jeweils mit den fünf Klassenbildern  $\mathbf{x}_{i \in \mathbf{I}_\mu}$  als positive Beispiele. Die signifikanten Regionen erhalten in der Tat hohe Werte auf den Decision



**Abbildung 6.2:** Für acht synthetische Geruchswörter wurden jeweils 5 Vertreter nach dem Rauschmodell aus Abschnitt 5.1.2 zufällig verrauscht. Die erste Spalte zeigt jeweils das unverrauschte Signal in schwarz-weißer Darstellung, gefolgt von den Rauschmustern. Die rot-gelben Farben stellen hohe Werte dar, grau-schwarz niedrige.



**Abbildung 6.3:** Die erste Spalte zeigt die Modulaufteilung für die synthetischen Muster, von der zweiten Spalte an sind in der ersten Zeile vier Beispiele für das DMO-Training gezeigt, konkret die Kodewörter (von links nach rechts) *A*, *BC*, *ACD* und *CE*. In der zweiten Zeile sind die zu dieser Klasse gelernten Gewichtsvektoren des DMO-Klassifikators zu sehen. Gelernt wurden jeweils die fünf Bilder der Klasse gegen die verbleibenden 45 Bilder (siehe auch Abb. 6.2).

Images  $\mathbf{d}_\mu$ , am vierten Beispiel (Kodewort *CE*) kann man weiterhin ganz explizit sehen, wie besonders exklusive Regionen (hier Modul *E*) im Ergebnis eine wesentlich höhere Klassensignifikanz zugewiesen bekommen als die strikt-kombinatorischen Anteile des Kodewortes. Selbst wenn wir also implizit nach strikten kombinatorischen Codes suchen, würden wir vermutlich dennoch vorhandene Single-Units herausarbeiten.

Glomeruli mit Aktivität haben auf dem Uptake-Material eine erwartete Ausdehnung von etwa (3x3) Pixel, da die Aufnahmen aus unterschiedlichen Tierpräparaten gemittelt werden und es keine Registrierung gibt, die eine exakte Überlagerung gewährleistet. Finden wir im Bildmaterial also einzelne Pixel mit hohem Uptake in

einer nichtaktiven Nachbarschaft, so handelt es sich vermutlich um eine spontane Aktivität. Um solche Signale zu eliminieren, verwendet man typischerweise morphologische Filter [32].

Die resultierenden Decision Images wurden daher mit Hilfe eines morphologischen Öffnens (eine Erosion mit anschließender Dilatation) in eine modulare Form gebracht, so dass wir auch auf dem Ergebnis unserer Realdatenanalyse Module formulieren und mit den Ausgangsdaten vergleichen konnten [73].

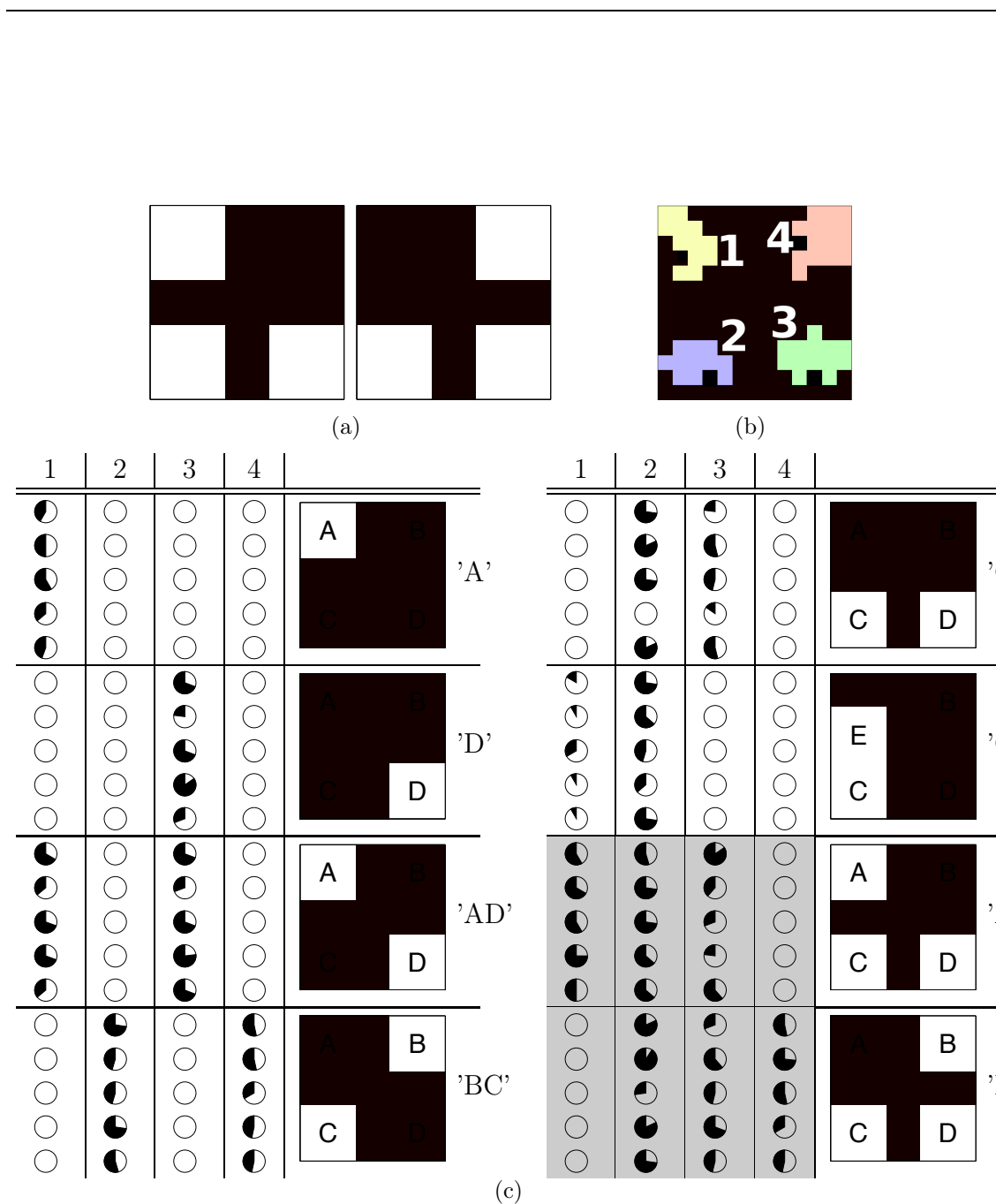
Wie verhält es sich aber nun mit den *multiplen* kombinatorischen Kodes, also den Klassen, die sich aus zwei oder mehr strikten Kombinationen zusammensetzen? Können wir mit Hilfe der Decision Images auch diese Kodes herausarbeiten? Wir haben es an dem synthetischen Modell versucht, indem wir uns einen zweifachen kombinatorischen Kode gewählt und gelernt haben. Es soll vor allem untersucht werden, ob sich Module ergeben, anhand derer wir identifizieren können, ob und aus wieviel strikten Wörtern sich der Kode zusammensetzt.

Es sollen zwei der strikten Kodewörter zu einer Klasse zusammengefasst und das dazugehörige Decision Image gelernt werden. Ausgewählt wurden die Worte *ACD* und *BCD*, beide decken zusammen die gesamte Modulkarte ab und haben viele Gemeinsamkeiten mit den übrigen Kodewörtern. Für die Klassifizierungsgüte sind die beiden Worte potentiell eine schlechte Wahl, da sie sich bereits alleine durch ihren Aktivierungsumfang von drei aktiven Modulen von den anderen unterscheiden lassen. Evaluieren wollten wir allerdings vor allem, wie deutlich sich die mehrfache Kodierung in den Modulkarten widerspiegelt.

In Abbildung 6.4.a sind die beiden Kodewörter noch einmal gezeigt, die wir zu einer einzigen Klasse zusammengefasst haben. Abbildung 6.4.b zeigt das auf den 10 Klassenbildern und den verbleibenden 30 Nicht-Klassenbildern gelernte Decision Image in seiner modularisierten Form. Es haben sich nach der morphologischen Filterung 4 Module ergeben, die hier unterschiedlich eingefärbt sind.

Unter Verwendung des Decision Images haben wir nun die Aktivität innerhalb der Module für alle am Lernen beteiligten Bilder berechnet und in eine tabellarische Form gebracht. Abbildung 6.4.c zeigt die Modulaktivitäten für alle 40 Bilder.

Man kann bereits erahnen, dass die Interpretation der Ergebnisse vermutlich nicht so einfach wird. Es ist aber einzusehen, dass wir auch in dieser Kodierungsvariante das Decision Image die beteiligten Module lernt. Somit erhalten wir in diesem Fall ebenfalls Kandidaten für potentielle Buchstaben, über welche sich Kodewörter definieren lassen. Auch wenn wir das hier nicht explizit gezeigt haben, wie schon bei den strikten kombinatorischen Kodes würde sich auch in den Decision Images eine Single-Unit durchsetzen, gefolgt von potentiellen strikten (also klassenexklusive) Kombinationen der beteiligten Module.



**Abbildung 6.4:** Übersicht zur Auswertung der Aktivität auf Kandidatenregionen. Die Datensätze mit den Mustern der in (a) dargestellten Typen wurden positiv gelabelt. In (b) ist die in Regionen aufgeteilte Lösung des Double-Min-Over-Klassifikators dargestellt. Die Tabelle (c) stellt dar, wie groß die Aktivität auf den einzelnen Datensätzen in den vier Regionen ist. Die grau hinterlegten Einträge sind diejenigen, die bei der Klassifikation positiv gelabelt waren.

Was sind nun die praktischen Implikationen unseres Modells? Vermutlich können wir nicht endgültig herausarbeiten, welches der drei Kodierungsmodelle vorliegt, zumindest nicht, wenn wir z.B. eine exklusive, dominante Region finden. Denn vielleicht eliminiert diese Single-Unit andere ebenfalls klassentypische Modulkombinationen. Umgekehrt jedoch sollten wir Hinweise auf kombinierte Module für jedes dieser Modelle finden, sofern sie auf den gegebenen Daten existieren.

## 6.2 Kodierungsmodelle auf den Uptake-Geruchsdaten

### 6.2.1 Multiple kombinatorische Geruchskodes

Wir wollen dieselbe Analyse wie auf den synthetischen Daten nun auf den realen Daten ausführen. Hierzu haben wir für jede der OC- und OFD-spezifischen Geruchsklassen Decision Images (DIs) lernen lassen und auf den resultierenden DIs die zugehörigen Modulkarten berechnet. Diese Modulkarten wurden dann dazu verwendet, um die Uptake-Bilder der entsprechenden Klasse in ihre Modulaktivitäten zu zerlegen.

**Praktisches Beispiel: Geruchsqualität „Apfel“** Die Klasse *apple* ist in unseren Daten durch 10 Uptake-Bilder vertreten, darunter einige Duplikate von Substanzen, die in unterschiedlichen Konzentrationen aufgenommen wurden. Insgesamt sind es 6 verschiedene Duftstoffe, die hier zu der Modulkarte beitragen. Dies ist insofern interessant, dass wir ebenfalls prüfen können, ob die Bilder von identischen Substanzen tatsächlich auch durch identische Kodewörter repräsentiert werden.

Analog zu den künstlichen Bildern haben wir also das DI berechnet und modularisiert. Wir erhalten 10 Module, die potentiellen Kandidaten für die in diesem kombinatorischen Kode verwendeten Buchstaben entsprechen. Bleibt eigentlich nur noch die Frage, aus welchen Kombinationen sich die konkreten Substanzen der Klasse aus den Modulen zusammensetzen lassen.

In Abbildung 6.5 sieht man nun die Modultorten, wie wir sie eben bereits auf den synthetischen Daten verwendet haben. Je voller die Torte, desto höher die Aktivierung des entsprechenden Moduls durch die der jeweiligen Zeile zugeordneten Substanz. Erfreulich ist erst einmal, dass sich in der Tat alle Duplikate offensichtlich recht konsistent durch die Modulkarten rekonstruieren lassen. Die drei Acetone sind ebenso wie die beiden ethyl-tragenden Duplikate durch relativ identische Modultorten vertreten.



| A | B | C | D | E | F | G | H | I | J | Duftstoff      |
|---|---|---|---|---|---|---|---|---|---|----------------|
| ○ | ◐ | ○ | ◐ | ● | ○ | ○ | ○ | ◐ | ○ | acetone        |
| ○ | ◐ | ○ | ◐ | ● | ○ | ○ | ○ | ● | ◐ | acetone        |
| ○ | ◐ | ○ | ● | ● | ○ | ○ | ○ | ◐ | ◐ | acetone        |
| ◐ | ◐ | ◐ | ◐ | ○ | ◐ | ◐ | ● | ◐ | ◐ | ethyl caproate |
| ◐ | ◐ | ◐ | ● | ○ | ◐ | ○ | ◐ | ○ | ◐ | ethyl caproate |
| ○ | ◐ | ● | ● | ◐ | ◐ | ◐ | ● | ◐ | ● | ethyl valerate |
| ○ | ● | ● | ◐ | ○ | ◐ | ● | ● | ◐ | ● | ethyl valerate |
| ● | ◐ | ○ | ○ | ○ | ● | ○ | ◐ | ○ | ◐ | hexyl acetate  |
| ◐ | ◐ | ○ | ○ | ◐ | ● | ◐ | ● | ● | ◐ | nerolidol      |
| ◐ | ◐ | ○ | ○ | ○ | ○ | ◐ | ● | ○ | ○ | neryl acetate  |

**Abbildung 6.5:** Modularisierung der OC-spezifischen Klasse *apple*. Auf der Basis des modularisierten Decision Images für *apple* ist hier für alle 10 Module (A-J) durch Tortendiagramme dargestellt, wieviel Prozent des jeweiligen Moduls (Spalten) durch den gegebenen Duftstoff (Zeile) aktiviert wird. Welches der vorgestellten Kodierungsmodelle hier zu tragen kommt ist leider nicht erkennbar.

Ansonsten allerdings ist es augenscheinlich, dass unsere Modulkarte keine eindeutige Gruppierung ermöglicht. Es handelt sich hier ganz offensichtlich um keinen zweifachen kombinatorischen Code, ebenso wie wir wenig Anzeichen von einem einzelnen (strikten) Code finden können. Vielmehr scheint jede der Substanzen ihre eigenen Modulwörter zu besitzen, über die sich die Substanz von der Nicht-Klasse unterscheiden lässt. Einzig die Module *B*, *I* und *J* scheinen so etwas wie eine klassenumfassende Aktivität zu zeigen. *H* zeigt zumindest für alle Substanzen außer Aceton recht hohe Aktivität.

Nun wollen wir prüfen, wie sich diese Modulaktivität zu den tatsächlichen Uptake-Bildern verhält. In Abbildung 6.6 sind nun die Module, die wir aus dem DI der Klasse *apple* gelernt haben, in ihrem räumlichen Kontext auf den Bulbus projiziert. Die Buchstaben in Abb. 6.6.a entsprechen den Modulbezeichnungen, die wir auch in Abbildung 6.5 für die Spalten verwendet haben. Man kann gut sehen, wie die meisten Substanzen untereinander durchaus sehr unterschiedliche Aktivierungsmuster aufzeigen, wie jedoch alle einen spezifischen Anteil ihrer Energie in Modulbereichen des DI besitzen.

## 6.2.2 Substanzspezifische Unterräume

Wir scheinen also doch nicht so schnell etwas über das konkrete Kodierungsmodell der Geruchswahrnehmung aussagen zu können. Immerhin jedoch gibt es starke Indizien, dass es kein einfacher, sondern vielmehr ein mehrfach kombinatorischer Code sein muss. Im Extrem wäre dies ein Code, bei dem einfach jede Riechsubstanz

[illegible]

---

In Abbildung 6.7 ist erneut der Unterschied zwischen OFD (links) und OC (links) skizziert. Da die Rezeptoren im Bulbus entlang ihrer biochemischen Struktur organisiert zu sein scheinen [37, 69], ist die OFD vermutlich für einen Maschinenerler die einfachere Aufgabe auf den gegebenen Daten: Um ähnliche Signale möglichst eindeutig und robust voneinander zu trennen, sind die sehr sensitiven Tuning-Kurven der Rezeptoren [70, 71] sicherlich das Werkzeug der Wahl, um den Kontrast zwischen zwei nahezu identischen Duftstoffen zu erhöhen. Die zugehörigen Rezeptoren werden wir vermutlich auf den Uptake-Daten in einem räumlichen Zusammenhang finden, was zu signifikanten Profilen auch auf den Decision Images führt. Diese Hypothese scheint sich bereits in der Kreuzvalidierung aus Kapitel 5 bestätigt zu haben.

Wie verhalten sich nun aber die Geruchsqualitäten verglichen mit den chemischen Eigenschaften? Grundsätzlich gibt es genau drei Möglichkeiten:

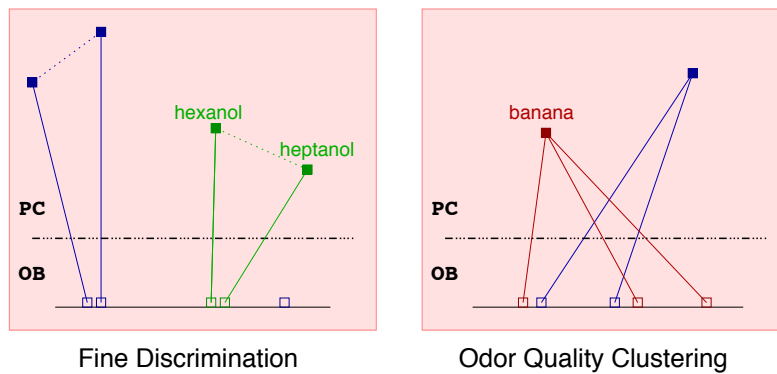
1. OC-spezifische Glomeruli sind eineindeutig OFD-spezifischen zuzuordnen.
2. Ein Teil der OC-spezifischen Glomeruli ist auch OFD-spezifisch.
3. Keiner der OFD-spezifischen Glomeruli ist auch OC-spezifisch.

Im ersten Fall wäre der zu beobachtende Zusammenhang leicht zu motivieren. Vermutlich könnte man dann sogar die OC-spezifischen Glomeruli der Geruchsklassen direkt den OFD-spezifischen Regionen zuordnen. Mit anderen Worten würden wir dann erwarten, dass das Kodewort für eine gegebene Substanz grundsätzlich durch seine OFD-spezifischen Module geprägt ist und die OC-spezifischen Bereiche echte Untermengen dieser Kodewörter sind. In diesem Fall würde das Riechsystem die Geruchsqualität aus der biochemischen Molekülinformation ableiten.

Würden wir also ein OC- und ein OFD-spezifisches Profil mit einem zugehörigen Uptake-Bild maskieren und die beiden resultierenden Riechwörter übereinander legen, so würden wir – gegeben Variante 1 – erwarten, dass es OC-Anteile nur dort gibt, wo sich auch OFD-Anteile befinden. Die Glomeruli, die die chemischen Eigenschaften des Moleküls gut herausarbeiten können, wären dann also auch verantwortlich für die wahrgenommene Qualität.

Sollte die Information für die wahrgenommene Qualität tatsächlich durch die Kombination von chemischen Eigenschaften der Duftstoffe kodiert sein, so müsste diese Information nämlich auch Teil des chemischen Musters sein; die z.B. an Alkohole gekoppelten Qualitäten auf ihren Decision Images müssten auch eine echte Schnittmenge mit dem Alkohol-Pattern besitzen. In anderen Worten, die Qualitäten würden aus Teilen der chemischen Muster zusammengesetzt werden.

Dies führt uns direkt zu der zweiten Hypothese: Nur ein Teil der Glomeruli, die nützlich sind, den OFD-Task zu lösen, also ähnliche Substanzen auseinander zu halten, sind auch prädiktiv für das Erkennen von typischen Gerüchen (dem OC-



**Abbildung 6.7:** Links: Schematische Darstellung der sog. *Fine Discrimination*. Chemisch ähnliche Substanzen wie z.B. Heptanol und Hexanol sollen unterscheidbar sein. Rechts: Schematische Darstellung des sog. *Odor Clustering*. Auch chemisch unterschiedliche Substanzen können gleich riechen. OB bezeichnet den Riechkolben, PC deutet hier die Riechrinde und höhere kortikale Regionen an.

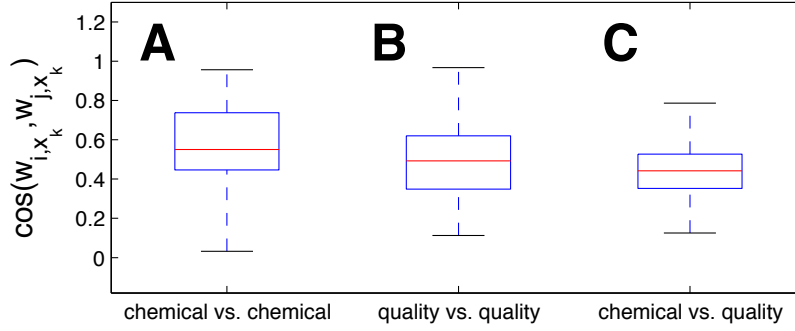
Task).

Die verbleibende Annahme führt uns zu der Möglichkeit, dass es komplett andere Bereiche sind, die die Qualitäten am Besten beschreiben. Diese Vermutung ist insofern naheliegend, dass die beiden Prozesse gegenläufige Informationen brauchen: OFD funktioniert am besten auf den Rezeptorantworten, die auf chemisch sehr ähnliche Substanzen mit unterschiedlichem Tuning reagieren. Für eine Art Riechkontrast, bei dem nahezu gleichartige Substanzen dennoch unterschiedlich wahrgenommen werden können, besitzen alle Glomeruli, die auf diese Klasse von Stoffen reagieren, eine hohe Relevanz und werden im Pattern für die entsprechende Klasse vermehrt auftreten.

Mit anderen Worten, Geruchsqualität wäre eine Eigenschaft, die kombinatorisch aus einem komplett anderen „Alphabet“ zusammengesetzt wäre.

Um die realen Daten auf diese Möglichkeiten zu überprüfen, haben wir die Ähnlichkeiten jeweils zweier *Riechwörter* miteinander verglichen. Riechwörter sind hier die substanzbezogenen Uptake-Bilder, maskiert mit einem der Decision Images, wie wir es in Abschnitt 5.2.2 vorgestellt haben. Konkret haben wir für jedes der Uptake-Bilder  $\mathbf{x}_i$  die zugeordneten Decision Images maskiert und bezüglich dieses Stimulus miteinander verglichen.

Im Falle, dass zwei Spuren  $\mathbf{s}_{\mu,i}$  und  $\mathbf{s}_{\nu,i}$  auf  $\mathbf{x}_i$  in dieselbe Richtung zeigen (dann gilt  $\cos(\mathbf{s}_{\mu,i}, \mathbf{s}_{\nu,i}) = 1$ ), gehen wir davon aus, dass beide Klassen potentiell durch das gleiche *Riechwort* beschrieben werden. Sind die beiden maximal unterschiedlich, also stehen die Vektoren senkrecht aufeinander ( $\cos(\mathbf{s}_{\mu,i}, \mathbf{s}_{\nu,i}) = 0$ ), so gehen wir davon aus, dass die beiden Klassen – zumindest bezüglich dieser Substanz – durch kein sehr ähnliches Wort beschrieben werden.



**Abbildung 6.8:** Winkelverteilung zwischen Paaren von Spuren  $s_{\mu_1,i}, s_{\mu_2,i}$ . **A:** Paarweiser Winkelvergleich zwischen chemischen Eigenschaften  $\cos(s_{\mu,i}^C, s_{\nu,i}^C)$ . **B:** Paarweiser Vergleich zwischen Geruchsqualitäten  $\cos(s_{\mu,i}^O, s_{\nu,i}^O)$ . **C:** Gemischter Vergleich zwischen Geruchsqualitäten und chemischen Eigenschaften  $\cos(s_{\mu,i}^C, s_{\nu,i}^O)$  jeweils für alle möglichen Paarungen  $(\mu, \nu)$  mit  $\mathbf{x}_{\{i \mid i \in I_\mu \wedge i \in I_\nu\}}$ . Im Zentrum der Boxplots liegt der Median. Die obere und untere Grenze ist beschrieben durch das obere und untere Quartil. Die Whisker sind beschrieben durch die 1.5-fache Interquartilabstände.

In Abbildung 6.8 ist die eben beschriebene Winkelverteilung dargestellt. Wir haben uns drei unterschiedliche Paarungen angeschaut: chemisches Riechwort gegen chemisches Riechwort (a), Qualität gegen Qualität (b) und die Mischung aus beidem, also den Winkel zwischen den Riechwörtern einer Geruchsqualität und einer chemischen Eigenschaft (c), alle Paarungen jeweils relativ zu einem Uptake-Bild.

Wie wir es implizit bereits vermutet haben, bestätigt sich die Annahme, dass die maskierten Decision Images für die verglichenen chemischen Klassen eher hohe Kosinuswerte ergeben. Dies begründet sich sicher auch darin, dass viele der Klassen durch logische Abhängigkeiten zwangsläufig starke Überlappungen haben. Ein *primärer Alkohol* ist immer auch ein *Alkohol*, entsprechend sind die Regionen, die signifikant für den primären Alkohol sind, sicher auch nützlich, um Alkohole zu erkennen, zumindest die *primären* Klassenmitglieder.

Ähnlich, wenn auch nicht so intuitiv, verhält es sich auch bei den Geruchsqualitäten. Riecht eine Substanz nach *apple*, so bescheinigen Probanden für nahezu jeder dieser Substanzen ebenfalls eine *fruity*-Note. Besonders für die Geruchsqualitäten ist unser Wissen über solche Abhängigkeiten aber relativ gering (siehe dazu auch Kapitel 4). Solche Zusammenhänge scheinen aber wesentlich seltener in dieser Konfiguration aufzutreten, die Kosinuswerte erscheinen in Abbildung 6.8 insgesamt etwas geringer.

In der Tat konnten wir feststellen, dass im paarweisen Vergleich der Spuren, also der maskierten Decision Images, von chemischen Eigenschaften und Geruchsqualitäten der systematisch größte Unterschied besteht. Dies spricht für die weiter oben bereits herausgearbeitete Annahme, dass die Kodierung der Geruchsquali-

tät vermutlich mehr ist als einfach nur eine Teilmenge der jeweiligen chemischen Substanzeigenschaften.

In Abbildung 6.9 sind drei Beispiele solcher Paarungen aufgezeigt, um ein Gefühl dafür zu vermitteln, wie groß der Überlapp für einen bestimmten Kosinuswert überhaupt ist. Das erste Beispiel (A) zeigt einen sehr niedrigen Kosinuswert von nur 0.164, der beinahe einer Orthogonalität der Pattern entspricht, d.h. verwendet man die Aktivität der Substanz *3-Hexanone* über seine Stoffeigenschaften *small aliphatic ketone* und *ethereal*, so stellt man fest, dass diese beiden Klassen nahezu keinen gemeinsamen Anteil bezüglich der Substanz haben. Abbildung 6.9.B zeigt *Hexanal* als Maske für die chemische Eigenschaft *small aliphatic aldehyde* und die Geruchsqualität *penetrating*. Es ergibt sich ein Kosinuswert von 0.686, der ein relativ typisches Bild abgibt. Die meisten der übrigen Paarungen ergeben einen niedrigeren Kosinuswert, der Median liegt leicht über 0.4, wie man Abbildung 6.8.C entnehmen kann.

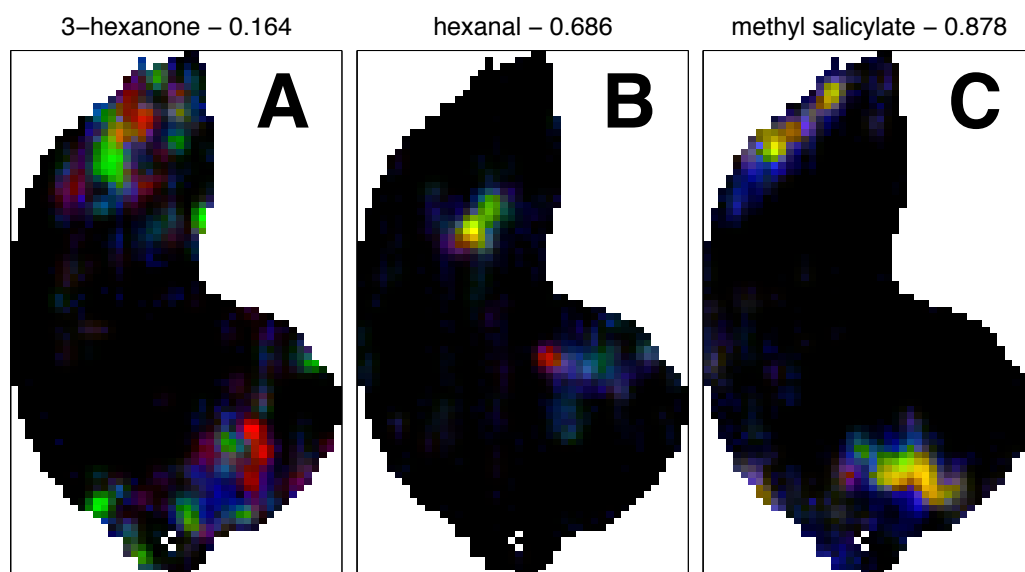
Beispiel C in Abbildung 6.9 zeigt letztlich eines der wenigen Beispiele, bei denen wir zwischen diesen beiden Typen von Substanzeigenschaften eine hohe Überlappung von 0.878 beobachten können. In diesem Fall handelt es sich um die Klassen *phenols* und *wintergreen*, die mit der Substanz *Methyl Salicylate* maskiert wurden. Es liegt nahe, hier eine enge Verbindung zwischen der Geruchsqualität und der chemischen Eigenschaft zu vermuten.

Offensichtlich sind für die Entscheidungsfindung der Decision Images substantiell andere Bereiche für die beiden Aufgaben (OC und OFD) geeignet. Benutzen die Muster für OC jedoch grundsätzlich andere Bereiche (Buchstaben des kombinatorischen Kodes) als die OFD Decision Images, so sollte sich über alle Klassen auch eine gewisse Aufspaltung der Glomeruli in OC und OFD Glomeruli ergeben.

Um dies zu überprüfen, haben wir diesmal nicht paarweise verglichen sondern haben gleich *alle* Klassenspuren einer Substanz überlagert und daraufhin untersucht, ob sich tatsächlich taskspezifische, disjunkte Regionen ergeben. Abbildung 6.10 zeigt dies anhand von einigen Beispielen aus mehreren Stoffklassen. In Anlehnung an Abbildung 5.8 ist im Blaukanal das Uptake-Bild der jeweiligen Substanz zu sehen. Weiterhin haben wir in den Grünkanal die Maximalwerte der maskierten chemischen Eigenschaften des Stoffes geschrieben und im Rotkanal die Maximalwerte der Geruchsqualitäten dieses Stoffes. In den sechs Plots ist der Blaukanal jeweils mit der Rot-Grün-Mischung der anderen beiden Kanäle überlagert.

Es gibt zwei grundsätzliche Eigenschaften, die auf allen Beispielbildern recht deutlich zu erkennen sind:

1. Es sind nur noch wenige Bereiche der Substanzen wirklich blau.
2. Es gibt insgesamt nur wenige gelbe Regionen, die meisten überlagerten blauen Bereiche sind rot oder grün.

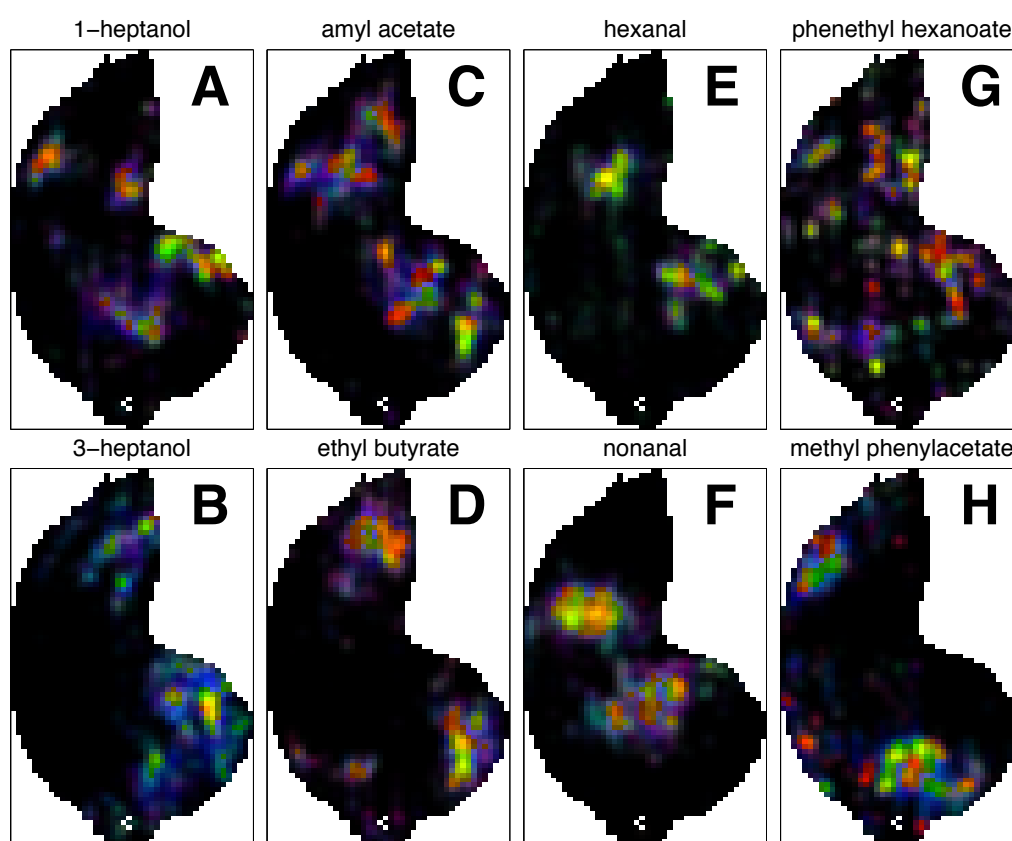


**Abbildung 6.9:** Beispiele für unterschiedliche Kosinuswerte. **A:** Die maskierten Decision Images für die Klassen *small aliphatic ketones* und *ethereal* zeigen für 3-Hexanon einen sehr niedrigen Kosinuswert (0.164), d.h. nahezu keine Überlappung. **B:** *Small aliphatic aldehyde* und *penetrating* maskiert mit *Hexanal* erreicht 0.686, die meisten Paarungen erreichen niedrigere Werte, zeigen also noch weniger Überlappung. **C:** *Phenols* und *wintergreen* maskiert mit *Methyl Salicylate* erreichen 0.878, ein außergewöhnlich hoher Wert, welche eine starke Bindung von Wert und Merkmal andeutet.

Da es kaum noch blaue Regionen auf den überlagerten Bildern gibt, scheint unsere Zerlegung entlang der hier betrachteten Klassen bereits recht gut das gesamte Aktivierungsmuster abzudecken. Es hätte durchaus sein können, dass wir durch unsere Decision Images nur Bruchteile des gesamten Uptake-Bildes erklären können.

Die geringe Durchmischung illustriert erneut die Vermutung, die wir bereits anhand der paarweisen Untersuchung der maskierten Profile aufgestellt haben: Die Regionen für die OC- und OFD-spezifischen Decision Images scheinen sogar recht erstaunlich disjunkt zueinander zu sein. Lediglich *3-Heptanol* und *Hexanal* zeigen kaum Rotanteile. Beide Substanzen haben aber auch kein sehr ausgeprägtes Riechprofil.

Man sollte es hier ruhig noch einmal betonen: Natürlich finden wir für nahezu alle Substanzen Bereiche, in denen sich die prädiktiven Muster, also die maskierten Decision Images, beider Gruppen überlagern, wie wir es in Abbildung 6.10 z.B. besonders ausgeprägt bei den beiden Aldehyden *Hexanal* (E) und *Nonanal* (F) sehen. Das besondere hier ist vielmehr, dass wir in den unabhängig voneinander berechneten maskierten Anteilen sehr stark verzahnte aber doch disjunkte Bereiche von Regionen sehen, die entweder nur prädiktiv sind für die chemischen Eigenschaften (grün) oder nur für die Geruchsqualitäten (rot).



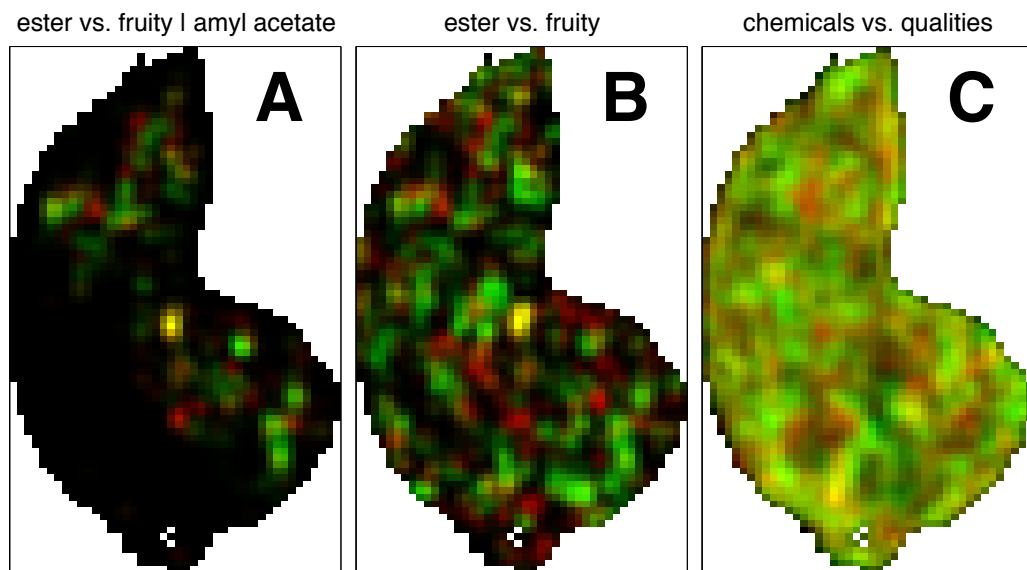
**Abbildung 6.10:** RGB-Plot mit der positiven Uptake-Energie der Substanzen im Blaukanal, den Maxima der maskierten OFD-spezifischen Decision Images im Grünkanal und den Maxima der maskierten OC-spezifischen Decision Images im Rotkanal. Die Spalten zeigen jeweils zwei Beispielsubstanzen aus den Gruppen der (von links nach rechts) alcohols (*1-Heptanol* und *3-Heptanol*), aliphatic esters (*Amyl Acetate*, *Ethyl Butyrate*), aldehydes (*Hexanal*, *Nonanal*) und aromatic esters (*Phenyl Hexanoate*, *Methyl Phenylacetate*).

### 6.2.3 Chemische Eigenschaften vs. Geruchsqualitäten

Leider lassen sich die maskierten Bilder der Einzelsubstanzen bezüglich ihrer jeweiligen Klassen nur sehr eingeschränkt miteinander vergleichen. Nicht nur, dass es kaum ein Paar von Substanzen mit genau identischen Eigenschaften gibt, jeder Duftstoff hat ja seinen eigenen, eindeutigen kombinatorischen Code und damit eine Maske, die auf einem bestimmten Unterraum gültig ist und somit mit dem maskierten Muster einer anderen Substanz nicht oder nur sehr eingeschränkt zu überlagern ist.

Handelt es sich aber tatsächlich um eine strukturelle Eigenschaft des Bulbus und nicht nur um ein substanzspezifisches Artefakt, so sollte sich diese Organisation von OFD- und OC-prädiktiven Regionen auch auf die unmaskierten, bulbusweiten





**Abbildung 6.11:** RG-plot der OFD-spezifischen (grün) gegen die OC-spezifischen Decision Images (rot). **A:** Die DIs für *ester* und *banana* maskiert bzgl. Amyl Acetate. **B:** Die unmaskierten DIs für *ester* und *banana* **C:** Alle 36 OFD-spezifischen und 44 OC-spezifischen Muster unmaskiert überlagert.

Decision Images übertragen lassen.

In Abbildung 6.11.A sind – maskiert durch *Amyl Acetate* – zwei Klassen übereinandergestellt, die eigentlich maximal ähnlich sein müssten. Wir haben es bereits in Kapitel 5 als Beispiel aufgeführt: *ester* und *fruity*. Es gibt zwar durchaus Substanzen, die keine Ester sind aber fruchtig riechen, aber umgekehrt riechen wirklich alle Ester fruchtig. Somit sollten Ester-spezifische Rezeptoren ebenfalls prädiktiv für die Qualität *fruity* sein und sich diese Muster entsprechend auch überlappen. Tatsächlich gibt es auch einige wenige Punkte, die einen deutlichen Überlapp aufweisen. Die meisten Regionen hingegen scheinen zwar dicht beieinander zu liegen, sind aber grundsätzlich alle räumlich voneinander getrennt.

Dieser Eindruck verstärkt sich sogar noch, betrachten wir die unmaskierten Decision Images für die beiden Klassen: In Abbildung 6.11.B sind die Decision Images für *ester* und *fruity* komplett (d.h. ohne Maskierung) übereinander gelegt. Auch hier sieht man eine ganz deutliche Trennung der Regionen.

In Abbildung 6.11.C schließlich haben wir die Maximalwerte über *alle* OFD- und OC-spezifischen Decision Images (36 chemische Eigenschaften (grün) und 44 Geruchsqualitäten (rot)) ermittelt. Und dennoch gibt es in der überlagerten Mischung kaum gelbe Flächen. Ebenso gibt es nahezu keine schwarzen Bereiche auf der Riechkolbenoberfläche. Auch dies scheint ein Hinweis darauf zu sein, dass die Abdeckung unserer Muster beinahe den gesamten Bulbus umfasst.

Offensichtlich scheinen sich die Decision Images aller hier betrachteten chemischen Klassen im direkten Vergleich zu denen aller Geruchsqualitäten grundsätzlich unterschiedlich zu organisieren.

Man sollte es ruhig noch einmal unterstreichen, dass diese deutliche Differenzierung zwischen diesen beiden Gruppierungen erstaunlich ist, da die Konturen trotz des komplett unabhängigen Lernprozesses relativ scharf sind und alleine schon die Messungenauigkeit wesentlich unschärfere Klassengrenzen hätte erwarten lassen.

### 6.3 Diskussion

In diesem Kapitel haben wir gezeigt, wie wir mit den Decision Images auf synthetischen Daten verschiedene Kodierungsmodelle erkennen können, sofern sie in den Daten existieren. Diesen „Nebeneffekt“ haben wir genutzt, um auch die Uptake-Daten auf diese Kodierungsmodelle zu testen.

Es scheinen aber tatsächlich keine der einfachen Kodierungen (Single-Units, strikte und zweifache kombinatorische Codes) auf den Daten zu existieren. Dies ist insofern ein interessantes Ergebnis, da man eigentlich einen fundierten, modulbasierten Code für die Riechkolben-Aktivitäten erwartet hätte. Unsere Ergebnisse scheinen darauf hinzuweisen, dass sich die Kodierung im wesentlichen auf eine eindeutige Abbildung der Substanzen im Sinne eines Fingerabdrucks zu beschränken scheint.

Wir haben weiterhin festgestellt, dass bei vielen Geruchsstoffen die prädiktiven Regionen für bestimmte Qualitäten räumlich sehr klar getrennt sind von allen chemischen Merkmalsklassen, die wir im Rahmen dieser Arbeit untersucht haben. Selbst bei Eigenschaften, die klar an molekulare Eigenschaften gebunden scheinen, wie z.B. der fruchtige Charakter von Estern, finden wir immer einen gewissen Anteil, der sich eben genau nicht an das chemische Muster anpasst.

Faszinierenderweise lässt sich diese Beobachtung auch klassen- und substanzübergreifend wiederholen: Selbst das Überlagern aller 80 unmaskierten Decision Images offenbart diese räumliche Trennung. Es scheint also in der Tat einige Rezeptoren zu geben, die durch ihr spezifisches Tuning eher OC-spezifisch sind, während andere eher ein OFD-spezifisches Tuning besitzen.

Diese Erkenntnisse passen gut zu theoretischen [86] und experimentellen [24, 45] Arbeiten zu diesem Thema, die ebenfalls über eine funktionelle Trennung von OFD und OC spekulieren. Besonders die experimentellen Arbeiten sind interessant, da es hier um den Pyriform Cortex geht, also die Riechrinde. Die Riechrinde teilt sich funktional in zwei Regionen, die beide den identischen Input aus dem Bulbus

---

erhalten. Es konnte nachgewiesen werden, dass der vordere Teil verantwortlich ist für die OFD, während sich der hintere Teil OC-spezifisch verhält [52].

Eine der spannenden Fragen, die sich direkt daraus ableitet, ist die Frage, inwieweit die Gewichtungen der Riechrinde bzgl. des Bulbus unseren Decision Images entsprechen. Es handelt sich bei den Decision Images schließlich genau um eine solche gewichtete Kombination der Bulbus-Aktivitäten, um diese beiden Tasks möglichst zuverlässig zu lösen.



# 7 Zusammenfassung und Ausblick

## Biologische Systeme und Maschinelernen

Im ersten Teil dieser Arbeit wurden durch eine realistische Neuronensimulation erzeugte virtuelle Elektrodenaufnahmen mittels ICA analysiert. Es wurden dabei sehr gute Ergebnisse erzielt, sowohl bei der Spike Detection als auch beim Spike Sorting. Mit alternativen Verfahren müssen typischerweise erst einmal Spikes im Elektrodensignal gefunden werden (Detection), um sie anschließend den richtigen Zellen zuweisen zu können (Sorting). Durch die ICA kehren sich diese beiden Prozesse um, denn das Entmischen der Elektrodenkanäle entspricht weitestgehend dem Spike Sorting. Da das Ergebnis einem potentiellen intrazellulären Signal entspricht, ist auch das Detektionsproblem dadurch quasi gelöst.

Obwohl sich die extrazellulären Potentiale der Zellen auf den Elektroden, nicht rein linear aufsummieren, scheint das Sorting mittels ICA grundsätzlich möglich zu sein. Ein konsequenter nächster Schritt wäre es daher, diesen Ansatz auch für die Entmischung *realer* Multielektroden-Aufnahmen einzusetzen. Die Einzelzell-Simulationen von GENESIS sind zwar sehr realistisch, dennoch sollte die Übertragbarkeit auf tatsächliche Laborumgebungen genau evaluiert werden, vor allem da speziell die Simulation der Elektrode, des Mediums und der extrazellulären Signalpropagierung stark vereinfacht werden musste.

Es hat sich weiterhin gezeigt, dass das hier verwendete Elektrodenlayout nicht optimal ist. Durch die Orientierung der Elektroden senkrecht entlang der z-Achse können zwei Zellen untrennbar für die ICA vermischt werden, sobald sie sich nur auf gleicher Höhe zur Elektrode befinden. Mit einer besseren räumlichen Abdeckung der Elektroden wäre hier ein noch besseres Ergebnis möglich gewesen. Mit einem solchen optimalen Elektrodenlayout wäre es dann potentiell sogar möglich, die Zellen im Raum zu lokalisieren. In Kombination mit einer online-ICA könnte eine komplett neuartige Neuronennavigation entworfen werden, die besonders im Bereich der Tiefenhirnstimulation in Echtzeit aus den Elektrodenaufnahmen eine Karte des Elektrodenumfeldes erzeugen und dort Position und Aktivität jeder sichtbaren Nervenzelle darstellen könnte.

### Wahrnehmungsprozesse aus Sicht des Maschinellen Lernens

Nachdem im ersten Teil dieser Arbeit eher die Lösung einer konkreten Fragestellung im Vordergrund stand, wurde in Kapitel 3 ein Analyserahmen vorgestellt, der es möglich macht, Wahrnehmungsprozesse alleine auf der Basis psychophysikalischer Daten – also verbaler Beschreibungen durch Probanden – zu analysieren.

Es wurde ein Experiment vorgestellt, in dem Probanden an einem Computer nacheinander Farben mit Beschreibern zu vergleichen hatten. Insgesamt wurden mehr als 25.000 paarweise Bewertungen zwischen 90 Farbstimuli und 16 Farbbezeichnern abgegeben. Statt die so entstandene Beschreibungsmatrix ausschließlich bezüglich der Farben und ihrer Ähnlichkeiten zu analysieren, wurde dieselbe Matrix ebenfalls dazu verwendet, eine Projektion des Wahrnehmungsraumes zu erzeugen. Dies entspricht einer Darstellung der Ähnlichkeiten zwischen den verbalen Bezeichnern.

Beide Varianten wurden mittels Multidimensional Scaling in eine niedrigdimensionale Darstellung gebracht. In beiden Fällen werden exakt die gleichen Daten analysiert, es wird jedoch sehr deutlich, dass im Falle der Interpretation auf der Karte des Merkmalsraumes *zwangsläufig* subjektiv zu erfolgen hat. Im speziellen Fall des Farbsehens kann dies aufgrund unserer gefestigten Vorstellungen und der guten Visualisierbarkeit der Stimuli sogar brauchbare Ergebnisse liefern. In der Regel ist dies aber, vor allem bei schlecht verstandenen Prozessen, kaum zu erwarten.

Hier wird allerdings aufgezeigt, dass die zweite Variante, die Analyse des Beschreibungsraumes, unserer Intuition des Wahrnehmungsraumes sehr nahe kommt. Dies ist auch nicht erstaunlich, da es sich um die *objektive* Auswertung der *subjektiven* Beschreibungen der Probanden handelt. Für Wahrnehmungsprozesse, die komplexer sind, wie zum Beispiel das Riechen oder die Attraktivität von Personen und Gegenständen, kann eine solche Analyse des Wahrnehmungsraumes potentiell hilfreiche Anhaltspunkte zum weiteren Verständnis liefern.

Einem dieser komplexeren Prozesse – nämlich dem Geruchssinn – widmete sich die zweite Hälfte dieser Arbeit. Im ersten Teil wurde eine Geruchsdatenbank, die uns Dank einer Kooperation mit dem California Institute of Technology zur Verfügung stand, unter Verwendung des vorgestellten Analyserahmens untersucht. Eine erste Abschätzung der Komplexität dieses Wahrnehmungsraumes ergab, dass das olfaktorische System nur durch einen hochdimensionalen Raum vernünftig beschrieben werden kann. Diese 32 Dimensionen, die durch die Analyse herausgearbeitet wurden, widersprechen zwar den traditionellen Ansätzen, die einen niedrigdimensionalen Sinn postulieren, unsere Ergebnisse werden aber durch jüngere Arbeiten in ihren Kernaussagen bestätigt [103].

Es wurden Selbstorganisierende Karten verwendet, um diesen hochdimensionalen

---

Zusammenhang auch auf einer zweidimensionalen Karte darzustellen. Auf diesen Karten war es nun erstmals möglich, Geruchsqualitäten zu quantifizieren. Es wurde an dem prominenten Beispiel der metabolischen Zyklen von Stickstoff und Schwefel gezeigt, wie wir bisher nur vermutete Zusammenhänge nun auf diesen Karten darstellen können. Dies beschränkt sich aber nicht auf dieses Beispiel. Mit der hier vorgestellten Technik können beliebige Stoffklassen entlang ihrer Wahrnehmungsqualitäten dargestellt werden.

## Der kombinatorische Kode des Riechens

Mit diesen Karten kann das Riechen aber nicht erschöpfend erklärt werden. Zu groß sind die Wissenslücken speziell bezüglich der hierarchischen Beziehung der Bezeichner zueinander: Es erscheint zwar intuitiv sinnvoll, *fruity* als einen übergeordneter Bezeichner zu *apple* zu vermuten, aber wie verhält sich z.B. *cherry* oder *citrus* dazu? Wie ist der Zusammenhang zwischen *fruity* und *sweet* zu beschreiben?

Es fehlt der sensorische Kontext zu den Beschreibern, eine Brücke zwischen Merkmals- und Wahrnehmungsräumen. Diese Lücke konnte – dank der Kooperation mit der UC Irvine – durch sensorische Informationen – Bilder der Rezeptoraktivität – geschlossen werden. Diese Daten wurden mit den Geruchsbeschreibern in Beziehung gesetzt und die sich daraus ergebenden Zusammenhänge untersucht. Es konnte gezeigt werden, dass mit Hilfe der in Kapitel 5 vorgestellten Decision Images die Aktivitätsbilder entlang chemischer Klasseninformation bedeutungsvoll zerlegt werden können. Diese Faktorisierung ist nicht auf chemische Klassen beschränkt und kann ebenso auch auf Geruchsklassen angewendet werden.

Auf diese Art und Weise lassen sich nicht nur Kandidaten für das „Geruchsalphabet“ ermitteln, es ist auch möglich *natürliche* Kodewörter des Riechkolbens zu visualisieren. Dies liefert ganz neue Einblicke in die *Organisation* dieses kombinatorischen Kodes, für dessen „Existenzbeweis“ 2005 der Nobelpreis für Medizin verliehen wurde<sup>1</sup>.

Da die komplette Analyse sämtlicher Bilder relativ zu allen Klassen den Rahmen dieser Arbeit gesprengt hätte, wurde diese systematische Zerlegung instruktiv für einige ausgewählte chemische Klassen und Geruchsklassen dargestellt. Die systematische Aufarbeitung des gesamten Datenmaterials ist aber ein konsequenter nächster Schritt.

Im letzten Teil dieser Arbeit wird schließlich auf synthetischen Daten demonstriert, dass durch die Decision Images bereits einige einfache Kodierungsmodelle der Geruchswahrnehmung implizit getestet wurden. Es hat sich allerdings gezeigt, dass

---

<sup>1</sup>[http://nobelprize.org/nobel\\_prizes/medicine/laureates/2004/press.html](http://nobelprize.org/nobel_prizes/medicine/laureates/2004/press.html)

keines dieser Modelle auf den Uptake-Bildern bestätigt werden kann. Dieses Erkenntnis scheint die Hypothese zu stärken, dass der Geruchswahrnehmung zwar ein kombinatorischer Kode zugrundeliegt, selbiger sich aber auf die eindeutige Kodierung jeder einzelnen Substanz – im Sinne eines Fingerabdruckes – beschränkt.

Die Decision Images geben sogar Grund zu der Annahme, dass auf der Basis dieser gut trennbaren Einzelsubstanz-Kodierung auch die beiden zentralen Aufgaben des Geruchssinns – Odorant Fine Discrimination und Odor Clustering – durch geeignete Verschaltung der Eingabe gelöst werden können. Genau eine solche Auswahl von geeigneten Gewichten hierzu liefern die Decision Images. Diese offensichtliche, funktionale Trennung von OFD- und OC-spezifischen Regionen auf den Decision Images findet ihren Sinn vermutlich in einer Inputselektion, wie sie auch im Pyriform Cortex, also der Riechrinde, stattfindet [24, 45].

Etwas formalisierter betrachtet gibt es also zwei Thesen, die sich aus dieser Arbeit bezüglich der olfaktorischen Sinneswahrnehmungen ableiten lassen:

Der kombinatorische Kode im Bulbus Olfactorius beschränkt sich auf einen eineindeutigen Fingerabdruck jedes einzelnen Riechmoleküls.

und

Geruchsqualität wird im Riechsystem unabhängig von den untersuchten chemischen Substanzeigenschaften kodiert.

Letzteres konnte im Rahmen dieser Arbeit leider nicht erschöpfend untersucht werden, es gibt jedoch bereits Verhaltensexperimente, die die These nahelegen, dass die Geruchsähnlichkeit eines Duftstoffes anhand der Aktivitätsmuster vorhergesagt werden kann [102]. Diese Kodierung scheint jedoch übergeordnet, also im Pyriform Cortex, stattzufinden, so dass sie – wie in dieser Arbeit gezeigt – nicht durch einfache und mehrfache kombinatorische Codes auf der Riechkolbenebene nachgewiesen werden können.

Vermutlich sind die Geruchsqualitäten in der Tat nicht über die chemischen Zusammenhänge definiert. Dies erscheint insofern plausibel, da ja insbesondere die Qualitäten erklärt werden sollen, die eben gerade nicht an eine chemische Eigenschaft gebunden sind. Der Pyriform Cortex scheint also seine beiden Module – ähnlich wie unsere Decision Images – derart mit dem Bulbus zu verbinden, dass der eine Teil die OC-Aufgaben, der andere die OFD-Aufgaben lösen kann. Da diese Aufgaben gegenläufig sind, verwundert es nicht, dass die dafür am besten geeigneten Glomeruli durch nahezu disjunkte Mengen beschrieben werden.

In den kommenden Jahren wird es sich zeigen, wieviele experimentelle Entsprechungen die hier vorgestellten Decision Images im realen Riechsystem haben. Die



---

Hoffnung ist, dass die gemachten Vorhersagen experimentelle Entwürfe gewinnbringend unterstützen können. Die Komplexität der Bilder verhindert, über eine große Menge von Daten geeignete Muster und Zusammenhänge zu erkennen. Wie jedoch in dieser Arbeit mehrfach demonstriert wurde, können die hier vorgestellten Decision Images genau dies leisten.



# Danksagung

Beginnen möchte ich meine Danksagungen bei Ulrich G. Hofmann, der mich nicht nur an die Geruchsforschung herangeführt hat, sondern auch die Idee zu dem ICA-Projekt hatte. Wo wäre ich jetzt ohne Dich! Auch James Bower und Christine Chee-Ruiter bin ich dankbar für alle Ideen und Einsichten, die sie mir ermöglicht haben. Mindestens ebenso danke ich Michael Leon und vor allem Brett Johnson. Beide haben viel dazu beigetragen, dass es trotz aller Hürden immer wieder voran ging.

Weiterhin bedanke ich mich im Rahmen der in dieser Arbeit vorgestellten Projekte bei Hanna Sharp für ihre Mitarbeit am ICA-Projekt, Martin Haker für seine wesentlichen Beiträge zum CPP, bei Ingrid Braenne und Elena Schuh, durch die ich überhaupt erst auf die Daten aus Irvine gestoßen bin, Sebastian Riemann und Anja Teehankee für ihre Arbeit an dem Visualisierungstool und der Datenanalyse, und im Rahmen der weiteren Arbeiten auf den Geruchsdaten bei Georg Oechsler, Martin Schröder und last but not least bei Marianne Jacob. Ohne Euch wäre in dieser Arbeit nichts so wie es jetzt ist!

Ich bedanke mich ebenfalls bei allen meinen Kolleginnen und Kollegen für die konstruktive Atmosphäre und den immer stärkeren Zusammenhalt. Vor allem bei Martin Böhme, mit dem ich nun schon seit Jahren Seite an Seite in Forschung und Lehre einen gemeinsamen Weg beschreite. Besonders im Endspurt meiner Arbeit haben mich aber auch Michael Dorr, Kai Labusch, Martin Haker, Fabian Timm und Sascha Klement weit über das selbstverständliche hinaus unterstützt. Dafür lasse ich Euch demnächst mal wieder eine Kicker-Partie gewinnen.

Mein ganz besonderer Dank gilt natürlich Thomas Martinetz, für viel Vertrauen und Freiheiten und auch für seinen Einsatz als Doktorvater und Mentor – besonders auf der Zielgeraden dieser Arbeit.

Und schließlich danke ich auch allen anderen Menschen in meinem Leben, ohne die ich jetzt nicht hier wäre.

Besonders Dir, Claudia.

Lübeck, im März 2008





# A Methoden und Algorithmen

## A.1 Multidimensional Scaling

*Multidimensional Scaling* (MDS) ist eine numerische Methode, um zu einer gegebenen Ähnlichkeitsmatrix  $\Delta \in \mathbb{R}^{n \times n}$  passende metrische Punkte  $\mathbf{x}_i, i = 1 \dots, n$  in einem  $p$ -dimensionalen Raum zu berechnen. MDS verfolgt dabei die intuitive Strategie, für jeden Punkt  $\mathbf{x}_i$  die Position im Raum zu finden, an der seine Distanzen zu den übrigen  $(n - 1)$  Punkten möglichst denen aus  $\Delta$  entsprechen. [51].

Dies wird für MDS wie folgt formuliert:

Gegeben sei also eine Ähnlichkeitsmatrix  $\Delta$  von  $n$  Objekten mit

$$\Delta = \begin{bmatrix} \delta_{11} & \cdots & \delta_{1n} \\ \vdots & & \vdots \\ \delta_{n1} & \cdots & \delta_{nn} \end{bmatrix}$$

Gesucht sind – für eine gewünschte Einbettungsdimension  $p$  – eine Menge von  $n$  Punkten  $\mathbf{x}_1, \dots, \mathbf{x}_n \in \mathbb{R}^p$ , die optimal zu der Ähnlichkeitsmatrix  $\Delta$  passen.

Wählen wir also  $n$  beliebige Punkte  $\mathbf{x}_i \in \mathbb{R}^p$  mit

$$\begin{aligned} \mathbf{x}_1 &= (x_{11} \quad \dots \quad x_{1p}) \\ &\vdots \\ \mathbf{x}_n &= (x_{n1} \quad \dots \quad x_{np}) \end{aligned}$$

Dann gibt es eine zugehörige Distanzmatrix  $\mathbf{D}$ , die die Distanzen dieser Punkte zueinander enthält.  $\mathbf{D}$  ist also definiert als

$$\mathbf{D} = \begin{bmatrix} d_{11} & \cdots & d_{1n} \\ \vdots & & \vdots \\ d_{n1} & \cdots & d_{nn} \end{bmatrix}$$

wobei die Distanzen  $d_{ij}$  sich durch eine beliebige Metrik, z.B. der euklidischen

Distanz ableiten lassen

$$d_{ij} = d_e(x_i, x_j) = \sqrt{\sum_{k=1}^p (x_{ik} - x_{jk})^2}$$

Wählen wir die  $\mathbf{x}_i$  zufällig, wird  $\mathbf{D}$  die gesuchte Matrix  $\Delta$  noch nicht entsprechen. Wir suchen nun also eine Funktion, die die Ähnlichkeiten möglichst gut auf die Distanzen projiziert, also

$$\min \sum_{i=1}^n \sum_{j=1}^p [f(\delta_{ij}) - d_{ij}]^2$$

Kruskal [51] hat eine sogenannte **stress**-Funktion wie folgt definiert

$$\mathbf{stress} = \sqrt{\frac{\sum_i \sum_j [f(\delta_{ij}) - d_{ij}]^2}{\sum_i \sum_j d_{ij}^2}}$$

Vorstellen kann man sich den *stress* als Federn, die zwischen allen Punkten gespannt sind. Wollen wir jetzt einen Punkt bewegen, so bewegen wir ihn in die Richtung, in die die Spannung von den meisten Federn am stärksten zieht. In den Kapiteln 3 und 4 haben wir diese Standardfunktion minimiert. Es gibt unterdessen einige Varianten des MDS, die sich zentral nur in der Formulierung der *stress*-Funktion unterscheiden [85].

---

## A.2 Self-Organizing Maps

*Self-Organizing Maps* (SOMs) oder auch Kohonen-Karten werden verwendet, um einen hochdimensionalen Eingaberaum  $\mathbf{X} \in \mathbb{R}^{p \times d}$  durch eine menschenlesbare Karte  $\mathcal{M}$  darzustellen. Daher sind SOMs typischerweise zwei- bis dreidimensional. Soweit dies mit der vorgegebenen Topologie möglich ist, sind die SOMs topologieerhaltend, d.h. Nachbarn auf der Karte  $\mathcal{M}$  sind auch Nachbarn im Datenraum  $\mathbf{X}$ .

Die Karte selbst besteht aus  $n$  Codebook-Vektoren  $\mathbf{m}_i \in \mathbb{R}^d, i = 1, \dots, n$ , die auf der Karte  $\mathcal{M}$  angeordnet sind. Diese Vektoren  $\mathbf{m}_i$  befinden sich als Knoten in einem festen niedrigdimensionalen Gitter und werden diese Ordnung auch während des Lernprozesses nicht mehr verlassen. Für die Topologie der Karte sind verschiedene Varianten denkbar, in dieser Arbeit z.B. liegen die Knoten in einer toroiden 2D-Anordnung. D.h. die Knoten liegen in einem regulären 2D-Gitter, in welchem auch die Knoten auf linker und rechter Seite, sowie oben und unten miteinander verbunden sind.

Das SOM wird iterativ trainiert. Zu Beginn werden die Codebook-Vektoren meist zufällig initialisiert.

Zu Beginn jeder Lernrunde  $t$  wird ein Trainingsvektor  $\mathbf{x}(t) \in \mathbf{X}$  ausgewählt. Zu  $\mathbf{x}(t)$  wird der nächstliegende Codebook-Vektor  $\mathbf{m}_c(t)$  berechnet mit

$$\|\mathbf{x} - \mathbf{m}_c\| = \min_i \|\mathbf{x} - \mathbf{m}_i\|$$

Der sog. Siegerknoten  $\mathbf{m}_c(t)$  lernt nun den Vektor  $\mathbf{x}(t)$ . Ebenso lernen auch seine topologischen Nachbarn  $\mathbf{m}_i$  auf  $\mathcal{M}$ , wobei der Lernzuwachs durch eine vordefinierte Nachbarschaftsfunktion  $h_{ci}$  gewichtet wird. Meist ist  $h_{ci}$  eine Gaußfunktion mit

$$h_{ci} = \exp\left(-\frac{\|\mathbf{r}_c - \mathbf{r}_i\|^2}{2\sigma^2}\right)$$

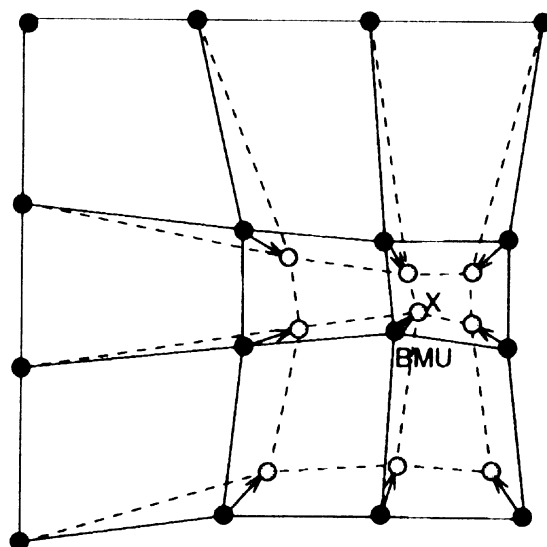
wobei  $\mathbf{r}_c, \mathbf{r}_i \in \mathbb{R}^2$  die festen Gitterkoordinaten der Knoten  $\mathbf{m}_c$  und  $\mathbf{m}_i$  sind, und  $\sigma$  die Breite der Nachbarschaft beschreibt.

Die Kohonen-Lernregel für jeden Codebook-Vektor  $\mathbf{m}_i$  lautet

$$\mathbf{m}_i(t+1) = \mathbf{m}_i(t) + \alpha(t)h_{ci}(t)[\mathbf{x}(t) - \mathbf{m}_i(t)], \quad (\text{A.1})$$

wobei  $\alpha(t)$  die Lernrate im Schritt  $t$  beschreibt.

Im Rahmen des Lernvorganges werden typischerweise Nachbarschaftsradius  $h_{ci}$  und Lernrate  $\alpha$  rundenweise verändert. Begonnen wird meist mit einer breiten Nachbarschaft und einer hohen Lernrate, um die Knoten initial an den Daten  $\mathbf{X}$  aus-



**Abbildung A.1: Competitive Learning des SOMs.** Der Trainingsvektor  $\mathbf{x}(t)$  ist durch ein Kreuz markiert. Die ausgefüllten Punkte zeigen die Codebook-Vektoren  $\mathbf{m}_i$  im Zeitpunkt  $t$ , die hohlen Punkte zeigen die neuen Positionen nach dem Lernschritt. BMU markiert den Siegerknoten  $\mathbf{m}_c$ . Die übrigen Knoten lernen in Abhängigkeit zu ihrer Nachbarschaftsdistanz zu  $\mathbf{m}_c$ . *Picture taken from [46]*

zurichten. Später, wenn sich die Karte an den feinen Strukturen ausrichten soll, müssen beide Parameter immer kleiner werden.

Am Ende jeder Runde testet der Algorithmus, ob das System bereits in einen stabilen Zustand konvergiert ist. In diesem Fall endet der Lernprozeß.



---

## A.3 DoubleMinOver Lernalgorithmus

Der *DoubleMinOver*-Lernalgorithmus (DMO) ist ein überwachtes Lernverfahren, welches ein lineares Klassifizierungsproblem löst. Er stellt eine Erweiterung des Perzeptron-Algorithmus für linear separable Punktmengen dar. Während die Perzeptron-Lernregel eine beliebige Lösung aus einer typischerweise unendlich großen Lösungsschar lernt, konvergiert der DMO-Algorithmus gegen die lineare Klassifikationsebene, die zu beiden Klassen den maximalen Abstand hat [64].

Sei also eine Menge  $\mathcal{D}$  von  $n$  gelabelten und linear separablen Datenpunkten gegeben mit

$$\mathcal{D} = \{ \mathbf{x}_i, y \}_{i=1}^n, \quad \mathbf{x}_i \in \mathbb{R}^d, y_i \in \{+1, -1\}$$

Gegeben sind weiter ein initialer (meist zufälliger) Gewichtsvektor  $\mathbf{w}_0$  und die Trainingslänge  $t_{\max}$ .

Nun werden die zugehörigen Indexmengen  $\mathbf{I}^+$  und  $\mathbf{I}^-$  der beiden Klassen gebildet mit

$$\mathbf{I}^+ = \{i \mid y_i = +1\}, \quad \mathbf{I}^- = \{i \mid y_i = -1\}$$

In jeder Iteration werden nun zunächst die beiden „ungünstigsten“ Ausgaben berechnet: Der Punkt  $\mathbf{x}_i$  aus Klasse  $+1$  mit dem kleinsten Skalarprodukt  $y_{\min}^{+1}$  zu  $\mathbf{w}$ , und der Punkt  $\mathbf{x}_i$  aus Klasse  $-1$  mit dem größten Skalarprodukt  $y_{\max}^{-1}$  zu  $\mathbf{w}$ . Somit ergibt sich

$$y_{\min}^{+1} = \min_{i \in \mathbf{I}^+} (\mathbf{w}^T \mathbf{x}_i), \quad i_{\min}^{+1} = \arg \min_{i \in \mathbf{I}^+} (\mathbf{w}^T \mathbf{x}_i)$$

und

$$y_{\max}^{-1} = \max_{i \in \mathbf{I}^-} (\mathbf{w}^T \mathbf{x}_i), \quad i_{\max}^{-1} = \arg \max_{i \in \mathbf{I}^-} (\mathbf{w}^T \mathbf{x}_i)$$

Die DMO-Lernregel für  $\mathbf{w}$  lautet

$$\mathbf{w} = \mathbf{w} + \mathbf{x}_{i_{\min}^{+1}} - \mathbf{x}_{i_{\max}^{-1}}$$

Nach  $t_{\max}$  Iterationen wird schließlich noch der Schwellwert  $\theta$  genau zwischen den beiden Support-Vektoren berechnet mit

$$\theta = \frac{1}{2} \cdot (y_{\min}^{+1} + y_{\max}^{-1})$$

## A.4 Stratified Cross-Validation

Bei der *Stratified Cross-Validation* (SCV) geht es insbesondere um den Ausgleich von Größenunterschieden der beiden zu validierenden Klassen. Daher wird im Vergleich zur einfachen Cross-Validation der Validierungsfehler  $\varepsilon_{\text{SCV}}$  getrennt auf beiden Klassen ermittelt und erst am Ende der Validierung gewichtet zusammenge-rechnet [10].

Sei also eine Menge  $\mathcal{D}$  von  $n$  gelabelten Datenpunkten gegeben mit

$$\mathcal{D} = \{ (\mathbf{x}_i, y_i) \}_{i=1}^n, \quad \mathbf{x}_i \in \mathbb{R}^d, y_i \in \{+1, -1\}$$

Ebenso gegeben ist ein Parameter  $k$ , der die gewünschte Partitionierungsgröße an-gibt. Seien die Klassen durch ihre Indexmengen  $\mathbf{I}^+$  und  $\mathbf{I}^-$  beschrieben, mit

$$\mathbf{I}^+ = \{i \mid y_i = +1\}, \quad \mathbf{I}^- = \{i \mid y_i = -1\}$$

Wähle daraus  $k$  zufällige, gleichgroße Partitionierungen mit

$$(\mathbf{I}_1^+, \mathbf{I}_2^+, \dots, \mathbf{I}_k^+) \quad \text{bzw.} \quad (\mathbf{I}_1^-, \mathbf{I}_2^-, \dots, \mathbf{I}_k^-)$$

und

$$|\mathbf{I}_1^+| = |\mathbf{I}_2^+| = \dots = |\mathbf{I}_k^+| \quad \text{bzw.} \quad |\mathbf{I}_1^-| = |\mathbf{I}_2^-| = \dots = |\mathbf{I}_k^-|$$

Seien die Punktmengen der Partitionen  $\mu, \mu = 1, \dots, k$  definiert als

$$\mathcal{D}_\mu^+ = \{ (\mathbf{x}_i, y_i) \}_{i \in \mathbf{I}_\mu^+}, \quad \mathcal{D}_\mu^- = \{ (\mathbf{x}_j, y_j) \}_{j \in \mathbf{I}_\mu^-}$$

Nun trainiere den zu testenden Klassifikator  $k$ -mal, und zwar in Runde  $t$  mit

$$\mathcal{D}_t = \mathcal{D} \setminus (\mathcal{D}_t^+ \cup \mathcal{D}_t^-)$$

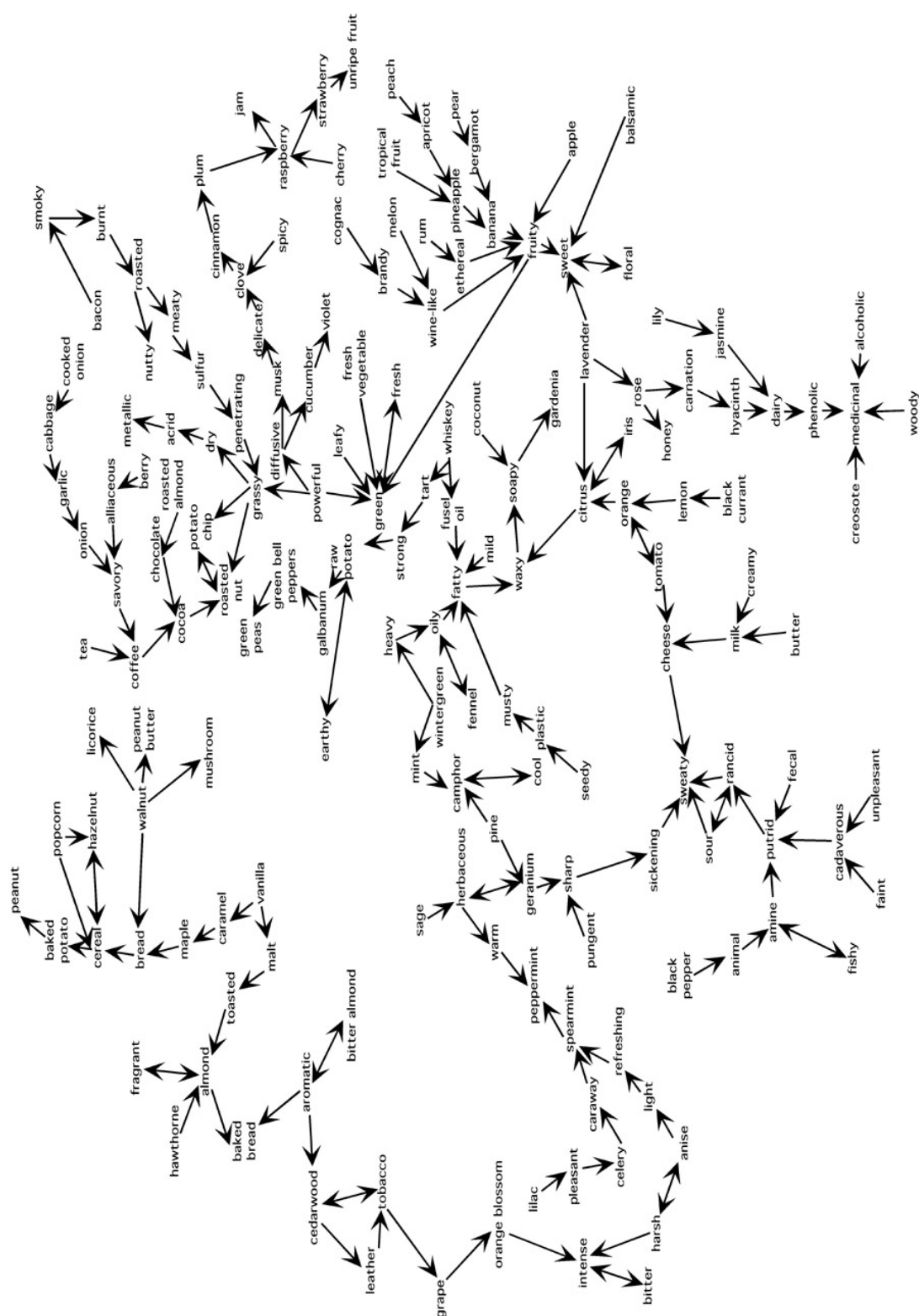
Das Ergebnis sei die Abbildung  $f_t$  mit  $f_t : \mathbb{R} \longrightarrow \{0, 1\}$ . Dann ergibt sich der balancierte Testfehler für Runde  $t$  als

$$\varepsilon_t = \frac{1}{2} \left[ \frac{1}{2|\mathbf{I}_t^+|} \sum_{i \in \mathbf{I}_t^+} |f_t(\mathbf{x}_i) - y_i| + \frac{1}{2|\mathbf{I}_t^-|} \sum_{j \in \mathbf{I}_t^-} |f_t(\mathbf{x}_j) - y_j| \right]$$

Der balancierte Validierungsfehler  $\varepsilon_{\text{SCV}}$  ergibt sich somit aus

$$\varepsilon_{\text{SCV}} = \frac{1}{k} \sum_{t=1}^k \varepsilon_t$$

## B Karten und Beschreiber



**Abbildung B.1:** Der gerichtete, ungewichtete Graph stellt Ähnlichkeiten zwischen Geruchsbeschreibern der Aldrich-Daten dar. Die Kanten verbinden jeweils einen Bezeichner mit seinem ähnlichsten Nachbarn, die Kantenlänge ist willkürlich gewählt. *Picture taken from [15].*

**Tabelle B.1: Aldrich Beschreiber für Geruchsqualität.** Geruchsbezeichner, mit denen Substanzen beschrieben wurden. Fett gedruckt sind die Bezeichner, die Substanzqualitäten in der Uptake-Datenbank beschreiben, zusätzlich kursiv sind die Bezeichner zu denen Decision Images gelernt wurden.

|                          |                       |                        |                      |                       |
|--------------------------|-----------------------|------------------------|----------------------|-----------------------|
| <i>putrid</i>            | roasted               | meaty                  | <i>burnt</i>         | <i>rancid</i>         |
| <i>pungent</i>           | <i>fatty</i>          | <i>butter</i>          | <i>cheese</i>        | <b>creamy</b>         |
| <i>oily</i>              | sour                  | <i>balsamic</i>        | anise                | (balsam)              |
| caramel                  | chocolate             | cinnamon               | honey                | <i>sweet</i>          |
| <i>vanilla</i>           | soapy                 | <i>waxy</i>            | <i>wine-like</i>     | coffee                |
| smoky                    | chemical              | <i>fruity</i>          | <i>apple</i>         | apricot               |
| <i>banana</i>            | <i>berry</i>          | cherry                 | coconut              | <i>grape</i>          |
| grapefruit               | jam                   | melon                  | peach                | <i>pear</i>           |
| <i>pineapple</i>         | plum                  | quince                 | raspberry            | <b>strawberry</b>     |
| <i>citrus</i>            | lemon                 | lime                   | orange               | <i>ethereal</i>       |
| <i>nutty</i>             | almond                | hazelnut               | peanut               | walnut                |
| <i>spicy</i>             | pepper                | medicinal              | <i>mint</i>          | <i>floral</i>         |
| blossom                  | carnation             | gardenia               | geranium             | hawthorne             |
| hyacinth                 | iris                  | jasmine                | jonquil              | lilac                 |
| lily                     | marigold              | narcissus              | <i>rose</i>          | <i>violet</i>         |
| <i>woody</i>             | <i>green</i>          | moosy                  | <i>vegetable</i>     | <i>herbaceous</i>     |
| caraway                  | sage                  | <b>earthy</b>          | <b>musty</b>         | <i>camphoraceous</i>  |
| sulfurous                | egg                   | cabbage                | metallic             | alliaceous            |
| onion                    | garlic                | animal                 | (pungent)            | <b>tart</b>           |
| leafy                    | strong                | <i>powerful</i>        | <i>fragrant</i>      | <b>aromatic</b>       |
| faint                    | popcorn               | potato chip            | toasted grain        | bread crust           |
| <i>heavy</i>             | <i>cocoa</i>          | cereal                 | bread                | odorless              |
| (anise)                  | <b>phenolic</b>       | harsh                  | bacon                | savory                |
| horseradish              | amber                 | dry                    | elegant              | incense               |
| oriental                 | egg yolk              | hard-boiled egg        | <i>penetrating</i>   | fennel                |
| mushroom                 | cadaverous            | gasoline               | <b>pleasant</b>      | <b>mild</b>           |
| bitter almond            | repulsive             | urine                  | quinoline            | rubbery               |
| fresh                    | fishy                 | <b>peppermint</b>      | cresylic             | <b>milk</b>           |
| <b>rum</b>               | <i>warm</i>           | <i>sharp</i>           | <i>sweaty</i>        | <b>spearmint</b>      |
| refreshing               | terpene               | cool                   | clove                | cassia                |
| lemon peel               | intense               | acid                   | raisin               | prune                 |
| musk                     | weak                  | <b>unpleasant</b>      | <b>baked potato</b>  | sauteed garlic        |
| clams                    | <b>orange blossom</b> | very strong            | fenu-greek           | licorice              |
| diffusive                | butyric               | roasted crude sugar    | mildew               | moldy                 |
| whiskey                  | <b>peanut butter</b>  | <b>new leather</b>     | <b>roasted nut</b>   | <b>grassy</b>         |
| <b>grilled chicken</b>   | tea                   | roasted barley         | boiled poultry       | <b>delicate</b>       |
| magnolia                 | plastic               | seedy                  | <b>light</b>         | <i>brandy</i>         |
| (sour)                   | burnt almond          | chamomile              | passion fruit        | dried fruit           |
| maple                    | butterscotch          | <b>tobacco</b>         | <b>leather</b>       | rhubarb               |
| skunk                    | <b>candy</b>          | raw potato             | <i>wintergreen</i>   | cognac                |
| mustard                  | <b>baked bread</b>    | ripe                   | <b>lavender</b>      | smoked sausage        |
| <b>toasted</b>           | <b>sickening</b>      | <b>alcoholic</b>       | (leafy)              | acrid                 |
| bitter                   | tropical fruit        | <b>unripe fruit</b>    | <b>hot sugar</b>     | <i>fecal</i>          |
| <b>fusel oil</b>         | mango                 | pine                   | turpentine           | <b>celery</b>         |
| grape skin               | green bell peppers    | green peas             | tomato leaves        | ammonia               |
| cedarwood                | blueberry             | rooty                  | creosote             | <b>clean</b>          |
| bergamot                 | <b>malt</b>           | <b>black currant</b>   | mercaptan            | galbanum              |
| roasted almond           | roasted peanut        | (gardenia)             | candy circus peanuts | dairy                 |
| buttermilk               | stinging              | <b>cucumber</b>        | watermelon           | <b>acrylic</b>        |
| (bread)                  | roasted corn          | boiled cabbage         | fried                | cooked onion          |
| cooked meat              | crackers              | wild                   | menthol              | rich                  |
| <b>brown</b>             | <b>tomato</b>         | <b>parmesan cheese</b> | <b>romano cheese</b> | <b>ricotta cheese</b> |
| green bean               | <b>sherry</b>         | amine                  | acetic               | saffron               |
| mothballs                | decayed               | bland                  | petroleum            | cauliflower           |
| <b>fermented soybean</b> | <b>lard</b>           | burnt caramel          | roasted coffee       | wet                   |
| orange peel              | <b>mandarin</b>       | flat                   |                      |                       |

**Tabelle B.2:** Alle 48 chemischen Klassen, in die die Uptake-Daten kategorisiert wurden. Fett gedruckt sind die Klassen, für deren Substanzen sowohl Uptake-Bilder wie auch Qualitätsbeschreibungen verfügbar waren, zusätzlich kursiv sind die Bezeichner zu denen Decision Images gelernt wurden.

|                                                     |                                                         |
|-----------------------------------------------------|---------------------------------------------------------|
| <i>explicit carboxylic acid</i>                     | <i>explicit and possible carboxylic acid</i>            |
| <i>possible carboxylic acid</i>                     | <i>small explicit carboxylic acid (5C or less)</i>      |
| <i>large explicit carboxylic acid (6C or more)</i>  | <i>alcohol (not phenol)</i>                             |
| <i>primary alcohol</i>                              | <i>small aliphatic primary alcohol (8C or less)</i>     |
| <i>large aliphatic primary alcohol (9C or more)</i> | <i>aliphatic (not alicyclic) secondary alcohol</i>      |
| <i>ester (not lactone)</i>                          | <i>aliphatic ester (not alicyclic)</i>                  |
| <i>small aliphatic ester (8C or less)</i>           | <i>large aliphatic ester (9C or more)</i>               |
| <i>aromatic ester</i>                               | <i>aldehyde</i>                                         |
| <i>aliphatic aldehyde</i>                           | <i>small aliphatic aldehyde (8C or less)</i>            |
| <i>large aliphatic aldehyde (9C or more)</i>        | <i>aromatic aldehyde</i>                                |
| <i>ether</i>                                        | <b>amine</b>                                            |
| <b>sulfur-containing compounds</b>                  | halogen-containing compounds                            |
| oxime                                               | <i>ketone</i>                                           |
| <i>aliphatic ketone (not alicyclic)</i>             | <i>small aliphatic ketone (8C or less)</i>              |
| <i>large aliphatic ketone (9C or more)</i>          | <i>aliphatic or alicyclic- mult.O-cont.funct.groups</i> |
| <b>aliphatic or alicyclic hydrocarbons</b>          | alkanes                                                 |
| alkenes                                             | alkynes                                                 |
| <i>aromatics</i>                                    | alkyl benzenes                                          |
| <i>aromatics with O-containing substituents</i>     | <i>phenols</i>                                          |
| <i>alicyclics</i>                                   | <i>polycyclics</i>                                      |
| <i>heterocyclics</i>                                | <i>heterocyclics with N in ring</i>                     |
| <i>pyrazines</i>                                    | <b>pyridines</b>                                        |
| heterocyclics with S in ring                        | <i>heterocyclics with O in ring</i>                     |
| <b>lactone</b>                                      | <i>furan</i>                                            |

## B.1 Übersetzungen

### B.1.1 chemische Klassen

| Englisch                                                        | Deutsch                                                                  |
|-----------------------------------------------------------------|--------------------------------------------------------------------------|
| explicit carboxylic acid                                        | eindeutige Carboxylsäure                                                 |
| explicit and possible carboxylic acid                           | eindeutige und mögliche Carboxylsäuren                                   |
| possible carboxylic acid                                        | mögliche Carboxylsäure                                                   |
| small explicit carboxylic acid (5C or less)                     | kurzkettige eindeutige Carboxylsäure (5C oder weniger)                   |
| large explicit carboxylic acid (6C or more)                     | langkettige eindeutige Carboxylsäure (6C oder mehr)                      |
| alcohol (not phenol)                                            | Alkohol (nicht Phenol)                                                   |
| primary alcohol                                                 | Primäralkohol                                                            |
| small aliphatic primary alcohol (8C or less)                    | kurzkettiger aliphatischer Primäralkohol (8C oder weniger)               |
| large aliphatic primary alcohol (9C or more)                    | langkettiger aliphatischer Primäralkohol (9C oder mehr)                  |
| aliphatic (not alicyclic) secondary alcohol                     | aliphatischer (nicht alizyklischer) Sekundäralkohol                      |
| ester (not lactone)                                             | Ester (nicht Lakton)                                                     |
| aliphatic ester (not alicyclic)                                 | aliphatischer (nicht alizyklischer) Ester                                |
| small aliphatic ester (8C or less)                              | kurzkettiger aliphatischer Ester (8C oder weniger)                       |
| large aliphatic ester (9C or more)                              | langkettiger aliphatischer Ester (9C oder mehr)                          |
| aromatic ester                                                  | aromatischer Ester                                                       |
| aldehyde                                                        | Aldehyd                                                                  |
| aliphatic aldehyde                                              | aliphatisches Aldehyd                                                    |
| small aliphatic aldehyde (8C or less)                           | kurzkettiges aliphatisches Aldehyd (8C oder weniger)                     |
| large aliphatic aldehyde (9C or more)                           | langkettiges aliphatisches Aldehyd (9C oder mehr)                        |
| aromatic aldehyde                                               | aromatisches Aldehyd                                                     |
| ether                                                           | Ester                                                                    |
| ketone                                                          | Keton                                                                    |
| aliphatic ketone (not alicyclic)                                | aliphatisches Keton                                                      |
| small aliphatic ketone (8C or less)                             | kurzkettiges aliphatisches Keton (8C oder weniger)                       |
| large aliphatic ketone (9C or more)                             | langkettiges aliphatisches Keton (9C oder mehr)                          |
| aliphatic or alicyclic- multiple O-containing functional groups | aliphatische oder alizyklische mehrere O enthaltende funktionale Gruppen |
| aromatics                                                       | Aromaten                                                                 |
| aromatics with O-containing substituents                        | Aromaten mit sauerstoffhaltigen Substituenten                            |
| phenols                                                         | Phenole                                                                  |
| alicyclics                                                      | alizyklische Verbindungen                                                |
| polycyclics                                                     | polyzyklische Verbindungen                                               |
| heterocyclics                                                   | heterozyklische Verbindungen                                             |
| heterocyclics with N in ring                                    | heterozyklische Verbindung mit N im Ring                                 |
| pyrazines                                                       | Pyrazine                                                                 |
| heterocyclics with O in ring                                    | heterozyklische Verbindung mit O im Ring                                 |
| furan                                                           | Furan                                                                    |

## B.1.2 Geruchsbeschreiber

| Englisch      | Deutsch                    |
|---------------|----------------------------|
| putrid        | verfault                   |
| burnt         | verbrannt                  |
| rancid        | ranzig                     |
| pungent       | stechend                   |
| fatty         | fettig                     |
| butter        | Butter                     |
| cheese        | Käse                       |
| oily          | ölig                       |
| balsamic      | balsamisch                 |
| sweet         | süßlich, lieblich          |
| vanilla       | Vanille                    |
| waxy          | wachsartig                 |
| wine-like     | Wein, weinartig            |
| fruity        | fruchtig                   |
| apple         | Apfel                      |
| banana        | Banane                     |
| berry         | Beeren                     |
| grape         | Trauben                    |
| pear          | Birne                      |
| pineapple     | Ananas                     |
| citrus        | zitronig                   |
| ethereal      | ätherisch, zart            |
| nutty         | nussig                     |
| spicy         | würzig                     |
| mint          | Minze                      |
| floral        | blumig                     |
| rose          | Rose                       |
| violet        | Veilchen                   |
| woody         | holzig, nach Wald riechend |
| green         | grün, sauer, herb          |
| vegetable     | Gemüse                     |
| herbaceous    | Kräuter                    |
| camphoraceous | Kampfer                    |
| powerful      | kräftig, intensiv          |
| fragrant      | duftend, wohlriechend      |
| heavy         | schwer, betäubend          |
| cocoa         | Kakao                      |
| penetrating   | durchdringend              |
| warm          | warm, mild                 |
| sharp         | scharf, beißend            |
| sweaty        | Schweiß                    |
| brandy        | Weinbrand                  |
| wintergreen   | Immergrün                  |
| fecal         | fäkal                      |



# Literaturverzeichnis

- [1] Aldrich, editor. *Flavor and Fragrances Catalog*. Sigma Aldrich Chemicals Company, Milwaukee, WI, 1996.
- [2] R. C. Araneda, A. D. Kini, and S. Firestein. The molecular receptive range of an odorant receptor. *Nat Neurosci*, 3:1248–1255, 2000.
- [3] S. Arctander. *Perfume and Flavor Chemicals (Aroma Chemicals)*. published by ed., Montclair, NJ, 1969.
- [4] R. Axel. The molecular logic of smell. *Scientific American*, 273(4):154–159, 1995.
- [5] T. Azetsu, E. Uchino, and N. Suetake. Blind separation and sound localization by using frequency-domain ica. *Soft Comput*, 11:185–192, 2007.
- [6] H. Barlow and P. Földiák. *Adaptation and decorrelation in the cortex*, pages 54–72. Addison-Wesley, Boston, 1989.
- [7] W. Bechtel and G. Graham. *A Companion to Cognitive Science*. Blackwells, 1998.
- [8] A. J. Bell and T. J. Sejnowski. An information-maximization approach to blind separation and blind deconvolution. *Neural Comput.*, 7(6):1129–1159, 1995.
- [9] J. M. Bower and D. Beeman. *The Book of GENESIS: Exploring Realistic Neural Models with the GEneral NEural SIMulation System*. Springer, New York, 2nd edition, 1998.
- [10] L. Breiman, J.H. Friedman, R.A. Olshen, and C.J. Stone. *Classification and Regression Trees*. Wadsworth International Group, 1984.
- [11] G. D. Brown, S. Yamada, and T. J. Sejnowski. Independent components analysis at the neural cocktail party. *Trends in Neuroscience*, 24(1):54–63, 2001.
- [12] L. Buck and R. Axel. A novel multigene family may encode odorant receptors: a molecular basis for odor reception. *Cell*, 65(1):175–187, 1991.
- [13] L. B. Buck. The molecular architecture of odor and pheromone sensing in mammals. *Cell*, 100(6):611–618, 2000.
- [14] W. S. Cain. History of research on smell. *E. C. Carterette and M. P. Friedman (Eds.), Handbook of Perception*, VIA: Tasting and Smelling:197–243, 1978.

- [15] C. W. J. Chee-Ruiter. *The Biological Sense of Smell: Olfactory Search Behavior and a Metabolic View for Olfactory Perception*. PhD thesis, California Institute of Technology, Pasadena, CA, 2000.
- [16] T. A. Cleland, A. Morse, E. L. Yue, and C. Linster. Behavioral models of odor similarity. *Behav Neurosci*, 2(116):222–231, 2002.
- [17] P. Comon. Independent component analysis: A new concept? *Signal Processing*, 36:287–314, 1994.
- [18] A. Dravnieks. Odor quality: Semantically generated multidimensional profiles are stable. *Science*, 218:799–801, 1982.
- [19] A. Dravnieks. *Atlas of Odor Character Profiles*. Data Series DS 61. ASTM, Philadelphia, PA, 1985.
- [20] H. Farahbod, B. A. Johnson, S. S. Minami, and M. Leon. Chemotopic representations of aromatic odorants in the rat olfactory bulb. *J Comp Neurol*, 496:350–366, 2006.
- [21] R. W. Friedrich and S. I. Korsching. Chemotopic, combinatorial, and non-combinatorial odorant representations in the olfactory bulb revealed using a voltage-sensitive axon tracer. *J Neurosci*, 18:9977–9988, 1998.
- [22] T. Furia and N. Bellanca. *Fenaroli’s Handbook of Flavor Ingredients*. CRC Press, Boca Raton, FL, 1994.
- [23] P. A. Godfrey, B. Malnic, and L. B. Buck. The mouse olfactory receptor gene family. *Proc Natl Acad Sci*, 101(7):2156–2161, 2004.
- [24] J. A. Gottfried, J. S. Winston, and R. J. Dolan. Dissociable codes of odor quality and odorant structure in human piriform cortex. *Neuron*, 49:467–479, 2006.
- [25] M. Haker. *CPP - Color Perception Project*. Study thesis, University of Luebeck, Germany, 2004.
- [26] M. Haker, A. Madany Mamlouk, and T. Martinetz. Perception space analysis: From color vision to olfaction. Poster presented at Computational Neuroscience Meeting (CNS2006) in Edinburgh, 2006.
- [27] H. Henning. *Der Geruch*. Barth, Leipzig, 1916.
- [28] S. L. Ho, B. A. Johnson, A. L. Chen, and M. Leon. Differential responses to branched and unsaturated aliphatic hydrocarbons in the rat olfactory system. *J Comp Neurol*, 499:519–532, 2006.
- [29] S. L. Ho, B. A. Johnson, and M. Leon. Long hydrocarbon chains serve as unique molecular features recognized by ventral glomeruli of the rat olfactory bulb. *J Comp Neurol*, 498:16–30, 2006.
- [30] A. Hyvärinen, J. Karhunen, and E. Oja. *Independent Component Analysis*. John Wiley & Sons, New York, 2001.
- [31] A. Hyvärinen and E. Oja. Independent component analysis: Algorithms and applications. *Neural Networks*, 13(4-5):411–430, 2000.

- [32] B. Jaehne. *Digitale Bildverarbeitung*. Springer, New Jersey, 4th edition, 1997.
- [33] A. K. Jain and R. C. Dubes. *Algorithms for clustering data*. Prentice Hall Advanced Reference Series: Computer Science, New Jersey, 1988.
- [34] B. A. Johnson, H. Farahbod, and M. Leon. Interactions between odorant functional group and hydrocarbon structure influence activity in glomerular response modules in the rat olfactory bulb. *J Comp Neurol*, 483:205–216, 2005.
- [35] B. A. Johnson, H. Farahbod, S. Saber, and M. Leon. Effects of functional group position on spatial representations of aliphatic odorants in the rat olfactory bulb. *J Comp Neurol*, 483:192–204, 2005.
- [36] B. A. Johnson, H. Farahbod, Z. Xu, S. Saber, and M. Leon. Local and global chemotopic organization: general features of the glomerular representations of aliphatic odorants differing in carbon number. *J Comp Neurol*, 480:234–249, 2004.
- [37] B. A. Johnson, S. L. Ho, Z. Xu, J. S. Yihan, S. Yip, E. E. Hingco, and M. Leon. Functional mapping of the rat olfactory bulb using diverse odorants reveals modular responses to functional groups and hydrocarbon structural features. *J Comp Neurol*, 449:180–194, 2002.
- [38] B. A. Johnson and M. Leon. Modular representations of odorants in the glomerular layer of the rat olfactory bulb and the effects of stimulus concentration. *J Comp Neurol*, 422:496–509, 2000.
- [39] B. A. Johnson and M. Leon. Odorant molecular length: One aspect of the olfactory code. *J Comp Neurol*, 426:330–338, 2000.
- [40] B. A. Johnson, J. Ong, K. Lee, S. L. Ho, S. Arguello, and M. Leon. Effects of double and triple bonds on the spatial representations of odorants in the rat olfactory bulb. *J Comp Neurol*, page in press, 2006.
- [41] B. A. Johnson, C. C. Woo, E. E. Hingco, K. L. Pham, and M. Leon. Multidimensional chemotopic responses to n-aliphatic acid odorants in the rat olfactory bulb. *J Comp Neurol*, 409:529–548, 1999.
- [42] B. A. Johnson, C. C. Woo, and M. Leon. Spatial coding of odorant features in the glomerular layer of the rat olfactory bulb. *J Comp Neurol*, 393:457–471, 1998.
- [43] B. A. Johnson, Z. Xu, J. Kwok, P. Pancoast, J. Ong, and M. Leon. Differential specificity in the glomerular response profiles for alicyclic, bicyclic and heterocyclic odorants. *J Comp Neurol*, 499:1–16, 2006.
- [44] I. T. Jolliffe. *Principal Component Analysis*. Springer-Verlag, New York, 1986.
- [45] M. Kadohisa and D. Wilson. Separate encoding of identity and similarity of complex familiar odors in piriform cortex. *Proc Natl Acad Sci*, 103:15206–15211, 2006.

- [46] S. Kaski. *Data Exploration using Self-Organizing Maps*. D.sc. (tech), Helsinki University of Technology, Finland, March 1997. Acta Polytechnica Scandinavica, Mathematics, Computing and Management in Engineering Series No. 82.
- [47] L. M. Kay. Two minds about odors. *PNAS*, 101(51):17569–17570, 2004.
- [48] C. Kennedy, M.H. Des Rosiers, O. Sakurada, M. Shinohara, M. Reivich, J.W. Jehle, and L. Sokoloff. Metabolic mapping of the primary visual system of the monkey by means of the autoradiographic [14c] deoxyglucose technique. *PNAS*, 73(11):4230–4234, 1976.
- [49] S. Klement, A. Madany Mamlouk, and T. Martinetz. Reliability of cross-validation for svms in high-dimensional, low sample size scenarios. submitted to DAGM2008.
- [50] M.L. Kringelbach, N. Jenkinson, S.L.F. Owen, and T.Z. Aziz. Translational principles of deep brain stimulation. *Nature Reviews Neuroscience*, 8:623–635, 2007.
- [51] J. B. Kruskal and M. Wish. *Multidimensional Scaling*. Sage Publications, Beverly Hills, CA, 1978.
- [52] M. Leon and B. Johnson. Functional units in the olfactory system. *Proc Natl Acad Sci*, 103:14985–14986, 2006.
- [53] M. W. Levine. *Fundamentals of Sensation and Perception*. Oxford University Press, New York, 3rd edition, 2001.
- [54] C. Linster, B. A. Johnson, E. Yue, A. Morse, Z. Xu, E. E. Hingco, Y. Choi, M. Choi, A. Messiha, and M. Leon. Perceptual correlates of neural representations evoked by odorant enantiomers. *J. Neurosci.*, 21(24):9837–9843, 2001.
- [55] F. Lorenz. *Lineare Algebra 1*. Spektrum Akademischer Verlag, 4th edition, 2003.
- [56] A. Madany Mamlouk. *Quantifying Olfactory Perception*. Diploma thesis, University of Luebeck, Germany, 2002.
- [57] A. Madany Mamlouk, C. Chee-Ruiter, U. G. Hofmann, and J. M. Bower. Quantifying olfactory perception: Mapping olfactory perception space by using multidimensional scaling and self-organizing maps. *Neurocomputing*, 52-54:591–597, 2003.
- [58] A. Madany Mamlouk and T. Martinetz. On the dimensions of the olfactory perception space. *Neurocomputing*, 58-60:1019–1025, 2004.
- [59] A. Madany Mamlouk, H. Sharp, K. M. L. Menne, U. G. Hofmann, and T. Martinetz. Unsupervised spike sorting with ica and its evaluation using genesis simulations. *Neurocomputing*, 65-66:275–282, 2005.
- [60] A. Madany Mamlouk, A. Teehanke, E. Schuh, T. Martinetz, M. Leon, and B. A. Johnson. Predictions of human odor descriptors and odorant chemical

- classes from glomerular response patterns in the rat olfactory bulb, 2006. Abstr. ECRO Meeting 17, page 234.
- [61] B. Malnic, P. A. Godfrey, and L. B. Buck. The human olfactory receptor gene family. *Proc Natl Acad Sci*, 101(18):2584–2589, 2004.
- [62] B. Malnic, J. Hirono, T. Sato, and L. B. Buck. Combinatorial receptor codes for odors. *Cell*, 96:713–723, 1999.
- [63] T. Martinetz. Maxminover: A simple incremental learning procedure for support-vector-classification. *Proc of IJCNN 2004*, pages 2065–2070, 2004.
- [64] T. Martinetz, K. Labusch, and D. Schneegaß. SoftDoubleMinOver: A Simple Procedure for Maximum Margin Classification. *ICANN*, pages 301–306, 2005.
- [65] K. M. L. Menne, A. Folkers, R. Maex, T. Malina, and U. G. Hofmann. Test of spike sorting algorithms on the basis of simulated data. *Neurocomputing*, 44-46:1119–1126, 2002.
- [66] K. M. L. Menne, T. Malina, A. Folkers, and U. G. Hofmann. Biologically realistic simulation of a part of hippocampal ca3: Generation of testdata for the evaluation of spike detection algorithms. *5th GWAL*, pages 17–25, 2002.
- [67] K. M. L. Menne, C. K. E. Moll, S. Kondra, M. Bär, P. Detemple, A. K. Engel, and U. G. Hofmann. On-line 32 channel signal processing and integrated database improve navigation during cranial stereotactic surgeries. *International Congress Series*, 1268:443–448, 2004.
- [68] P. Mombaerts, F. Wang, C. Dulac, S. K. Chao, A. Nemes, M. Mendelsohn, J. Edmondson, and R. Axel. Visualizing an olfactory sensory map. *Cell*, 87(4):675–686, 1996.
- [69] K. Mori. Grouping of odorant receptors: odour maps in the mammalian olfactory bulb. *Biochem Soc Trans.*, 31(1):134–136, 2003.
- [70] K. Mori, H. Nagao, and Y. Yoshihara. The olfactory bulb: Coding and processing of odor molecule information. *Science*, 286:711–715, 1999.
- [71] K. Mori, Y. K. Takahashi, K. M. Igarashi, and M. Yamaguchi. Maps of odorant molecular features in the mammalian olfactory bulb. *Physiol Rev.*, 86(2):409–433, 2006.
- [72] J. Ngai, M. M. Dowling, L. Buck, R. Axel, and A. Chess. The family of genes encoding odorant receptors in the channel catfish. *Cell*, 72(5):657–666, 1993.
- [73] G. Oechsler. *Kodierungsmodelle der Geruchswahrnehmung*. Master thesis, University of Luebeck, Germany, 2006.
- [74] J. Parkinson. An essay on the shaking palsy. 1817. (reproduced). *J Neuropsychiatry Clin Neurosci*, 14(2):223–236, 2002.
- [75] K. J. Ressler, S. L. Sullivan, and L. B. Buck. A zonal organization of odorant receptor gene expression in the olfactory epithelium. *Cell*, 73:597–609, 1993.

- [76] K. J. Ressler, S. L. Sullivan, and L. B. Buck. Information coding in the olfactory system: evidence for a stereotyped and highly organized epitope map in the olfactory bulb. *Cell*, 79:1245–1255, 1994.
- [77] H. Ritter, T. Martinetz, and K. Schulten. *Neural Computation and Self-Organizing Maps: An Introduction*. Addison-Wesley, Massachusetts, 1992.
- [78] S. Rösch. Die Kennzeichnung der Farben. *Physikalische Zeitschrift*, 29:83–91, 1928.
- [79] K. J. Rossiter. Structure-odor relationships. *Chem. Rev.*, pages 3201–3240, 1996.
- [80] Harvey Richard Schiffman. *Sensation and Perception - An Integrated Approach*. John Wiley & Sons, Inc., New York, 4th edition, 1996.
- [81] S. S. Schiffman. Contributions to the physicochemical dimensions of odor: A psychophysical approach. *Ann N Y Acad Sci.*, 237:164–183, 1974.
- [82] H. Schuhart, K.M.L. Menne, and U.G. Hofmann. Grogra and genesis: in computo grown neurons used for realistic compartmental modeling. In: Poster and Abstract, CNS\*2002, Chicago.
- [83] C. E. Shannon. Programming a computer for playing chess. *Philosophical Magazine*, 41 (7), 1950.
- [84] F. R. Sharp, J. S. Kauer, and G. M. Shepherd. Laminar analysis of 2-deoxyglucose uptake in olfactory bulb and olfactory cortex of rabbit and rat. *J Neurophysiol*, 40:800–813, 1977.
- [85] R. N. Shepard, A. K. Romney, and S. B. Nerlove. *Multidimensional Scaling: Theory and Applications in the Behavioural Sciences, Volume I – Theory*. Seminar Press, New York, 1972.
- [86] E. Sivan and N. Kopel. Mechanism and circuitry for clustering and fine discrimination of odors in insects. *Proc Natl Acad Sci USA*, 101(51):17861–6, 2004.
- [87] W.B. Stewart, J.S. Kauer, and G.M. Shepherd. Functional organization of rat olfactory bulb analyzed by the 2-deoxyglucose method. *J. Comp. Neurol.*, 185(4):715–734, 1979.
- [88] Y. K. Takahashi, M. Kurosaki, S. Hirono, and K. Mori. Topographic representation of odorant molecular features in the rat olfactory bulb. *J Neurophysiol*, 92:2413–2427, 2004.
- [89] Holm Tetens. *Geist, Gehirn, Maschine - Philosophische Versuche über ihren Zusammenhang*. Reclam, 1994.
- [90] R.D. Traub, J.G. Jefferys, M.A. Whittington, and M. Toth. A branching dendritic model of a rodent ca3 pyramidal neurone. *J Physiol*, 481:79–95, 1994.

- [91] R.D. Traub and R. Miles. Pyramidal cell-to-inhibitory cell spike transduction explicable by active dendritic conductances in inhibitory cell. *J Comput Neurosci*, 2:291 – 298, 1995.
- [92] Alan Turing. Computing machinery and intelligence. *Mind*, LIX(236):433–460, 1950.
- [93] N. Uchida, Y. K. Takahashi, M. Tanifuji, and K. Mori. Odor maps in the mammalian olfactory bulb: domain organization and odorant structural features. *Nat Neurosci.*, 3(10):1035–1043, 2000.
- [94] V. Vapnik. *Statistical Learning Theory*. Wiley, New York, 1998.
- [95] J. Vesanto, J. Himberg, E. Alhoniemi, and J. Parhankangas. SOM Toolbox for Matlab 5. Report A57, Helsinki University of Technology, Neural Networks Research Centre, Espoo, Finland, 2000.
- [96] A. Walthelm and A. Madany Mamlouk. Multisensoric Active Spatial Environment Exploration and Modeling. *Dynamische Perzeption*, pages 141–146, 2000.
- [97] D. L. Wells and P. G. Hepper. The discrimination of dog odours by humans. *Perception*, 29:111–115, 2000.
- [98] F Wichmann, A Graf, E P Simoncelli, H Bülthoff, and B Schölkopf. Machine learning applied to perception: Decision images for gender classification. *Advances in Neural Information Processing Systems*, 17:1489–1496, 2005.
- [99] F. Wood, M. J. Black, C. Vargas-Irwin, M. Fellows, and J. P. Donoghue. On the variability of manual spike sorting. *IEEE TBME*, 51(6):912–918, 2004.
- [100] M. H. Woskow. *Multidimensional Scaling of Odors*. PhD thesis, University of California Los Angeles, Los Angeles, CA, 1964.
- [101] Y.Gilad, C.D.Bustamante, D.Lancet, and S.Paabo. Natural selection on the olfactory receptor gene family in humans and chimpanzees. *Am J Hum Genet.*, 73(3):489–501, 2003.
- [102] S. L. Youngentob, B. A. Johnson, M. Leon, P. R. Sheehe, and P. F. Kent. Predicting odorant quality perceptions from multidimensional scaling of olfactory bulb glomerular activity patterns. *Behav Neurosci*, page in press, 2006.
- [103] M. Zarzo and D. T. Stanton. Identification of latent variables in a semantic odor profile database using principal component analysis. *Chem Senses*, 31:713–724, 2006.
- [104] Z. Zou, F. Li, and L. B. Buck. Odor maps in the olfactory cortex. *PNAS*, 102(21):7724–7729, 2005.





# Lebenslauf

**Dipl.-Inf. Amir Madany Mamlouk**

## **Persönliche Daten:**

Name: Madany Mamlouk  
Vorname: Amir  
Geburtsdatum: 12. März 1975  
Geburtsort: Berlin  
Staatsangehörigkeit: deutsch  
  
Adresse: Engelsgrube 1-17  
23552 Lübeck  
Tel.: 0451 / 612 95 41  
Email: amir@madany.de

## **Bildungsweg:**

1981 – 1985 Christian-Morgenstern Grundschule, Berlin - Spandau  
1985 – 1987 Grundschule am Amalienhof, Berlin - Spandau  
09/1987 – 07/1994 Kant-Gymnasium, Berlin-Spandau  
Abschluss: Allgemeine Hochschulreife  
07/1994 – 07/1995 Grundwehrdienst  
10/1995 – 09/1996 Studium an der Humboldt-Universität zu Berlin  
im Fach Mathematik  
10/1996 – 07/2002 Studium an der Universität zu Lübeck im Fach Informatik  
mit dem Nebenfach Bioinformatik  
Abschluss: Diplom-Informatiker  
Diplomprüfung im Juli 2002, Note: „mit Auszeichnung“  
07/2002 – heute Wissenschaftlicher Mitarbeiter am Institut für Neuro- und  
Bioinformatik an der Universität zu Lübeck



# Eigene Publikationen

A. Madany Mamlouk, C. Chee-Ruiter, U. G. Hofmann, and J. M. Bower. Quantifying olfactory perception: Mapping olfactory perception space by using multidimensional scaling and self-organizing maps. *Neurocomputing*, 52-54:591–597, 2003.

A. Madany Mamlouk, J. T. Kim, E. Barth, M. Brauckmann, and T. Martinetz. One-class classification with subgaussians. *DAGM 2003, Lecture Notes Comput. Sci.*, 2781:346–353, 2003.

A. Madany Mamlouk and T. Martinetz. On the dimensions of the olfactory perception space. *Neurocomputing*, 58-60:1019–1025, 2004.

A. Madany Mamlouk, H. Sharp, K. M. L. Menne, U. G. Hofmann, and T. Martinetz. Unsupervised spike sorting with ica and its evaluation using genesis simulations. *Neurocomputing*, 65-66:275–282, 2005.

T. Martinetz, A. Madany Mamlouk, and C. Mota. Fast and Easy Computation of Approximate Smallest Enclosing Balls. *Proc. SIBGRAPI*, pages 163–170, 2006.

S. Klement, A. Madany Mamlouk, and T. Martinetz. Reliability of cross-validation for SVMs in high-dimensional, low sample size scenarios. accepted at ICANN2008.