

Birgit Alber
birgit.alber@univr.it

Einführung in die Phonologie des Deutschen¹
(2. Fassung, Wintersemester 2005/2006)

1. Einleitung

Die Phonologie beschäftigt sich mit der lautlichen Struktur der Sprache. Bevor wir jedoch damit beginnen, diese Struktur kennenzulernen, wollen wir uns zuerst ein paar allgemeinen Fragen zuwenden.

Wenn kleine Kinder sprechen lernen, dann gibt es etwa ab dem 2. Monat ein Stadium, in dem wir Erwachsenen zwar hören, dass das Kind anfängt, vor sich hinzuplappern, es fällt uns aber schwer, in dem Geplappere einzelne Sprachlaute auszumachen. Es gurr und gurgelt, und schmatzt und schnalzt, und ab und zu glaubt man auch, einzelne Laute erkennen zu können, aber in dieser ersten Lallphase produziert das Kind eher noch Geräusche als Laute. Ein paar Monate nach dieser Phase, etwa ab dem 6. Monat, beginnt das Kind nun Lautsequenzen von sich zu geben, die wie *amma, mamma, papa* ... klingen. Das Kind will damit noch nicht seine Mutter oder seinen Vater rufen, [a] [m] und [p] gehören einfach zu den Lauten, die ein Kind als erstes aussprechen lernt. Es ist daher kein Zufall, dass die Namen, mit denen Kinder ihre Eltern bezeichnen, in vielen Sprachen der Welt genau diese Laute enthalten. Hier ist eine Liste von sehr verschiedenen Sprachen und das Wort, mit dem Kinder dieser Sprache ihre Mutter rufen:

- (1) Deutsch: Mama, Mami
Italienisch: mamma
Englisch: mum, mummy
Polnisch: mama
Estnisch: mamma
Swahili: mama
Maori: māmā (ausgesprochen: maama)

Das Kind hat in dieser Phase damit begonnen, Laute voneinander zu unterscheiden. Es wird nun im Laufe seiner Entwicklung immer neue Laute dazulernen, bis es schließlich alle Laute seiner Sprache kennt und verwenden kann.

Warum ist aber die menschliche Sprache so beschaffen, dass wir Ketten von einzelnen Lauten bilden? Warum bleiben wir nicht bei der ersten Lallphase und verwenden als Bezeichnung für BAUM ein kräftiges Schmatzen? Im Grunde ist es doch viel anstrengender, sich Laut für Laut anzueignen. Immerhin brauchen kleine Kinder mehrere Monate, um ihre Zunge so weit in den Griff zu bekommen, dass sie Silbenketten bilden können.

Der Grund für die Verwendung von Lautsequenzen, im Gegensatz zu Geräuschen, liegt darin, dass auch die lautliche Struktur der Sprache, wie Sprache überhaupt, auf **diskreten Einheiten** beruht. So wie wir Wörter aus Morphemen zusammensetzen

¹ Wer dieses Material im Unterricht verwenden will, muss die Quelle angeben (z.B. Als: Birgit Alber (2005), *Einführung in die Phonologie des Deutschen*, ms. Università di Verona.)

und Sätze aus syntaktischen Phrasen bauen, so bestehen auch lautliche Strukturen aus klar voneinander unterscheidbaren Einheiten.

Wie in der Morphologie und Syntax, so beschäftigen sich auch Phonetik und Phonologie damit, wie diese Einheiten – Phone und Phoneme – beschaffen sind und nach welchen Regeln sie sich miteinander verbinden. Im 2. Kapitel dieser Einführung werden wir uns vor allem mit den Einheiten der Phonologie beschäftigen. Wir werden sehen, wie Laute klassifiziert werden und wie man Phone und Phoneme definieren kann.

Übung 1

Aus wievielen Lauten bestehen die folgenden Wörter? In welche Fällen entsprechen mehrere Buchstaben einem Laut? In welchen Fällen entsprechen mehrere Laute einem Buchstaben?

Schale
Zopf
Buch
Hexe
Mitte
Schuhe

Aber warum sind diskrete Einheiten so wichtig für die menschliche Sprache? Wenn menschliche Sprache nicht über diskrete Einheiten verfügen würde, dann müssten wir uns für jeden Begriff ein neues Geräusch merken. Wenn wir an die Einträge denken, die sogar ein kleines Wörterbuch hat, dann kommt da eine ganz schöne Anzahl von Geräuschen zusammen. Wenn wir unsere Wörter hingegen aus Lautsequenzen bilden, dann entlasten wir unser Gedächtnis um einiges. Eine Sprache wie Deutsch verfügt über etwa 40 Phoneme. Mit diesen Lauten können wir eine ziemlich große, im Prinzip unendliche Anzahl von Wörtern bilden, wir benötigen dazu aber nur eine endliche Anzahl von Lauteinheiten (nämlich 40 Phoneme).

Die Artikulation von Sprachlauten braucht große Präzision von Seiten der Sprechorgane und unser Ohr muss außerdem lernen, die einzelnen Laute voneinander zu unterscheiden. Die Hauptarbeit leistet aber eigentlich unser Gehirn. Denn obwohl wir oft undeutlich sprechen, obwohl oft ein Laut mit dem anderen verschmilzt, obwohl im Hintergrund oft viele Geräusche stören – dennoch gelingt es uns immer wieder, zu verstehen, was unsere Kommunikationspartner gesagt haben, welche Laute sie aneinandergereiht haben. Unsere Gehirn scheint irgendwie dafür bestimmt zu sein, ein Kontinuum von Lauten in Einzellaute zu unterteilen. Wir scheinen eine Prädisposition für die Analyse des Lautstroms in seine Einzelteile zu haben. Diese Tatsache zeigt uns auch, dass wir es bei der Zusammensetzung aus diskreten Einheiten mit einer sehr tiefen, universellen Charakteristik der menschlichen Sprache zu tun haben.

Wir haben bis jetzt nur von Lauten als Einheiten der phonologischen Struktur gesprochen. Es gibt aber auch andere Einheiten der lautlichen Struktur der Sprache, bei denen es sich nicht unbedingt um Einzellaute handelt. Die Rede ist hier von "größeren" Einheiten wie der Silbe, die ihrerseits aus Einzellauten besteht. Diese Einheiten werden im vierten Kapitel dieser Einführung ("Prosodische Phonologie") behandelt.

Wie in Syntax und Morphologie, so unterliegt auch die Aneinanderreihung von lautlichen Einheiten bestimmten **Regeln**. Nicht jede Sequenz von Lauten ist wohlgeformt. Eine Sequenz wie *sparf* könnte es im Deutschen geben, eine Sequenz wie *fpars* hingegen nicht. Dazu kommt noch, dass bestimmte Laute in gewissen Positionen **phonologischen Prozessen** unterliegen. Betrachten wir die folgenden Beispiele:²

- (2) Das Wort <Räte> wird ausgesprochen als *Rä[t]e*
 Das Wort <Räder> wird ausgesprochen als *Rä[d]er*

Das Wort <Rat> wird ausgesprochen als *Ra[tʰ]*
 Das Wort <Rad> wird ausgesprochen als *Ra[tʰ]* !!

Wir sehen, obwohl der letzte Buchstabe des Wortes *Rad* ein <d> ist, wird der letzte Laut des Wortes als ein stimmloses, aspiriertes [tʰ] ausgesprochen. Und das, obwohl wir im Plural von *Rad*, in *Räder*, ein stimmhaftes [d] haben. Die Aussprache von *Rad* ist also genau dieselbe wie die von *Rat*. Die Aussprache der beiden Wörter unterscheidet sich nur im Plural, nicht aber im Singular. Der Grund dafür ist ein phonologischer Prozess des Deutschen, den man **Auslautverhärtung** nennt. Wir werden uns diesen Prozess noch im Detail anschauen, im Augenblick soll nur gesagt werden, dass er dazu führen kann, dass stimmhafte Konsonanten am Ende des Wortes stimmlos werden.

Übung 2

Wie werden wohl die folgenden Wörter des Deutschen ausgesprochen werden, wenn man den Prozess der Auslautverhärtung berücksichtigt?

Tod
 Lob
 aktiv
 Kind
 lieb
 sag!
 Job

Phonologische Prozesse sind den Sprechern nicht bewusst. Sie werden von ihnen automatisch angewandt, und es ist für die Sprecher fast unmöglich, sie zu unterdrücken, wenn sie z.B. eine Fremdsprache sprechen. So ist es für deutsche Sprecher sehr schwierig, die Regel der Auslautverhärtung *nicht* anzuwenden, wenn sie z.B. Englisch sprechen. Da wird dann aus einem *jo[b]* leicht ein *Jo[pʰ]*.

Interessanterweise haben die Laute, die von der Auslautverhärtung betroffen sind ([b, d, g, v ...], alle etwas gemeinsam. Es handelt sich hier immer um Laute, bei denen der Luftstrom, der bei ihrer Produktion entsteht, behindert wird. Entweder wird er ganz

² Grapheme (Buchstaben) schreibt man zwischen spitzen Klammern (<...>), Phone werden zwischen eckige Klammern gesetzt ([...]). Phoneme (zu dem Unterschied zwischen Phonen und Phonemen kommen wir noch!) schreibt man zwischen Schrägstrichen (/.../).

blockiert, wie bei [b, d, g] oder er wird so stark blockiert, dass ein Reibegeräusch entsteht (wie bei [v]). Diese Klasse von Lauten nennt man **Obstruenten**.

In jeder Sprache gibt es phonologische Prozesse und wir werden uns im 3. Kapitel mit den phonologischen Prozessen des Deutschen befassen.

Eine Frage, mit der sich Phonologen intensiv beschäftigen ist die nach den Charakteristiken von phonologischen Prozessen und Regeln. Wir werden uns im 4. Kapitel vor allem mit zwei Aspekten beschäftigen.

Erstens wollen wir uns in diesem Einführungstext immer wieder fragen, was an den phonologischen Prozessen und Regeln, die wir untersuchen, **sprachspezifisch** ist und was hingegen **universell** ist.

So wissen wir z.B., dass die Auslautverhärtung ein typischer Prozess des Deutschen ist, der in anderen Sprachen der Welt nicht auftaucht. Andererseits hat dieser Prozess auch universelle Charakteristiken. Er betrifft eine Klasse von Lauten (die Obstruenten) und es scheint in allen Sprachen der Welt so zu sein, dass Prozesse immer Klassen von Lauten betreffen und nicht eine Liste von zufällig zusammengewürfelten Einzellauten.

Ein zweiter Aspekt von universellen Regeln und Prinzipien, den wir vor allem im 4. Kapitel diskutieren werden, betrifft ihre **Verletzbarkeit**. Gelten universelle phonologische Prinzipien immer, in allen Sprachen der Welt? Wenn ein Prinzip universell ist, sollte es seiner Definition nach eigentlich immer und überall gelten. Seit etwa zehn Jahren gibt es in der Phonologie jedoch eine Theorie, die sogenannte **Optimalitätstheorie**, die einen ziemlich radikalen Standpunkt vertritt. Sie behauptet, dass alle phonologischen Bedingungen universelle Gültigkeit haben. Wie sollen wir dann aber die Unterschiede zwischen den einzelnen Sprachen erklären? Die Hypothese der Optimalitätstheorie ist, dass alle phonologischen Bedingungen universell sind, dass diese Bedingungen jedoch in bestimmten Sprachen, in bestimmten Kontexten verletzt werden können. Ausserdem verwenden Sprachen verschiedene Strategien, um diese Bedingungen zu erfüllen.

Schauen wir uns ein Beispiel an, damit dieser Vorschlag etwas klarer wird. Es gibt eine universelle Bedingung, die besagt, dass eine Silbe mit einem Konsonanten beginnen sollte. Die Details dieser Bedingung werden wir noch im 4. Kapitel genauer untersuchen, hier formulieren wir diese Bedingung einfach folgendermaßen:

- (3) Silbenanlautbedingung:
Silben müssen mit einem Konsonanten beginnen

Wir sehen, dass diese Bedingung wahr sein muss, wenn wir uns überlegen, dass das Wort *Amerika* wohl immer als *A.me.ri.ka* silbifiziert werden wird und nie als *Am.er.ik.a*.

Das Deutsche verfolgt eine besondere Strategie, um diese Bedingung zu befolgen. Wenn Wörter mit einem Vokal beginnen, dann wird vor diesen Vokal ein Konsonant eingefügt. Bei diesem Konsonanten handelt es sich um einen **glottalen Plosiv**. Das ist ein Konsonant, der gebildet wird, indem man die Öffnung zwischen den Stimmbändern kurz verschließt, und dann die Luft auf einmal herausströmen läßt. Versucht, leicht zu husten, dann produziert ihr genau diesen Laut. Das phonetische Symbol für den glottalen Plosiv ist [ʔ]. Man hört diesen Laut sehr deutlich in einem Satz wie dem folgenden, bei dem jedes Wort mit einem Vokal beginnt:

- (4) [ʔ]Anton [ʔ]aß [ʔ]am [ʔ]Abend [ʔ]immer [ʔ]Auflauf

Wird die Silbenanlautbedingung im Deutschen immer befolgt? Nein, es gibt Fälle, in denen sie verletzt wird. Das geschieht zum Beispiel, wenn sich die vokalinitiale Silbe im Wortinneren befindet und nicht betont ist. Ein Beispiel dafür ist das Verb *sehen*. Die phonetische Transkription zeigt uns die Aussprache dieses Wortes:³

- (5) [ze:.ən] <sehen>

Bei dem ersten Laut handelt es sich um ein stimmhaftes [z]. Dann folgt ein langes [e:] und anschließend ein weiterer, e-ähnlicher Laut, ein sogenanntes **schwa**. Beachtet, dass das <h> nicht ausgesprochen wird. Es dient nur dazu, uns zu zeigen, dass der vorangehende Vokal lang ist.

Die zweite Silbe dieses Wortes beginnt auch mit einem Vokal. Trotzdem wird hier kein glottaler Plosiv eingefügt. Die Silbenanlautbedingung kann also im Deutschen nur in bestimmten Fällen befolgt werden.

Im Italienischen ist die Sache wieder etwas anders. Auf den ersten Blick sieht es so aus, als würde die Silbenanlautbedingung im Italienischen keine besonders große Rolle spielen. Schließlich gibt es viele Wörter, die mit einem Vokal beginnen, wie *amore*, *agguato*, *o*, *e*, *inverno* u.s.w. Vor diese Wörter werden keine Konsonanten gesetzt, die ersten Silben beginnen also mit einem Vokal. Trotzdem sieht man den Einfluss der Silbenanlautbedingung auch im Italienischen. Wenn diese Wörter nämlich nach Wörtern mit einem finalen Konsonanten erscheinen, dann kann man folgendes beobachten:

- (6) co.n a.mo.re
del.l' in.ver.no
i.n ag.gua.to

Lest diese Beispiele langsam, wie bei einem Abzählreim (*in una conta*). Ihr werdet beobachten, dass die letzten Konsonanten des ersten Wortes mit dem ersten Vokal des zweiten Wortes silbifiziert werden. Das Italienische verwendet also eine andere Strategie als das Deutsche, um Silben mit einem Anfangskonsonanten zu versehen: es "zieht" den Konsonanten eines vorangehenden Wortes in die vokalinitiale Silbe hinüber. Durch diese Strategie kann es sogar zu Fällen von Ambiguität kommen, wie z.B. in den folgenden Sätzen:

- (7) Ambiguität, die sich durch Resilbifizierung ergibt:
- a. La spiaggia greca aveva u.n' a.re.na strepitosa
a'. La spiaggia greca aveva u.n a.re.na strepitosa
- b. Il veterinario provò con u.n' a.ve.na diversa
b'. Il veterinario provò con u.na ve.na diversa

Wenn wir diese Sätze nur hören, aber nicht lesen, können wir nicht unterscheiden, ob es sich um die Bedeutung in a. oder a', b. oder b' handelt. Die Sätze klingen gleich, da das Italienische die Grenze zwischen *un* und dem darauffolgenden, vokalinitialen Wort nicht markiert, sondern Resilbifizierung des letzten Konsonanten von *un* erlaubt.

³ Die Silbengrenze bezeichne ich mit einem Punkt "."

Wir haben in diesen Beispielen drei Dinge gesehen. Erstens scheint es in der Tat eine universelle Silbenanlautbedingung zu geben, deren Einfluss sich sowohl im Deutschen als auch im Italienischen bemerkbar macht. Diese Bedingung ist jedoch verletzbar: eine "schöne" Silbe sollte mit einem Konsonanten beginnen, aber im Deutschen ist das im Inneren des Wortes nicht immer möglich und im Italienischen ist das auch am Anfang des Wortes nicht immer möglich. Beide Sprachen versuchen jedoch – soweit möglich – die Silbenanlautbedingung zu befolgen. Dazu verfolgen Deutsch und Italienisch unterschiedliche Strategien: im Deutschen wird ein Konsonant eingesetzt, im Italienischen wird über Wortgrenzen hinweg silbifiziert.

Diese drei Charakteristiken von phonologischen Bedingungen (Universalität, Verletzbarkeit, sprachspezifische Strategien, um sie zu erfüllen) werden wir im Kapitel über die Silbenstruktur noch genauer diskutieren.

Fassen wir nun kurz zusammen, was bisher gesagt wurde. Wir haben in dieser Einleitung einige Charakteristiken des Lautsystems der menschlichen Sprache gesehen, mit denen wir uns hier befassen werden:

- das Lautsystem der menschlichen Sprache besteht aus diskreten Einheiten. In Kapitel 2 werden wir sehen, wie diese Laut-Einheiten aussehen, welche Arten von Lauten es gibt und welche Laute im Deutschen verwendet werden. Wir werden in diesem Kapitel auch die Unterschiede zwischen den Einheiten **Phon** und **Phonem** diskutieren.
- das Lautsystem der menschlichen Sprache unterliegt regelhaften Prozessen. In Kapitel 3 werden wir uns vor allem mit den phonologischen Prozessen beschäftigen, die wir im Deutschen finden.
- das Lautsystem jeder menschlichen Sprache hat sowohl universelle als auch sprachspezifische Charakteristiken. Mit diesem Aspekt, und der Behandlung dieses Aspektes im Rahmen der Optimalitätstheorie werden wir uns in Kapitel 4 beschäftigen.

2. Diskrete Einheiten in der Phonologie

2.1 Phone – Exkurs zur Phonetik

Bevor wir die Einheiten vorstellen, mit denen sich die Phonologie befasst, müssen wir uns zuerst ein wenig mit den physikalischen Eigenschaften von Lauten im Allgemeinen beschäftigen.

Mit den physikalischen, materiellen Charakteristiken der Sprachlaute beschäftigt sich das linguistische Teilgebiet der **Phonetik**. Die Phonetik untersucht die akustischen Eigenschaften von Sprachlauten (**akustische Phonetik**), sie untersucht, wie Sprachlaute gehört werden (**auditive Phonetik**) und schließlich wie Laute von den Sprechorganen produziert werden (**artikulatorische Phonetik**).

Im Gegensatz zur Phonetik beschäftigt sich die **Phonologie** mit der (bedeutungsunterscheidenden) Funktion, die Laute in einer Sprache haben. Sie fragt danach, wie die bedeutungsunterscheidenden Einheiten (die **Phoneme**, s. nächstes Kapitel) in einer Sprache organisiert sind, wieviele Phoneme es in einer Sprache gibt,

oder auch welche Phoneminventare es im Allgemeinen in den Sprachen der Welt gibt und welche nicht.

Wir werden uns im nächsten Abschnitt vor allem mit artikulatorischer Phonetik beschäftigen, da die Laute in der Phonologie zum größten Teil nach artikulatorischen Aspekten klassifiziert werden. Es wird also vor allem darum gehen, wie unsere Sprechorgane die einzelnen Laute produzieren.

Zur Illustration, wie die akustische Phonetik arbeitet, wollen wir uns jedoch ein Spektrogramm ansehen und verstehen, wie ein Laut physikalisch gesehen aussieht. Spektrogramme sind Darstellungen des physikalischen Sprachlautes.

Sprachlaute sind, wie andere Geräusche auch, Schallwellen, die gemessen werden können. Die Frequenz einer Welle gibt die Tonhöhe an, die Amplitude die Lautstärke. Sprachlaute - vor allem Vokale - sind nicht einfache, sondern komplexe periodische Wellen, d.h., sie bestehen aus mehreren überlappenden periodischen Wellen. Viele Laute haben deshalb mehrere Frequenzen. Vokale sind z.B. durch ein ganz bestimmtes Frequenzspektrum charakterisiert.

Wie sieht nun so ein Spektrogramm aus? Die x-Achse stellt die Zeit dar, die y-Achse die Frequenz der akustischen Welle. Die Intensität der Schwärzung gibt die Lautstärke des Lautes an.

Lo spettrogramma che dovrebbe apparire su questa pagina lo trovate sulla homepage di Peter Ladefoged

<http://hctv.humnet.ucla.edu/departments/linguistics/VowelsandConsonants/course/chapter8/8.7.htm>

oppure, in forma cartacea, alla fotocopisteria Rapida

Das Spektrogramm von Peter Ladefoged ist die Aufzeichnung der englischen Wörter *a bab*, *a dad* und *a gag*. Es zeigt uns einige interessante Dinge des physikalischen Sprachlautes:

- Vokale sind durch mehrere schwarze "Ringe" gekennzeichnet. Diese Ringe sind die charakteristischen Frequenzen des betreffenden Vokals. Die Frequenzen eines spezifischen Vokals werden auch als seine Formantenstruktur bezeichnet.
- Konsonanten wie [b], [d] und [g] sind gekennzeichnet durch - Stille! Wenn ihr genau hinschaut, so erscheint über den Konsonanten fast gar keine Schwärzung.
- worin unterscheiden sich dann die einzelnen Konsonanten? Wenn wir z.B. das erste [b] von *a bab* mit dem ersten [g] von *a gag* vergleichen, dann sehen wir, dass der Unterschied vor allem in der Formantenstruktur des vorhergehenden Vokals besteht: die zweiten und dritten Formanten des schwa vor dem [g] bilden eine Art von Kurve, die sie vor [b] nicht bilden.

[b], [d] und [g] sind **Plosive** (sie werden auch **Verschlusslaute** genannt). Sie werden gebildet, indem die Mundhöhle an einer bestimmten Stelle geschlossen wird. Im Mund wird so ein (Luft)druck aufgebaut, der dann plötzlich gelöst wird, sobald die Blockierung gelöst wird. Beim Plosiv [b] zum Beispiel wird der Luftstrom bei den Lippen blockiert. [b], [d] und [g] sind stimmhafte Laute, bei ihrer Produktion vibrieren die Stimmbänder. Diese Vibration kann man auch ganz leicht auf den Spektrogrammen sehen. Bei der Produktion der stimmlosen Plosive [p], [t] und [k] hingegen herrscht vollkommene Stille. Andere Konsonanten, wie z.B. [f] und [s] produzieren Geräusche, die man auf dem Spektrogramm sehen kann.

In so kurzen Sequenzen wie auf dem vorliegenden Spektrogramm lassen sich einzelne Segmente noch ziemlich gut voneinander unterscheiden. Wenn wir jedoch ganze Sätze aufzeichnen, dann ist es oft schwierig zu verstehen, wo ein Laut aufhört und wo der andere anfängt. Der Grund dafür ist, dass die Lautsequenz physikalisch gesehen ein Kontinuum bildet, ein Laut ist nicht immer klar von dem anderen abgegrenzt. Dennoch haben wir die Intuition, dass ein Wort aus einzelnen, klar voneinander unterscheidbaren Lauten besteht.

2.1.1 Artikulatorische Phonetik – Die Sprechwerkzeuge

In der Phonologie werden Laute traditionell nach artikulatorischen Gesichtspunkten klassifiziert. Wie werden nun aber Sprachlaute produziert?

Schauen wir uns zuerst die Sprechwerkzeuge an, jene Organe, mit denen wir die Laute unserer Sprache produzieren:

- die Lunge: bei der Erzeugung von Lauten funktioniert die Lunge wie eine Pumpe, die den Luftstrom liefert. In den meisten Sprachen der Welt werden Laute mit einem Luftstrom gebildet, der von der Lunge, über die Mundhöhle nach außen dringt (egressiver, pulmonarer Luftstrom). Es gibt aber auch in einigen Sprachen Laute, die anders gebildet werden. Ein Beispiel dafür sind die Klick-Laute in bestimmten afrikanischen Sprachen (z.B. Xhosa und Zulu). Einige dieser Laute ähneln dem Zungenschmalzen eines Kutschers, der seine Pferde antreibt, oder dem Schmatzen, wenn man sich auf Entfernung ein Küsschen zuwirft.⁴

⁴ Auf der homepage des Phonetikers Peter Ladefoged kann man sich diese *clicks* anhören und sogar ein Röntgenbild von einem Sprecher sehen, der einen Klick produziert:

- Lippen (*labbra*)
- Zähne (*denti*)
- Alveolen (*alveoli*)
- Palatum (harter Gaumen) (*palato duro*)
- Velum (weicher Gaumen) (*palato molle*)
- Uvula (Zäpfchen) (*ugola*)
- die Zunge:
 - Apex (Zungenspitze) (*apice*)
 - Dorsum (Zungenrücken) (*dorso*)
 - Zungenwurzel (*radice*)
- Glottis (*glottide*): ist eigentlich kein Organ, sondern der Raum zwischen den Stimmbändern
 - sie ist offen → bei stimmlosen Lauten
 - sie ist geschlossen → beim glottalen Plosiv [ʔ]
 - sie ist geschlossen, aber die Stimmbänder vibrieren → bei stimmhaften Lauten

Teilweise Stimmhaftigkeit

Oft sind Laute nicht durchgehend stimmhaft bzw. stimmlos. Im Deutschen z.B.: bei einem "b" am Wortanfang fangen die Stimmbänder erst gegen Ende hin an zu vibrieren. Das erweckt dann bei Nicht-Muttersprachlern oft den Eindruck, man habe "p" gesagt.
z.B. in "Baum", "bis", "Bar"

Aspiration:

nach der Explosion des Konsonanten bleiben die Stimmbänder noch für eine kurze Zeit offen. Luft fließt ohne Behinderung durch die Glottis. Erst danach setzt wieder Stimmhaftigkeit (des darauffolgenden Vokals) ein.

Deutsche Konsonanten am Wortanfang unterscheiden sich also nicht so sehr in der Stimmhaftigkeit, sondern eher als teilweise stimmhaft vs. aspiriert

z.B. Paar vs. Bar = [p^h]aar vs. [b]ar

Erster Hinweis auf die Phonologie! →

Die Aspiration ist im Deutschen nicht **distinktiv**, d.h. man kann keine zwei Wörter finden, die sich nur durch die Aspiration eines Konsonanten unterscheiden, z.B.

*	[p ^h]irk =	langweilige Linguistikvorlesungen
*	[p]irk =	Blätter der Birke

Es gibt aber Sprachen, wo die Aspiration **distinktiv** ist, d.h., sie kann dazu verwendet werden, Wörter voneinander zu unterscheiden. z.B. Hindi:

[p]al	=	'sich um jemanden kümmern'
[p ^h]al	=	'Messerklänge'

2.1.2 Die Klassifikation der Konsonanten

Konsonanten werden allgemein nach ihrem Artikulationsort und nach ihrer Artikulationsart klassifiziert. Außerdem sagt man bei Plosiven und Frikativen noch dazu, ob sie stimmhaft oder stimmlos sind.

z.B. [p] ist ein	stimmloser	bilabialer	Verschlußlaut
		↓	↓
		Artikulationsort	Artikulationsart

Artikulationsart:

Es gibt verschiedene Arten, einen Konsonanten zu produzieren: man kann den Luftstrom dabei an irgendeiner Stelle im Mund komplett blockieren, teilweise blockieren, kaum blockieren. Man kann die Luft durch den Mund entweichen lassen, oder durch die Nase.

Plosive (auch Verschlusslaute genannt, engl.: *stops* oder *plosives*, it.: *occlusive*): zwei Sprechorgane blockieren den Luftstrom gänzlich. z.B.: [p], [t], [k] etc.

Frikative (engl.: *fricatives*, it.: *fricative*): der Luftstrom wird teilweise blockiert, dabei entsteht ein reibendes Geräusch. z.B. [s], [f], [x] etc.

Nasale (engl.: *nasals*, it.: *nasali*): der Luftstrom wird in der Mundhöhle blockiert, aber die Luft kann durch die Nase entweichen (bei gesenktem Velum). z.B. [m], [n], [ŋ]

Laterale (engl.: *laterals*, it.: *lateral*): der Luftstrom wird in der Mundhöhle blockiert, aber die Luft kann seitlich entweichen. z.B. [l]

Vibranten (engl.: *trills* it: *vibranti*): ein bewegliches Sprechorgan vibriert: z.B. die Zungenspitze oder die Uvula. z.B. [r], [R]

Approximanten (engl.: *approximants*, it.: *aprossimanti*): der Luftstrom muss wie bei den Frikativen durch einen sehr engen Kanal, aber es entsteht kein Reibegeräusch. z.B. [j]

Konsonanten werden im IPA mit Hilfe von drei Merkmalen klassifiziert:
Stimmhaftigkeit, Artikulationsort, Artikulationsart:
z.B.: [p]: ein stimmloser, bilabialer Plosiv

Im Netz hat Peter Ladefoged ein IPA zur Verfügung gestellt, bei dem man die einzelnen Laute anklicken und dann ihre Aussprache hören kann: ihr findet das "sprechende" IPA unter
<http://hctv.humnet.ucla.edu/departments/linguistics/VowelsandConsonants/>

Schaut dort nach unter: → Vowels and Consonants
→ 1. Sounds and languages -The IPA chart sounds

IPA

Ihr könnt die Tabelle des Internationalen Phonetischen Alphabets an folgenden Stellen finden:

- im Internet, auf der Web-Seite der *International Phonetic Association*:
<http://www2.arts.gla.ac.uk/IPA/images/ipachart.gif>
- in jedem Einführungsbuch zur Phonologie oder Linguistik (z.B. Hall 2000, Nespor 1999)
- in der *copisteria Rapida*, wo es eine Papierfassung dieser *dispensa* gibt.

2.1.3 Klassifikation der Vokale

Das **Vokaltrapez** illustriert den Artikulationsort der Vokale.

Sie werden klassifiziert:

- nach der Höhe (**hoch, tief, mittel**)
- nach der horizontalen Position der Zunge (**vorne, hinten, zentral**)

Im Deutschen sind noch zwei weitere Kategorien wichtig:

- ob ein Vokal **gerundet** oder **ungerundet** ist: in den Sprachen der Welt sind im allgemeinen hintere Vokale gerundet und vordere Vokale ungerundet, Im Deutschen gibt es aber auch die eher ungewöhnlichen gerundeten Vordervokale [y, ʏ, ø, œ]
- ob ein Vokal **gespannt** oder **ungespannt** ist:

gespannt: [i, y, e, ø, u, o]

ungespannt: [ɪ, ʏ, ɛ, œ, ɔ]

Gespannte Vokale sind "im allgemeinen" lang, ungespannte kurz. Es gibt aber auch kurze gespannte Vokale, z.B. in den unbetonten Silben von Fremdwörtern, wie in Ph[i]losophie. Und es gibt einen langen, ungespannten Vokal, das [ɛ:], wie in "Ähre" = [ɛ:]. Vokallänge wird durch einen Doppelpunkt ausgedrückt: [i:]

Der Status von [a] im Deutschen ist kontrovers. Manche Wissenschaftler nehmen an, dass es auch beim [a] eine Unterscheidung zwischen kurzem ungespanntem [a] (wie in "Mann") und langem, gespanntem [ɑ:] gibt. Wir werden mit Wiese (1996) annehmen, dass es im Deutschen nur eine Unterscheidung zwischen langem [ɑ:] und kurzem [a] gibt, dass sich aber die beiden Laute nicht in ihrer Vokalqualität unterscheiden

Es gibt im Deutschen auch zwei zentrale Vokale:

- das sogenannte "schwa" = [ə], z.B. in Lieb[ə]
- das "vokalisierte r" = [ɐ], z.B. in wi[ɐ]

Beispiele zu den Vokalen

	vorne	zentral	hinten
hoch	gesp.: M[i:]te ungesp.: M[ɪ]tte gerundet gesp.: H[y:]te gerundet ungesp. H[ʏ]tte		gesp.: sp[u:]ken ungesp.: sp[ʊ]cken
mittel	gesp.: B[e:]t ungesp.: B[ɛ]tt [ɐ:]re gerundet gesp.: H[ø:]le gerundet ungesp.: H[œ]lle	schwa: lach[ə]n vok. "r": nu[ɐ]	gesp.: Sch[o:]ten ungesp.: Sch[ɔ]tten
tief		B[a:]n, B[a]nn	

2.2 Phoneme und ihre Realisierung

Das Phonem ist ein sehr wichtiger Begriff in der Phonologie. Vorneweg ein paar Phonemdefinitionen:

"Laute als sprachliche Einheiten, mit einer distinktiven Funktion in einem bestimmten Sprachsystem nennt man Phoneme" (Vater 1993: 44)

Trubetzkoy/Jakobson:

" ... das Phonem als die minimale bedeutungsdifferenzierende Lauteinheit"

Das Phonem ist eine abstrakte Einheit, nicht etwas Konkretes. Wir können keine Phoneme aussprechen, nur **Phone**. Die Phone sind die Realisierungen eines Phonems. Phoneme werden zwischen Schrägstrichen "/.../" geschrieben, Phone zwischen eckigen Klammern "[...]". Orthographische Zeichen schreiben wir zwischen spitzen Klammern "<...>".

Nehmen wir z.B. die beiden Phoneme /f/ und /p/ im Deutschen. /f/ und /p/ sind im Deutschen Phoneme, da sie die Fähigkeit haben, zwei Wörter mit unterschiedlicher Bedeutung voneinander zu unterscheiden, wie z.B. die beiden Wörter

(8) /f/ein ~ /p/ein <fein>, <Pein>

Es gibt Sprachen, wie das Koreanische, in denen /f/ und /p/ keine Wörter unterscheiden können, also nicht Phoneme sind.

/b/ und /p/ sind im Deutschen auch Phoneme, aber hier ist die Sache etwas komplizierter. /b/ und /p/ können auch Wörter voneinander unterscheiden, wie im folgenden Beispiel:

(9) Ra/p/e ~ Ra/b/e⁵ <Rappe> (*morello*), <Rabe>, (*corvo*)

Aber die Phoneme /b/ und /p/ werden im Deutschen oft sehr unterschiedlich realisiert, je nachdem, ob sie am Anfang des Wortes stehen, oder am Ende, oder zwischen zwei Vokalen usw. D.h., die Phone, die diese beiden Phoneme realisieren, ändern sich je nach Kontext:

- Im Wortinnern: <Rappe> und <Rabe>

Ra/p/e	Ra/b/e	← phonologische Ebene/Ebene der Phoneme
↓	↓	
Ra[p]e	Ra[b]e	← phonetische Ebene/Ebene der Phone

/p/ wird als [p] realisiert

/b/ wird als [b] realisiert

- am Wortanfang: <Pein> und <Bein>

/p/ein	/b/ein	← phonologische Ebene/Ebene der Phoneme
↓	↓	
[p ^h]ein	[b ^h]ein	← phonetische Ebene/Ebene der Phone

/p/ wird als aspiriertes [p^h] realisiert

/b/ wird als teilhaft stimmhaftes [b^h] realisiert

Präzisierung des Phonembegriffs:

Was bedeutet "distinktive" (bedeutungsunterscheidende) Funktion?

- Aspiration ist im Deutschen nicht distinktiv:

d.h. es gibt keine Wörter, die sich nur durch Aspiration unterscheiden

- im Hindi ist Aspiration distinktiv:

[p]al = 'sich um etwas kümmern'

[p^h]al = 'Messerklänge'

d.h., im Hindi ist das *Merkmal* [± aspiriert] distinktiv

Besser, als von einem "Phonem" zu sprechen ist es also, von **distinktiven Merkmalen** zu sprechen. Was wir so Phonem /p/ nennen ist eigentlich nur eine Menge von Merkmalen, die distinktiv sind.

Deshalb:

Jakobson & Halle (1956: 20) (aus: Ramers & Vater 1992:30):

"The distinctive features are aligned into simultaneous bundles called phonemes"

d.h.: Phoneme sind Bündel von distinktiven Merkmalen

Jakobson und Halle meinen damit, dass sich in jedem Phonem ein Bündel von Eigenschaften vereinigt, die, jede für sich, in der betreffenden Sprache Bedeutungen unterscheiden können. Nehmen wir z.B. das Phonem /p^h/ im Hindi. Es besteht aus einer ganzen Reihe von distinktiven Merkmalen, das, jedes für sich, Bedeutungen

⁵ Diese beiden Wörter unterscheiden sich aber auch durch die Länge der Wurzelvokale: das [a] in *Rappe* ist kurz, in *Rabe* ist es lang.

unterscheiden kann, z.B. [- stimmhaft] unterscheidet /p^h/ von /b^h/ und je nachdem, ob ein Wort /p^h/ oder /b^h/ enthält, kann sich seine Bedeutung ändern. Das Merkmal [+labial] wiederum unterscheidet /p^h/ von /t^h/, das nicht labial sondern dental ist. [+aspiriert] unterscheidet /p^h/ von /p/. Auch dieses Merkmal kann, wie wir gesehen haben, Bedeutungen unterscheiden.

[+aspiriert]	—	}	/p ^h /
[- stimmhaft]	—		
[+ labial]	—		
....	—		
....	—		

Stellt euch ein Phonem als eine Schachtel vor, in die alle distinktive Merkmale hineinkommen. Auf der Schachtel steht dann vielleicht /p/ und drinnen sind so Merkmale wie [-stimmhaft] (um dieses Phonem von /b/ zu unterscheiden). Wenn wir vom Deutschen sprechen, dann ist in der Schachtel kein Merkmal [-aspiriert], da dieses Merkmal im Deutschen nicht distinktiv ist. Im Hindi hingegen haben wir zwei Schachteln, eine für das Phonem /p/ und eine für das Phonem /p^h/. Und in der ersten Schachtel muss das distinktive Merkmal [-aspiriert] drinnen sein, in der zweiten das distinktive Merkmal [+aspiriert].

Ein großer Teil der Arbeit eines Phonologen besteht darin, festzustellen, welche Merkmale in einer bestimmten Sprache distinktiv (bedeutungsunterscheidend) sind und welche nicht. Dazu wird oft die Technik der Minimalpaare verwendet. Wenn man zwei Wörter mit verschiedener Bedeutung findet, die sich nur in einem Merkmal unterscheiden, dann muss dieses Merkmal distinktiv sein.

z.B. <Todes> und <totes> unterscheiden sich nur dadurch, dass wir im ersten Fall einen stimmhaften Konsonanten zwischen den beiden Vokalen haben, im zweiten Fall aber einen stimmlosen. Also muss [±stimmhaft] im Deutschen ein distinktives Merkmal sein.

In den nächsten Abschnitten werden wir versuchen, festzustellen, welche Phoneme es im Deutschen gibt, und wie diese jeweils realisiert werden.

2.3 Allophone

Es gibt auch Laute, die im Deutschen distinktiv sind, aber nicht in anderen Sprachen.

Die folgenden Beispiele illustrieren die Distribution der Laute [l] und [r] im Koreanischen:

Am Ende einer Silbe: immer [l]		Am Anfang einer Silbe: immer [r]	
mul	‘Wasser’	ru.pi	‘Rubin’
mul.ka.ma	‘ein Platz für Wasser’	ra.tio	‘Radio’
mal	‘Pferd’	ma.re	‘am Pferd’
mal.ka.ma	‘ein Platz für ein Pferd’	mu.re	‘am Wasser’
səul	‘Seoul’	pa.ri	‘vom Fuß’
il.kop	‘sieben’		
pal	‘Fuß’		

Welche Schwierigkeiten hat ein Koreaner bei der Aussprache des Wortes *Leere*?

Im Koreanischen sind [l] und [r] Allophone, d.h. sie sind kontextbedingte Realisierungen eines Phonems.

Im Deutschen sind z.B. [ç] und [x] Allophone eines Phonems.

~~~~~

### **Übung zum Begriff Allophon:**

Versucht, anhand von folgenden Beispielen, die Distribution des palatalen Frikativs [ç] (*ich-Laut*) und des velaren Frikativs [x] (*ach-Laut*) zu beschreiben.

1. Schritt: schreibt neben jedes Beispiel, ob es ein [ç] oder ein [x] enthält
2. Schritt: macht eine Liste der Laute, die vor [ç] auftreten und eine Liste der Laute, die vor [x] auftreten. Gebt auch an, ob sich vor den beiden Frikativen jeweils eine Morphemgrenze befindet oder nicht
3. Schritt: Beschreibt die Distribution von [ç] oder ein [x]:

[ç] tritt in folgendem Kontext auf: \_\_\_\_\_  
 [x] tritt in folgendem Kontext auf: \_\_\_\_\_

|           |                         |
|-----------|-------------------------|
| dich      | manch                   |
| Tuch      | Kuchen                  |
| Tücher    | Kuhchen (kleine Kuh)    |
| Loch      | fauchen                 |
| löchern   | Pfauchen (kleiner Pfau) |
| Dach      | Chemie                  |
| Tschechow | China                   |
| Milch     |                         |

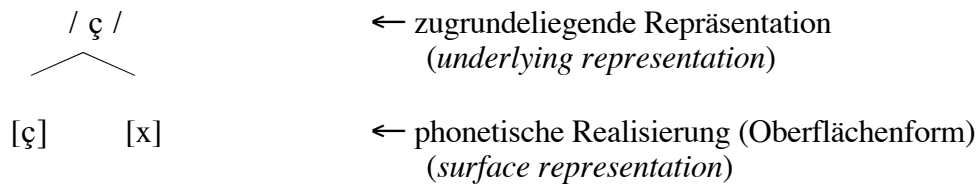
~~~~~

Ein Problem, das sich Phonologen oft stellen: wenn wir wissen, dass zwei Phone wie z.B. [ç] und [x] Allophone desselben Phonems sind, welches von den beiden ist dann das Phonem? /ç/ oder /x/?

Wurzel (1981) (s. Ramers & Vater 1992: 98):

[ç] ist zugrundeliegend, weil es die "freiere Lautvariante" ist, es steht in vielen unterschiedlichen Lautkontexten, [x] nur nach Intervokalen. (N.B.: es gibt aber auch gute Gegenargumente gegen diese Hypothese (s. Ramers & Vater 1992, Wiese 1996))

Wie schon bei den Allomorphen können wir von einer "zugrundeliegenden Form" sprechen, auf der die Phoneme dargestellt werden, und von einer "Oberflächenform", auf der die Realisierungen dieser Phoneme repräsentiert werden:



Man sagt: [ç] und [x] sind Allophone desselben Phonems /ç/, weil ihr Vorkommen kontextbedingt ist. Im Kontext "nacht Intervokalen" kommt [x] vor, in allen anderen Kontexten kommt [ç] vor.

2.4 Andere Arten von phonetischer Realisierung

Allophone haben eine komplementäre Distribution, aber es gibt auch andere Arten von phonetischer Realisierung:

- **freie Variation:** bei freier Variation kann man nicht sagen, dass in einem Kontext immer ein bestimmtes Phon vorkommt, in einem anderen aber ein anderes. Bei der freien Variation sagt der eine Sprecher mal dies, der andere das, oder ein und derselbe Sprecher benutzt einmal das eine Phon und dann wieder das andere. Im Deutschen gibt es folgende Fälle von freier Variation:

A. Aspiration am Wortende:
 Lo[p] oder Lo[p^h] ?

Manche Sprecher aspirieren am Wortende, andere nicht, und wahrscheinlich macht auch ein und derselbe Sprecher es einmal so und dann wieder anders.

B. vielleicht verschiedene Realisierungen des /r/ im Deutschen (aber nicht die Realisierung als [ʁ], die ist eindeutig durch den Kontext bestimmt)!

/r/ hat verschiedene phonetische Realisierungen:

/r/ → [r], [R], [ʀ]

z.T. ist die Variation regionaler Art, aber z.T. scheint sie auch vollkommen frei zu sein.

- dialektale Variation: in Bayern finden wir in einigen Gegenden [r]
 - deutsche Bühnenaussprache der dreißiger/vierziger Jahre: ein Schauspieler sollte damals möglichst viel das [r] benutzen. Hört mal genau hin, wenn Zarah Leander oder Lotte Lenya singt, oder bei Hitlerreden.

2.5 Die phonologische Regel

Eine phonologische Regel beschreibt phonologische Prozesse, z.B. (aber nicht nur) die phonetische Realisierung eines bestimmten Phonems. Sie hat die allgemeine Form:

A → B / C _ D

d.h.: "A" wird zu "B" im Kontext C _ D

z.B.: das Phonem /A/ wird nach C und vor D phonetisch als [B] realisiert

3. Phonologische Prozesse und Regeln

3.1 Assimilation der dorsalen Frikative

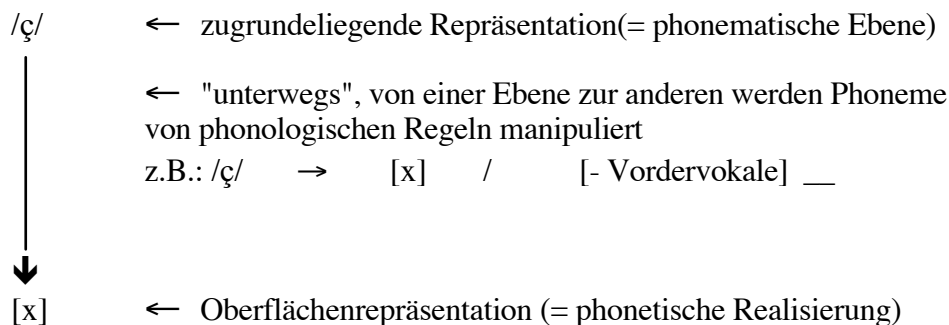
Die Distribution von [ç] und [x] kann man als **Assimilationsprozess** verstehen: der Frikativ wird normalerweise palatal (also weiter vorne als [x]) realisiert, aber wenn ein Vokal vor ihm steht, der weiter hinten ausgesprochen wird, dann gleicht sich der Frikativ an diese Aussprache an und wird auch weiter hinten (also velar) ausgesprochen:

/ç/ → [x] / [- Vordervokale] __

Preisfrage: muss man (mit Hilfe einer Regel) auch beschreiben, wann /ç/ als [ç] realisiert wird?

Nein: /ç/ → [ç] ist der *elsewhere-case*: man kann einfach sagen: "in allen anderen Fällen wird /ç/ als [ç] realisiert."

Noch einmal ein Exkurs zum Phonembegriff:



Die Ebene der Schrift: hat mit diesen beiden Ebenen direkt gar nichts zu tun. Wir schreiben meistens <ch> für diesen Laut, aber wir könnten morgen auch alle beschließen, ein Blümchen dafür zu malen. An der Sprache ändert das nichts.

Warum zwei Ebenen? So können wir erklären, was ein Sprecher des Deutschen über diese Laute *weiß*, was er in seinem Kopf gespeichert hat. Er weiß,

- dass das Wort für <Tuch> aus den Phonemen /t/, /u/ und /ç/ besteht.
- Und er hat eine Regel gespeichert, die ihm sagt, dass /ç/ in diesem Fall als [x] realisiert wird.

Aber er muss nicht wissen, dass das Wort <Tuch> ein [x] enthält, das Wort <Milch> aber [ç]. Das ergibt sich aus der phonologischen Regel.

Eine Theorie, die mit (zugrundeliegenden) Phonemen und Phonen (an der Oberfläche) arbeitet, kann erklären, warum die Struktur der Sprache ziemlich ökonomisch ist. Wir müssen uns von einem Wort nicht alle phonetischen Details merken, sondern nur die Phoneme, aus denen es gebildet ist - und die phonologischen Regeln, die für die ganze Sprache gelten. Wir müssen uns also gar nicht so viel merken - das ist der Grund, warum wir Sprache schnell lernen können und in unserem Hirn Platz für so viele Wörter haben.

3.2 Vokalisierung des /r/

Übung zur Vokalisierung des /r/

Versucht, anhand von folgenden Beispielen, die Distribution (= Verteilung) von [R] bzw. [ʀ] versus [ɐ] zu bestimmen. Der Unterschied zwischen [R] und [ʀ] ist nicht so wichtig, nur der Unterschied zwischen diesen beiden auf der einen Seite und [ɐ] auf der anderen.

a) Um welche Laute handelt es sich hier? Beschreibt kurz, wo im Mund und wie diese drei Laute gebildet werden (bei den Konsonanten genügt auch der IPA-Name):
[R] ist

[ʀ] ist ...

[ɐ] ist ...

b) Schreibt neben jedes der folgenden Wörter, ob das /r/, das es enthält, als [R] bzw. [ʀ] oder aber als [ɐ] ausgesprochen wird.

c) Beschreibt in Worten (kurz!) wann das /r/ vokalisiert wird (also wann es als [ɐ] realisiert wird). Falls ihr nicht herauskriegt, wann das passiert, dann schaut euch noch einmal genau das Beispiel zum Koreanischen an.

d) Konstruiert eine phonologische Regel (Regel der r-Vokalisierung), die das ausdrückt, was ihr in c) mit eigenen Worten beschrieben habt.

Tür	bitter	rein
Tor	Lehrer	Rolle
mehr	weiter	fahren
ver- (z.B. vergessen)	Wetter	mehrere
zer- (z.B. zerstreuen)		Leiterin
Herr	Wirt	
werden	fährt	

3.3 Nasalassimilation

Übung zur Nasalassimilation:

a) Was ist der (artikulatorische) Unterschied zwischen [n] und [ɲ]? Wie heißen die beiden Laute nach dem IPA?

b) beschreibt in Worten, wann /n/ als [n] realisiert wird und wann als [ɲ].

c) beschreibt dasselbe mit einer phonologischen Regel: wann wird das /n/ zum [ɲ]?

d) um welche Art von Prozess handelt es sich hier? Welchen ähnlichen Prozess gibt es sonst noch im Deutschen?

Ha[n]d	Ba[ŋ]k
ga[n]z	fi[ŋ]gieren
Ki[n]d	Fi[ŋ]k
Wa[n]d	de[ŋ]ken
Ha[n]s	Ta[ŋ]go
Wu[n]sch	U[ŋ]gar

~~~~~

Das ist ein Fall von **regressiver Assimilation** (ein Segment beeinflusst ein vorhergehendes Segment, d.h., Assimilation "rückwärts"), aber im Deutschen gibt es auch Fälle von **progressiver Assimilation** (Assimilation "vorwärts"), z.B. in der umgangssprachlichen Aussprache von

<geben> = [ge:bm̩]

In vielen Fällen kann ein /n/ vor [g] zuerst zu [ŋ] werden und **danach** fällt das [g], das den Assimilierungsprozess ausgelöst hat, weg. Z.B.:

si/n/g → zugrundeliegende Form

si[ŋ]g → Zwischenresultat nach Anwendung der Nasalassimilierungsregel

si[ŋ] → Resultat nach Anwendung der g- **Tilgungsregel**  
= wirkliche Aussprache

d.h.: Regeln können untereinander eine bestimmte Ordnung haben.

N.B.: die genaue Formulierung der g-Tilgungsregel ersparen wir uns hier, sie ist zu kompliziert (aber s. Wiese 1996:224). Grob gesagt fällt das [g] weg:

- vor einem schwa [ə] oder [ɐ]

- am Ende eines Morphems, aber nur, wenn danach kein betonter Vokal kommt

<Finger> → [fɪŋɐ]

<Bengel> → [bɛŋəl]

<länglich> → [lɛŋlɪç]

<Dinge> → [dɪŋə]

aber nicht: <fingieren> → [fɪŋɡɪ:rən]

### 3.4 Aspiration der Plosive

**Stark:** - Wortanfang, vor betontem Vokal  
(aber nur, wenn der Plosiv alleine im Onset (= Silbenanfang) steht)

z.B. <tot> [tʰ]ot

**Schwächer:** - vor unbetontem Vokal?

z.B. <Talent> [tʰ]alént

- am Wortende?

z.B. <tot>                      to[t<sup>h</sup>]

keine Aspiration:

- wenn der Plosiv nicht alleine im Onset steht

z.B. <Stab>                      S[t]ab

z.B. <Traum>                    [t]raum

### 3.5 Auslautverhärtung

~~~~~  
Übung zur Auslautverhärtung:

Was passiert mit Plosiven und Frikativen in den folgenden Beispielen ?

	Aussprache:	Graphie:
Lo[]es	Lo[]	<Lobes, Lob>
Ra[]es	Ra[]	<Rades, Rad>
Ta[]es	Ta[]	<Tages, Tag>
akti[]e	akti[]	<aktive, aktiv>
e[]ige	e[]ge	<ewige, ew' ge>
Ro[]e	Rö[]chen	<Rose, Röschen>
fra[]en	fra[]te	<fragen, fragte>

- Gebt an, wie die fehlenden Laute in den obigen Beispielen ausgesprochen werden
- Beschreibt in Worten, in welchem Kontext die Frikative und Plosive in diesen Beispielen stimmlos sind.
- Schreibt eine phonologische Regel, die angibt, wo Frikative und Plosive im Deutschen stimmlos werden
- Welches ist das zugrundeliegende Phonem in <Lobes, Lob>, /b/ oder /p/?
- Nehmt mal an, dass es sich hier nicht um einen Prozess der Auslautverhärtung handelt, sondern um einen Prozess der Inlauterweichung. Beschreibt in Worten diesen Prozess
- Wenn wir einen Prozess der Inlauterweichung annehmen, welche Probleme ergeben sich dann bei Wörtern wie <Rates, Rat>?

~~~~~

Regel für die Auslautverhärtung:

1. Versuch:

Plosive + Frikative    →    [-stimmhaft] / \_\_\_\_ (C)]<sub>σ</sub>

N.B.: (C) steht für einen optionalen Konsonanten. Es kann nämlich zwischen dem Laut, der stimmlos wird und dem Silbenende noch ein anderer Konsonant stehen, und trotzdem wird der besagte Laut stimmlos. z.B. /b/ in dem Wort <Abt>. Zwischen /b/ und dem Silbenende steht der Konsonant /t/. Trotzdem wird /b/ zu [p]:

(10) /abt/ → [apt]

eine bessere Regel, weil kürzer:

[+ **Obstruent**] → [-stimmhaft] / \_\_\_\_ (C)]<sub>σ</sub>

"Obstruenten" sind all jene Laute, bei denen der Luftstrom ernsthaft behindert wird, also Frikative und Plosive. Frikative und Plosive verhalten sich oft ähnlich (z.B. in so einem Prozess wie der Auslautverhärtung), deshalb fasst man sie zur Gruppe der Obstruenten zusammen.

Preisfrage: Warum heißt die Regel nicht

+ Obstruent → [-stimmhaft] / \_\_\_\_ (C)]<sub>σ</sub>

+ stimmhaft

### 3.6. Die Phoneme /s/ und /z/

~~~~~

ÜBUNG:

Schaut euch die folgenden Daten zur Distribution von [s] und [z] genau an und beantwortet dann die Fragen.

- | | | |
|----|------------------------|---------------------|
| a. | [ˈʁaɪsən] ~ [ˈʁaɪzən] | <reißen> ~ <reisen> |
| | [bi:kzam] | <biegsam> |
| | [ʃtœpsəl] | <Stöpsel> |
| b. | [ˈze:] | <See> |
| | [ˈzo:] | <so> |
| | [ˈzɪnt] | <sind> |
| c. | [ˈʃtal] | <Stall> |
| | [ˈʃpi:l] | <Spiel> |
| | [ˈʃtaɪn] | <Stein> |
| | [ˈsmokɪŋ] | <Smoking> |
| | [ˈsla:lɔm] | <Slalom> |
| d. | [ˈɡrɛ:.zə] [ˈɡra:s] | <Gräser>, <Gras> |
| | [ˈʁo:.zə], [ˈʁø:s.çən] | <Rose>, <Röschen> |

1. Handelt es sich bei [s] und [z] um zwei Phoneme oder Allophone eines einzigen Phonems? Schaut euch zur Beantwortung dieser Frage besonders die Daten in a. gut an.

2. Beschreibt die Distribution von [s] und [z] in Worten. In welchen Kontexten kommt immer nur [s] vor, in welchen Kontexten kommt immer nur [z] vor?

3. Beschreibt den Prozess, der wortinitiale Phone [z] generiert, mit Hilfe einer phonologischen Regel nach dem Schema A → B / C_ D

4. Wie ist das im Italienischen? Sucht nach italienischen Wörtern, in denen das <s> in den gleichen Kontexten steht wie wir sie hier im Deutschen haben. Beobachtet, was dort passiert.

~~~~~

Wenn ihr diese Übung in allen ihren Teilen beantwortet habt, dann solltet ihr Folgendes herausgefunden haben.

Die Beispiele in (a) zeigen uns, dass /s/ und /z/ im Deutschen Phoneme sind, denn sie können Minimalpaare wie <reißen> und <reisen> unterscheiden. Auch die Wörter <biegsam> und <Stöpsel> weisen auf einen bedeutungsdifferenzierenden Kontrast hin: in beiden Fällen steht vor dem [z], bzw. [s] ein Konsonant und dahinter ein Vokal, der Kontext ist in beiden Fällen also sehr ähnlich, doch im ersten Beispiel haben wir den stimmhaften Sibilanten, im zweiten Beispiel den stimmlosen Laut. Die Beispiele in (b) zeigen uns, dass die Phoneme /s/ und /z/ am Wortanfang vor Vokal jedoch immer als [z] realisiert werden. Wir können dieses Phänomen mit einer phonologischen Regel folgendermaßen ausdrücken:

(11) /s/ → [z]/ [w \_\_ V

Diese Regel besagt, dass das Phonem /s/ am Anfang eines Wortes und vor einem Vokal als stimmhaftes [z] realisiert wird. Die offene eckige Klammer mit der Etikette "W" steht für den Wortanfang, der Platzhalterstrich " \_\_ " zeigt uns den Kontext des betreffenden Phonems an und "V" steht für "Vokal". Müssen wir hier noch eine zweite phonologische Regel angeben, die besagt, dass das Phonem /z/ in demselben Kontext auch als [z] realisiert wird, also:

(12) /z/ → [z]/ [w \_\_ V

Nein, das ist natürlich nicht notwendig, denn die Realisierung von /z/ als [z] ist der "Normalfall" (im Englischen sagt man, der *default*-Fall). Im Normalfall nehmen wir an, dass ein Phonem ohne phonetischen Veränderungen realisiert wird.

Bei diesem phonologischen Prozess handelt es sich um einen Prozess der **Neutralisierung**: am Wortanfang vor Vokal werden die beiden Phoneme /s/ und /z/ zu einem einzigen Laut [z] neutralisiert, d.h., sie werden beide mit demselben Laut [z] realisiert:

(13) Neutralisierung der Phoneme /s/ und /z/ zu [z] am Wortanfang vor Vokal:

/s/ ↘  
[z]  
/z/ ↗

Das bedeutet, dass wir von einem Wort wie ['ze:] gar nicht sagen können, ob es zugrundeliegend das Phonem /s/ oder das Phonem /z/ hat. Beide Phoneme würden auf jeden Fall als [z] realisiert werden.

Die Beispiele in (c) zeigen uns einen weiteren Neutralisierungsprozess. Wir sehen, dass /s/ und /z/ am Wortanfang vor Konsonant als [ʃ] realisiert werden. Das gilt jedenfalls für native Wörter wie *Stall*, *Spiel*, *Stein*. Bei Fremdwörtern finden wir in diesem Kontext auch oft die Realisierung [s], wie z.B. in *Smoking* und *Slalom*. Wir können also für den nativen Wortschatz folgende Regel aufstellen:

- (14) /s/, /z/ → [ʃ] / [w] \_\_ C

Die Regel besagt, dass die Phoneme /s/ und /z/ am Wortanfang vor einem Konsonanten ("C") als [ʃ] realisiert werden. Da es sich um einen Neutralisierungsprozess handelt, wissen wir auch in diesem Fall bei einem Wort wie z.B. *Spiel* nicht, ob zugrundeliegend ein /s/ oder /z/ vorliegt. Es könnte zugrundeliegend sogar ein [ʃ] vorliegen, auch wenn die Orthographie uns eher an ein /s/ oder /z/ denken lässt. Was aber auch immer das zugrundeliegende Phonem sein mag, alle drei Phoneme können in diesem Kontext im nativen Wortschatz nur als [ʃ] realisiert werden:

- (15) Neutralisierung der Phoneme /s/, /z/, /ʃ/ zu [ʃ] am Wortanfang vor Konsonant:  
 /s/ ↘  
 /ʃ/ → [ʃ]  
 /z/ ↗

Die Beispiele in (d) zeigen uns einen Neutralisierungsprozess, den wir schon kennen, die Auslautverhärtung. Das Phonem /z/ wird am Silbenende als stimmloses [s] realisiert.

Vergleichen wir nun die Realisierung von /s/ und /z/ im Deutschen mit der Distribution derselben Laute im Italienischen. Sind [s] und [z] im Italienischen bedeutungsunterscheidende Einheiten, also Phoneme? Bevor wir diese Frage beantworten, müssen wir vorausschicken, dass sich die Realisierung der beiden Laute je nach regionaler Varietät ziemlich unterscheidet. So kann man davon ausgehen, dass /s/ und /z/ im Toskanischen zwei Phoneme sind, wie uns das folgende Minimalpaar zeigt (Canepari 1979):

- (16) chie[s]e ~ chie[z]e                      *fragte (passato remoto), Kirchen (pl.)*

Allerdings gibt es nicht viele Minimalpaare mit /s/ und /z/. Da das Toskanische oft auch als Maßstab für die italienische Standardaussprache benutzt wird, gehen die meisten Lehrbücher davon aus, dass das Italienische s/ und /z/ als Phoneme kontrastiert.

In den norditalienischen Varietäten wie sie z.B. im Veneto gesprochen werden ist es allerdings so, dass im Wortinneren, zwischen Vokalen, immer ein stimmhaftes [z] erscheint. Betrachtet die folgenden "norditalienischen" Aussprachen:

- (17) ca[z]a  
 ro[z]a  
 co[z]a  
 me[z]e

Wenn allerdings vor oder hinter dem Sibilanten im Wortinneren ein Konsonant steht, dann ist dieser in den Varietäten des Veneto stimmlos:<sup>6</sup>

- (18) sen[s]o  
 sal[s]a

<sup>6</sup> In den Varietäten des Trentino ist allerdings ein alveolarer Sibilant, der in der Wortmitte auf einen Konsonanten folgt oft stimmhaft. So hört man im Trentino oft die Aussprache *sen[z]o*.

re[s]to

Wie sieht es nun am Wortanfang aus? Sehen wir uns Beispiele an, in denen der auf den Sibilanten ein Vokal folgt:

- (19) [s]enza  
[s]olo  
[s]ale

In allen diesen Fällen ist der betreffende Laut stimmlos.  
Und wenn auf den Sibilanten ein Konsonant folgt?

- (20) [s]torto  
[s]pendere  
[s]carto  
[z]drammatizzarte  
[z]berla  
[z]gangherato

Wir sehen hier einen Assimilationsprozess: der Sibilant ist stimmlos, wenn ein stimmloser Konsonant folgt und stimmhaft, wenn ein stimmhafter Konsonant folgt.

Es gibt auch italienische Varietäten, in denen wir in diesem Kontext einen ähnlichen Prozess wie im Deutschen finden. So wird der Sibilant z.B. in den Dialekten der Campania vor einem Konsonanten zu [ʃ] und wir finden Aussprachen wie [ʃ]porco, [ʃ]catola.

Und wie sieht es am Wortende aus? Nun, das ist nicht so leicht zu sagen, da es im Italienischen kaum Wörter mit wortfinalen Konsonanten gibt. Aber in den wenigen Wörtern, wo der Sibilant am Ende des Wortes steht, wird er eindeutig stimmlos ausgesprochen:

- (21) lapi[s]  
ga[s]  
rebu[s]

### 3.7. Phonologische/Phonetische Merkmale

Die Merkmale sind die "Atome", aus denen Phone und Phoneme gemacht werden. z.B. kann das deutsche [p] folgendermaßen charakterisiert werden

[p] = [+ konsonantisch]  
[+ obstruent]  
[- stimmhaft]  
[+ labial]  
.....

- Merkmale können dazu benutzt werden, um Laute minimal zu unterscheiden:

[p] und [b] unterscheiden sich durch das Merkmal  
[± stimmhaft]

- Merkmale helfen dabei, **natürliche Klassen** zu charakterisieren. Natürliche Klassen sind Mengen von Lauten, die sich in phonologischen Prozessen gleich verhalten. So kann man zum Beispiel das Resultat der Auslautverhärtung ganz

einfach beschreiben, indem man sagt: alle Obstruenten bekommen das Merkmal [-stimmhaft].

Obstruent  $\rightarrow$  [-stimmhaft] / \_\_\_\_ (C)]<sub>σ</sub>

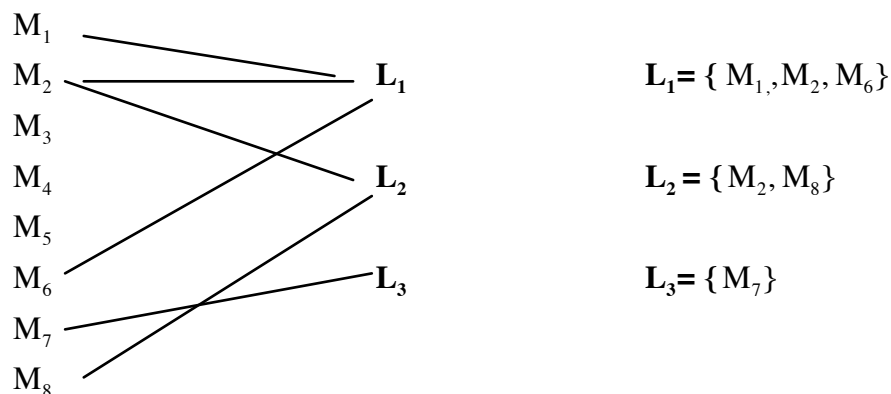
Wenn wir keine Merkmale hätten, müssten wir eine ganze Reihe von Regeln schreiben:

/b/  $\rightarrow$  [p] / \_\_\_\_ (C)]<sub>σ</sub>  
 /d/  $\rightarrow$  [t] / \_\_\_\_ (C)]<sub>σ</sub>  
 /g/  $\rightarrow$  [k] / \_\_\_\_ (C)]<sub>σ</sub>  
 /v/  $\rightarrow$  [f] / \_\_\_\_ (C)]<sub>σ</sub>  
 /z/  $\rightarrow$  [s] / \_\_\_\_ (C)]<sub>σ</sub>

- man kann annehmen, dass es ein universelles Inventar an Merkmalen gibt, aus dem sich jede Sprache bedient und verschiedene Merkmale zu bestimmten Lauten kombiniert:

universelles  
Merkmalsinventar

Laute einer bestimmten  
Sprache



- wenn sich eine Sprache ihre Merkmale ausgesucht hat, dann trifft sie aber noch eine Unterscheidung in "wichtigere" und "weniger wichtige" Merkmale. In jeder Sprache gibt es distinktive und nicht-distinktive Merkmale, d.h., Merkmale, die Wörter in ihrer Bedeutung voneinander unterscheiden können und andere, die das nicht können. Phoneme werden aus distinktiven Merkmalen gebaut.

z.B. [± aspiriert]  $\rightarrow$  ist distinktiv im Hindi  
 ist nicht distinktiv im Deutschen

Wir werden uns nun jene distinktiven Merkmale anschauen, die im Deutschen eine Rolle spielen. Die folgende Liste ist eine Untermenge der Liste von universalen Merkmalen in Hall (2000):

### Allgemeine Merkmale:

**[± konsonantisch]:** dieses Merkmal unterteilt die Laute in zwei große natürliche Klassen: Vokale und Konsonanten. Alle Laute, die [+konsonantisch] sind (also alle Konsonanten), sind durch eine Verengung des Artikulationsraumes zwischen Larynx und Lippen gekennzeichnet. So wird z.B. ein [p] produziert, indem der

Artikulationsraum bei den Lippen geschlossen wird. Laute, die [-konsonantisch] sind (also Vokale), werden ohne Verengung produziert. Wenn wir z.B. ein [a] produzieren, dann strömt die Luft ohne Verengung durch den Artikulationsraum.

**[± sonorantisch]:** Dieses Merkmal unterscheidet Obstruenten von Sonoranten. Obstruenten (also Frikative und Plosive) sind [-sonorantisch]. Sonoranten sind [+sonorantisch]. Zur Gruppe der Sonoranten gehören Nasale, Liquide (also r- und l-Laute) und Approximanten (im Deutschen [j]).

Der große Unterschied zwischen Sonoranten und Nicht-Sonoranten besteht darin, dass Sonoranten spontan stimmhaft sind, Nicht-Sonoranten jedoch nicht. Das bedeutet, dass Sonoranten normalerweise immer stimmhaft sind, während es bei den Obstruenten die typischen Paare von stimmhaften und stimmlosen Konsonanten gibt. Ein Sonorant wie [m], z.B., ist normalerweise stimmhaft. Ein Nicht-Sonorant wie [b] hingegen hat normalerweise einen stimmlosen "Bruder", nämlich [p].<sup>7</sup>

**Merkmale, die (vor allem) zur Klassifizierung von Konsonanten verwendet werden:**

**[± stimmhaft]:** dieses Merkmal trennt stimmhafte von stimmlosen Lauten. Es unterscheidet also Laute wie /b, d, m, a/ von Lauten wie /p, t, s/.

Merkmale, die die Artikulationsart (der Konsonanten) betreffen:

**[± kontinuierlich]:** dieses Merkmal ist wichtig, weil es Plosive von Frikativen unterscheidet. Bei Lauten mit dem Merkmal [-kontinuierlich] kommt es bei der Produktion zu einem kompletten Verschluss, es handelt sich dabei also um die Plosive. Frikative hingegen sind [+kontinuierlich], weil bei ihrer Produktion der Luftstrom kontinuierlich durch den Artikulationsraum fließt. Nasale werden normalerweise auch als [-kontinuierlich] klassifiziert, da sie auch mit einem kompletten Verschluss produziert werden.

**[± nasal]:** dieses Merkmal unterscheidet nasale Laute von nicht-nasalen (oralen) Lauten.

**[± lateral]:** dieses Merkmal unterscheidet die Laterale von allen anderen Lauten.

Merkmale, die den Artikulationsort (der Konsonanten) betreffen:

**[LABIAL]:** dieses Merkmal charakterisiert die labialen Laute wie z.B. /p, b, f, v, m/. Dieses Merkmal ist ein sogenanntes **privatives Merkmal**, das heißt, ein Laut hat entweder das Merkmal [LABIAL] oder er hat es nicht, aber es gibt nicht [+labiale] Laute und [-labiale] Laute. Man nimmt an, dass [LABIAL] ein privatives Merkmal ist, da sich die Laute, die dieses Merkmal besitzen, wie eine natürliche Klasse verhalten, die Merkmale, die es nicht besitzen bilden jedoch keine natürliche Klasse. Das bedeutet, dass es keine natürliche Klasse mit dem Merkmal [-labial] gibt. Da es keine natürliche Klasse dieser Art gibt, kann es auch das Merkmal [-labial] nicht geben.

---

<sup>7</sup> Es gibt Sprachen, in denen stimmlose Sonoranten vorkommen, aber es handelt sich dabei meist um allophonische Prozesse, die den Sonoranten in einem bestimmten Kontext stimmlos machen. Es gibt auch Sprachen, die z.B. nur ein [p], aber kein [b] haben. Wenn aber einer von den beiden Lauten im Lautinventar einer Sprache fehlt, dann ist das immer das [b]. Diese Tatsache zeigt uns, dass der "normale", unmarkierte Plosiv das stimmlose [p] sein muss, nicht das stimmhafte [b]. Denn in allen Sprachen finden wir [p], aber nicht in allen finden wir auch [b]. Die unmarkierte Charakteristik von Sonoranten ist also die Stimmhaftigkeit, die unmarkierte Charakteristik von Obstruenten hingegen ist die Stimmlosigkeit.

Mit anderen Worten, labiale Laute wie /p, b, f/ u.s.w. sind manchmal als Gruppe das Ziel eines phonologischen Prozesses (bilden also eine natürliche Klasse). Es gibt aber keinen phonologischen Prozess, der Nicht-Labiale als Gruppe betrifft.

**[KORONAL]:** dieses Merkmal charakterisiert im Deutschen dentale/alveolare Laute [d, t, s, z, n, l] und postalveolare Laute [ʃ].

[KORONAL] ist auch ein privatives Merkmal. Entweder besitzt ein Laut dieses Merkmal, oder er besitzt es nicht.

**[± anterior]:** dieses Merkmal kann im Deutschen dazu verwendet werden, um die alveolaren Sibilanten [s, z] von dem postalveolaren Sibilanten [ʃ] zu unterscheiden. Dieses Merkmal ist notwendig, da alle diese Sibilanten das Merkmal [KORONAL] besitzen.

**[DORSAL]:** zu den dorsalen Lauten gehören die palatalen, velaren und uvularen Konsonanten des Deutschen, also [ç, x, k, g, ŋ, R, ʁ]. Auch [DORSAL] ist ein privatives Merkmal.

### **Merkmale, die (vor allem) zur Klassifizierung von Vokalen verwendet werden:**

**[± rund]:** wir verwenden dieses Merkmal, um gerundete Vokale wie /y, œ, u, o/ zu beschreiben.

**[± hoch]:** wir verwenden dieses Merkmal, um hohe Vokale wie /i, y, u/ von allen anderen Vokalen zu unterscheiden

**[± tief]:** wir verwenden dieses Merkmal, um tiefe Vokale wie /a/ zu kennzeichnen. N.B.: mittlere Vokale wie /e, o, œ/ sind durch die Merkmale [- hoch] und [-tief] gekennzeichnet

**[± gespannt]:** dieses Merkmal unterscheidet gespannte Vokale wie /i:, u:, o:/ von ungespannten Vokalen wie /ɪ, ʊ, ɔ/.

**[± hinten]:** dieses Merkmal unterscheidet hintere Vokale wie /u:, o:/ von vorderen Vokalen wie /i:, e:/. Die zentralen Vokale werden meistens zu den [+hinten] Vokalen gezählt.

### **Übung zu den distinktiven Merkmalen:**

Gebt mindestens ein distinktives Merkmal an, das die folgenden Minimalpaare unterscheidet:

|                 |  |               |  |
|-----------------|--|---------------|--|
| reißen ~ reisen |  | Meere ~ Möhre |  |
| Bann ~ dann     |  | sie ~ See     |  |
| nackt ~ Nacht   |  | Paar ~ Bar    |  |
| Sünde ~ Sunde   |  | seine ~ Leine |  |
| Miete ~ Mitte   |  | Pein ~ fein   |  |
| Pein ~ mein     |  | Höhle ~ hohle |  |

**Übung 3**

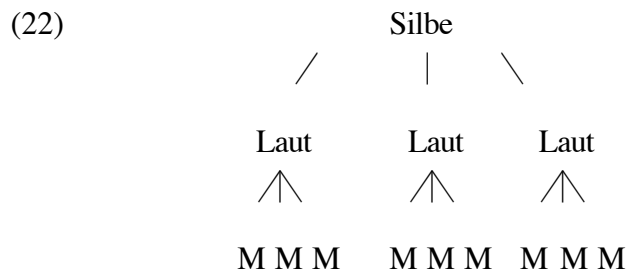
Wir haben die Regel der Auslautverhärtung folgendermaßen formuliert:

Obstruent → [-stimmhaft] / \_\_\_\_ (C)]<sub>σ</sub>

Überlegt euch nun, wie man diese Formulierung verbessern könnte. Wie kann man die natürliche Klasse der Segmente, die von dieser Regel betroffen sind, mit Hilfe von Merkmalen ausdrücken?

**4. Prosodische Phonologie und ihre Regeln****4.1 Die Silbe**

In der Phonologie werden Merkmale zu Lauten gebündelt, aber Laute und Merkmale sind nicht die einzigen phonologischen Einheiten. Laute werden wiederum zu einer größeren Einheit verbunden, der Silbe:



Die Sprecher einer Sprache haben eine gewisse Intuition, was Silben betrifft. Sie wissen meistens ziemlich genau, wo eine Silbe anfängt und wo sie aufhört und welche Silben in ihrer Sprache wohlgeformt sind und welche hingegen nicht möglich sind. So weiß ein Sprecher des Deutschen, dass *sparf* eine mögliche Silbe des Deutschen sein könnte, *sprf* jedoch nicht.

Diese Tatsachen sollte uns eigentlich ziemlich verwundern, denn es ist praktisch unmöglich, die Silbe einer Sprache zu messen. In einem Spektrogramm können wir also keine Silbengrenzen erkennen. Existiert ein sprachliches Element, das man nicht messen kann, überhaupt? Diese Unsichtbarkeit der Silbe hat einige Linguisten dazu geführt, zu behaupten, dass es die Kategorie der Silbe eigentlich nicht gibt. Wir werden jedoch gleich sehen, dass die Silbe eine große Rolle in unserem Sprachbewusstsein spielt. Wenn die Silbe also auch physikalisch gesehen unsichtbar ist, so scheint sie doch eine mentale Kategorie unserer Sprachkompetenz zu sein.

Dass Sprecher die Kategorie der Silbe besitzen, können wir z.B. beobachten, wenn Kinder einen Abzählreim aufsagen:

## (23) Ich-und-du-Mül-lers-Kuh-Mül-lers-E-sel-das-bist-du

Bei jeder Silbe zeigt das Kind auf einen anderen Spieler. Es verwendet also, auch wenn es sich dessen nicht bewusst ist, die Kategorie der Silbe zum Zählen der Mitspieler.

**Übung 4**

Im Italienischen scheinen Abzählreime etwas anders zu funktionieren. Zerlegt die folgenden Abzählreime in ihre rhythmischen Einheiten und versucht zu beschreiben, wie diese Einheiten aussehen. Ihr könnt auch einen Abzählreim aus eurer Kindheit verwenden:

Uno, due, tre  
chi non scappa c'è

Ambarabà, cici, cocò,  
tre galline sul comò,  
che facevano l' amore,  
con la figlia del dottore,  
il dottore s' ammalò,  
ambarabà, cici, cocò

Obwohl die Sprecher einer Sprache meistens ziemlich genau wissen, wo die Silbengrenzen in einem Wort liegen, so gibt es doch einige Zweifelsfälle. Im Deutschen ist es zum Beispiel nicht völlig klar, wie das folgende Wort in Silben zerlegt wird:

- (24) <Müller>  
a. [mʏl.lə]      oder  
b. [mʏ.lə]      ?

Wir wissen, dass das Wort "Müller" im Deutschen mit nur einem [l] ausgesprochen wird. Es sollte eigentlich wie in b. getrennt werden. Die Sprecher sind sich allerdings nicht ganz sicher. Einige ziehen es vor, das Wort wie in a. in zwei Silben zu zerlegen, wobei das [l] sowohl die erste Silbe schließt als auch die zweite beginnt. Das hängt vielleicht auch damit zusammen, dass die orthographische Norm die Silbentrennung <Mül.ler> vorschreibt.

Die Phonologen sprechen in diesen Fällen von **ambisilbischen Konsonanten**. Sie gehen davon aus, dass das Wort nur ein Segment [l] enthält, das aber in einem gewissen Sinne sowohl zur ersten als auch zur zweiten Silbe gehört. Man könnte sagen, dass die Silbengrenze in die Mitte des Segments fällt. Ambisilbische Konsonanten tauchen immer dann auf, wenn der Vokal einer Silbe betont und kurz ist und kein anderer Konsonant die Silbe schließt. Man transkribiert ambisilbische Konsonanten, indem man den Silbentrennungspunkt unter den betreffenden Konsonanten setzt.

Auch in Fällen wie den folgenden ist die Intuition der Sprecher nicht völlig eindeutig:

(25) Fen.ster oder  
Fens.ter ?

Gei.stes oder  
Geis.tes ?

Die orthographische Norm schreibt die Trennung *Fen.ster* bzw. *Gei.stes* vor, aber auch in diesen Fällen ist die Intuition der Sprecher nicht ganz klar.

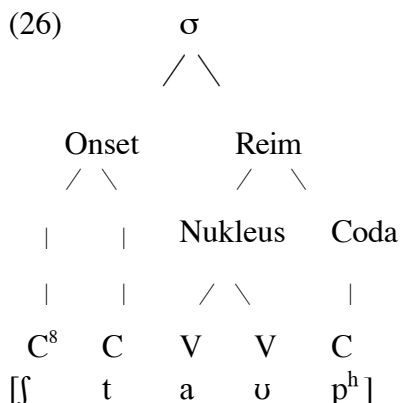
Weitere Evidenz für die Existenz der Silbe können wir darin finden, dass bestimmte phonologische Prozesse als Domäne ihrer Anwendung die Silbe haben. So haben wir z.B. gesehen, dass die r-Vokalisierung und die Auslautverhärtung am Ende einer Silbe stattfinden. Wenn es die Kategorie der Silbe nicht gäbe, wie sollte dann der Anwendungskontext dieser Prozesse beschrieben werden?

Zusammenfassend können wir sagen, dass es, auch wenn wir Silben physikalisch nicht messen können, doch einige Evidenz für die Existenz der Silbe gibt. Die Silbe scheint Teil unserer sprachlichen Kompetenz zu sein, denn wir haben eine gewisse Intuition, was diese Kategorie betrifft: wir wissen normalerweise, wo eine Silbe anfängt und wo sie aufhört, wie eine wohlgeformte Silbe in unserer Muttersprache aussieht, wir verwenden sie spontan in Abzählreimen. Außerdem können wir beobachten, dass die Silbe den Kontext für viele phonologische Prozesse ausmacht.

Die Silbe selbst besteht aus gewissen Untereinheiten. Diese innere Struktur der Silbe werden wir uns im nächsten Kapitel genauer anschauen.

#### 4.1.1 Die Struktur der Silbe

Eine Silbe ist nicht einfach eine Sequenz von Segmenten. Sie hat eine innere Struktur. Es gibt in der Phonologie mehrere Vorschläge dazu, wie die Struktur der Silbe genau aussehen soll. Wir werden das folgende, einfache Modell verwenden:



Jede Silbe einer natürlichen Sprache besteht aus mindestens einem **Nukleus** (oder **Silbengipfel**), bei dem es sich meistens um einen Vokal handelt. In allen Sprachen der Welt gibt es außerdem Silben, die einen **Onset** (oder **Silbenanlaut**) besitzen, d.h. Silben, in denen dem Nukleus einer oder mehrere Konsonanten vorangehen. In manchen Sprachen gibt es außerdem auch Silben, die eine **Coda** (oder

<sup>8</sup> "C" steht für "Konsonant" und "V" für "Vokal"

**Silbenauslaut**) besitzen, d.h., in denen auf den Nukleus noch einer oder mehrere Konsonanten folgen.

### Übung 5

Transkribiert folgende Wörter und gebt ihnen eine Silbenstruktur:

*Kraft, Nase, Onkel, Obst*

Wie kommen aber eigentlich die Linguisten dazu, anzunehmen, dass die Struktur der Silbe genau so und nicht anders aussieht? Warum nimmt man an, dass es eine Kategorie wie den Onset gibt, oder dass Nukleus und Coda zusammen eine Einheit bilden, die wir hier mit *Reim* bezeichnet haben? Im Folgenden werden wir kurz einige Argumente diskutieren, die zu der oben abgebildeten Struktur geführt haben.

Warum nehmen wir an, dass es eine Kategorie *Onset* gibt? Ein gutes Argument für die Existenz der Kategorie *Onset* ist die Tatsache, dass es bestimmte Regeln gibt, die sich genau auf diese Kategorie beziehen. Wenn zum Beispiel im Italienischen ein Onset aus drei Konsonanten besteht, dann muss der erste dieser Konsonanten immer ein [s] sein. Wir können also eine Silbe mit einem Onset wie *spr* in *spre.co* haben, aber es gibt keine Silbe die etwa mit *fpr* beginnt.

Es gibt auch etwas lustigere Evidenz für die Existenz des Onset. In **Versprechern** ist es oft so, dass die Sprecher verschiedene Laute miteinander vertauschen, wie im folgenden Beispiel:<sup>9</sup>

(27) mit dem Zinger feigen <= mit dem Finger zeigen

In diesem Beispiel wurden die beiden Onsets der Wörter *Finger* und *zeigen* vertauscht. Es scheint so zu sein, dass bei Versprechern dieser Art meist Untereinheiten der Silbe miteinander vertauscht werden, also Onsets mit Onsets, Nuklei mit Nuklei, Cudas mit Cudas, Reime mit Reimen, aber nicht zum Beispiel Onset und Nukleus zusammen, denn diese beiden Elemente bilden keine Einheit!

Warum nehmen wir eigentlich an, dass es eine Kategorie *Reim* gibt, die aus Nukleus und Coda besteht? Ein gutes Argument für die Existenz der Kategorie *Reim* ist die Tatsache, dass diese Kategorie in der Lyrik verwendet wird: der Reim eines Gedichtes besteht meistens aus den Einheiten Nukleus und Coda, wie wir im folgenden Gedicht von Goethe sehen können:

(28) Wer reitet so spät durch Nacht und **Wind**  
 Es ist der Vater mit seinem **Kind**  
 Er hat den Knaben wohl in dem **Arm**,  
 Er faßt ihn sicher, er hält ihn **warm**.

<sup>9</sup> Aus der Versprechersammlung von Helen Leuninger (s. <http://web.uni-frankfurt.de/fb10/LSLeuninger/pl/pl-neb.htm>)

Mein Sohn, was birgst du so bang dein Gesicht?  
 Siehst Vater, du den Erlkönig nicht!  
 Den Erlenkönig mit Kron' und Schweif?  
 Mein Sohn, es ist ein Nebelstreif.

(Die ersten beiden Strophen des Gedichtes *Erlkönig*, Johann Wolfgang von Goethe, 1749-1832)

Wir sehen hier, dass in zwei aufeinander folgenden Zeilen immer der Nukleus und die Coda der letzten Silbe (*ind, arm, icht, eif*), also der Reim, gleich bleiben.

Für die Existenz der Kategorie *Nukleus* spricht schon alleine die Tatsache, dass diese Einheit die einzige Untereinheit der Silbe ist, die immer realisiert werden muss. Im Deutschen kann der Nukleus durch einen (kurzen oder langen) Vokal, einen Diphthong, oder einen silbischen Konsonanten realisiert werden:

- (29) a. Kurzer Vokal: W[a].che  
 b. Langer Vokal: L[i:].be  
 c. Diphthong: L[au].be  
 d. Silbischer Konsonant: war.[t̩] <warten>  
                                   ne.[b̩] <neben>  
                                   Fle.[g̩] <Flegel>

Silbische Konsonanten wie in den Beispielen d. werden im Deutschen vor allem bei schneller Aussprache verwendet (die schnelle Aussprache stellt den Normalfall dar, nicht die Ausnahme!). Im Beispiel *ne.[b̩]* wurde außerdem der Nasal an den vorangehenden Plosiv assimiliert. Die schwas werden in diesen Wörtern nur bei sehr sorgfältiger, langsamer Aussprache gesprochen. Dann klingen diese Beispiele folgendermaßen:

- (30) war.[tən]  
       ne.[bən]  
       Fle.[gəl]

Auch für die Silbencoda gibt es Restriktionen, die uns zeigen, dass es eine Kategorie dieser Art innerhalb der Silbe geben muss. So können z.B. im Italienischen nur bestimmte Konsonanten in der Coda stehen. Ein Nasal in der Coda ist möglich, ein Plosiv hingegen nicht:

- (31) a. ba[n].da  
       bi[m].bo  
       b. \*ba[p].da  
           \*bi[t].bo

Außer den Restriktionen in Bezug auf die Silbencoda gibt es außerdem noch Evidenz aus dem Bereich der **Geheimssprachen** und **Sprachspiele** für die Existenz dieser Kategorie. Schauen wir uns das folgende Gedicht von Joachim Ringelnatz an:

## (32) Gedicht in Bi-Sprache

Ibich habibebi dibich,  
Lobittebi, sobi liebib.

Habist aubich dubi mibich  
Liebib? Neibin, vebirgibib.

Nabih obidebir febirn,  
Gobitt seibi dibir gubit.  
Meibin Hebirz habit gebirn  
Abin dibir gebirubiht.

*Joachim Ringelnatz (1883-1934)*

Dieses Liebesgedicht ist in der sogenannten Bi- (oder Be)-Sprache geschrieben, einer Geheimsprache, wie sie Kinder oft zum Spiel verwenden. Es gibt im Deutschen viele verschiedene Geheimsprachen dieser Art, die man unter den Name Hühnersprache, Löffelsprache, Gabelsprache, Erbsensprache oder Räubersprache finden kann. Wie funktioniert diese Sprache? Hier ist ein Tipp: die ersten Zeilen lauten in der "Übersetzung" folgendermaßen:

(33) Ich habe dich,  
Lotte, so lieb

Wie werden die Wörter in dieser Sprache gebildet? Wenn man die erste Zeile "zerlegt", dann bekommt man:

(34) I-bi-ch ha-bi-be-bi di-bi-ch  
Lo-bi-tte-bi, so-bi lie-bi-b

Schauen wir uns das Wort *di-bi-ch=dich* genauer an. Es wurde hier die Sequenz *-bi-* hinter den Nukleus und vor die Coda gesetzt. Wenn wir uns die anderen Wörter anschauen, so sehen wir, dass dort dasselbe passiert: immer kommt die Sequenz *-bi-* genau hinter dem Nukleus und vor der Coda. Am Beispiel *He-bi-rz=Herz* sieht man auch, dass das *-bi-* nicht etwa vor den letzten Konsonanten des Wortes gesetzt wird, sondern vor alle finale Konsonanten, also vor die Coda.

Nun, wie lautet der Rest des Gedichtes?

Dieses Beispiel aus dem Bereich der Geheimsprachen zeigt uns, dass die finalen Konsonanten einer Silbe eine Einheit bilden, die Coda. Denn diese Einheit bestimmt, an welche Stelle die Sequenz *-bi-* gesetzt wird.

### Übung 6

In den folgenden vier Zeilen habe ich den Anfang von Ringelnatz' Gedicht in die Hühnersprache übersetzt. Die Hühnersprache ist etwas komplizierter als die Bi-Sprache. Versucht herauskriegen, wie sie funktioniert! Für welche Unterheiten der Silbe liefert uns die Hühnersprache Evidenz?

Ihichlefich hahalefa behelefe dihichlefich  
Loholefo ttehelefe soholefo liehieblefieb

Hahastlefast auhauchlefauch duhulefu mihichlefiŋ  
 liehieblefiē? Neiheinlefein, veherlefer gihiblefiē

### Übung 7

Der folgende Vers ist ein sogenannter Schüttelreim:

1. "Die Boxer aus der Meisterklasse
2. boxten sich zu Kleisternasse
3. Und aus dem ganzen Massenkleister
4. erhebt sich stolz der Klassenmeister"

- a) Transkribiert das unterstrichene Wort in Zeile 1
- b) Zerlegt dieses Wort in Silben und gibt die Silbenstruktur der einzelnen Silben an (N.B. <ss>, ausgesprochen als [s] ist ein ambisilbisches Segment, d.h. es gehört sowohl zur Coda der vorangehenden Silbe, als auch zum Onset der folgenden)
- c) Erklärt das Muster, nach dem die ersten beiden Zeilen (1 und 2) des Schüttelreims gebildet werden. Verwendet dazu die Terminologie der Silbenstruktur.
- d) Erklärt, wie das Rezept für den Schüttelreim nach Zeile 2 weitergeht, d.h., wie Zeile 3 aus Zeile 1 und 2 abgeleitet wird. Ist eure Erklärung morphologischer oder phonologischer Art?

### 4.1.2 Die Sonoritätshierarchie

Wir haben in den letzten beiden Kapiteln gesehen, dass es verschiedene Argumente für die Existenz der Silbe und ihrer Untereinheiten Onset, Nukleus, Coda und Reim gibt. Ein weiterer Grund, an die Existenz der Silbe zu glauben, ist, dass es ein Prinzip gibt, das in allen Sprachen der Welt eine große Rolle bei der inneren Organisation der Silbe spielt. In den Sprachen der Welt sind Silben meistens so organisiert, dass die sonorsten Segmente im Zentrum (also im Nukleus) stehen, während die **Sonorität** zum Silbenrand hin abnimmt. Was ist hier mit Sonorität gemeint?<sup>10</sup>

Viele Linguisten nehmen an, dass sich Segmente auf einer Sonoritätsskala anordnen lassen. Wie dieses Skala genau aussieht ist noch nicht klar. Wir werden hier ein Modell nehmen, das von einigen Phonologen für das Deutsche vorgeschlagen wird (Wiese 1988, Vater 1996, Ramers&Vater 1995, Ramers 2001):

(35) **Sonoritätsskala:**

Plosive < Frikative < Nasale < l < r < Vokale

("<" bedeutet: "weniger sonor als)

<sup>10</sup> Der Begriff *Sonorität* darf nicht mit dem Begriff *Stimmhaftigkeit* (auf italienisch: *sonorità*) verwechselt werden! Er darf auch nicht mit dem Merkmal [ $\pm$  *sonorantisch*] verwechselt werden!

Nach dieser Skala sind Plosive weniger sonor als Frikative, Frikative sind weniger sonor als Nasale, diese sind wiederum weniger sonor als [l] und [l] ist weniger sonor als [r]. Vokale sind die sonorsten Elemente.

Nun kann man beobachten, dass in den Sprachen generell Silben so konstruiert sind, dass die Sonorität bis zum Nukleus hin ständig ansteigt und nach dem Nukleus wiederum bis zum Ende der Silbe hin sinkt. Vor dem Nukleus steht also nie ein sonoreres Element vor einem weniger sonoren Element, nach dem Nukleus ist es genau umgekehrt. Wir formulieren diese Beobachtung im Sonoritätsprinzip:

(36) **Sonoritätsprinzip:**

Die Sonorität nimmt bis zum Nukleus stetig zu und fällt nach dem Nukleus stetig ab

Wir können das am Beispiel von deutschen Silben sehen:

- (37) a. Kraft       $k < r < a > f > t$   
       b. Kerl       $k < \epsilon > r > l$   
       c. Knie       $K < n < i$ :

Im ersten Beispiel sehen wir, dass die Elemente im Onset nach dem Sonoritätsprinzip organisiert sind: der Plosiv /k/ ist weniger sonor als /r/, deshalb muss ein /k/ im Onset immer vor einem /r/ stehen, eine Silbe, die mit /rk/ beginnt, ist im Deutschen nicht möglich. In der Coda ist es genau umgekehrt, die Sonorität fällt nach dem Nukleus ständig ab. Deshalb ist die Sequenz /ft/ erlaubt, bei der der sonorere Frikativ vor dem weniger sonoren Plosiv steht. Im zweiten Beispiel, *Kerl*, sehen wir, dass /r/ sonor als /l/ sein muss, denn /r/ steht in der Coda näher am Nukleus als /l/. Das letzte Beispiel schließlich zeigt uns, dass Plosive weniger sonor als Nasale sind, denn das /k/ steht im Onset vor dem /n/.

Die Sonoritätshierarchie ist nicht unbedingt für alle Sprachen gleich und sie wird auch nicht immer in allen Sprachen befolgt. So gibt es z.B. im Russischen Wörter, bei denen das sonorere /r/ im Onset vor dem weniger sonoren /t/ steht:

- (38) rta      'Mund, Genitiv Sg.'

Auch im Deutschen gibt es einige Lautsequenzen, die nicht die Sonoritätshierarchie befolgen. So kann z.B. im Onset vor einem Plosiv noch ein [ʃ] oder ein [s] stehen:

- (39) Strumpf      [ʃ]trumpf  
       Skala      [s]kala

Sibilanten sind allerdings die einzigen Segmente, die in diesem Kontext stehen können. Wir werden in den nächsten Abschnitten noch sehen, welche Segmente sich im Deutschen in der Coda über die Sonoritätshierarchie hinwegsetzen können.

Trotzdem es Segmentsequenzen gibt, die sich nicht an die Sonoritätshierarchie halten, so kann man doch sagen, dass die Sonoritätshierarchie von den meisten Sprachen wenigstens größtenteils befolgt wird. Es handelt sich also um ein universelles Prinzip.

### 4.1.3 Universelle Tendenzen: unmarkierte Onsets und Codas

Außer der Befolgung der Sonoritätshierarchie gibt es auch noch andere universelle Tendenzen, die die Struktur der Silbe bestimmen. Schauen wir uns einige Beispiele an, die Hall (2000: 212) entnommen sind:

- |      |              |              |              |
|------|--------------|--------------|--------------|
| (40) | [vzdwuʃ]     | 'entlang'    | Polnisch     |
|      | [fspla.nout] | 'aufflammen' | Tschechisch  |
|      | [vptskvni]   | 'ich schäle' | Georgisch    |
| (41) | [a.lo.ha]    | 'Liebe'      | Hawaiianisch |
|      | [wa.hi.ne]   | 'Frau'       | Samoanisch   |
|      | [du:.xo]     | 'Hochzeit'   | Iraqw        |

In den ersten drei Beispielen sehen wir Wörter aus den Sprachen Polnisch, Tschechisch und Georgisch, in denen sehr viele Konsonanten im Onset der Silbe erlaubt sind. Nehmen wir z.B. das tschechische Beispiel *fspla.nout*: der Onset der ersten Silbe enthält vier Konsonanten! Die letzten drei Beispiele hingegen stammen aus Sprachen, in denen höchstens ein Konsonant im Onset stehen darf. Wenn wir nun sehr viele Sprachen untersuchen, dann können wir beobachten dass es einerseits Sprachen gibt, in denen nur ein Konsonant im Onset erlaubt ist, und andererseits Sprachen, die mehrere Konsonanten im Onset erlauben. Es gibt aber keine Sprachen, in denen obligatorisch mehrere Konsonanten im Onset stehen müssen! Mit anderen Worten können wir sagen, dass es anscheinend ein Prinzip gibt, das komplexe Onsets verbietet. Manche Sprachen, wie z.B. Hawaiianisch, Samoanisch und Iraqw befolgen dieses Prinzip, andere Sprachen, wie Polnisch, Tschechisch, Georgisch, befolgen es nicht. In der Linguistik spricht man in diesen Fällen von **markierten** und **unmarkierten** Strukturen. Ein einfacher Onset ist die unmarkierte Struktur, denn wir finden sie in allen Sprachen der Welt, auch in denen, die komplexe Onsets erlauben. Ein komplexer Onset ist die markierte Struktur, denn wir finden sie nicht in allen Sprachen, sondern nur in Sprachen wie Polnisch, Tschechisch, Georgisch.

In Bezug auf den Onset können wir eine weitere Beobachtung machen. Alle Sprachen der Welt erlauben Silben mit einem Onset. In einigen Sprachen gibt es auch Silben ohne Onset (z.B. im Hawaiianischen, wie uns die erste Silbe in *a.lo.ha* zeigt). In anderen Sprachen hingegen sind Silben ohne Onset verboten. Es gibt aber keine Sprache der Welt, in denen Silben *mit* einem Onset nicht möglich sind. Wir können also hier wieder sagen: die unmarkierte Silbe ist eine Silbe, die einen Onset besitzt; die markierte Silbe ist die Silbe ohne Onset.

Es ist nicht ganz einfach, den Begriff *Markiertheit* zu erklären. Unter unmarkierten Strukturen versteht man in der Linguistik Strukturen, die im Normalfall vorkommen. Markierte Strukturen kommen hingegen nur in besonderen Kontexten vor. Unter unmarkierter Struktur verstehen wir hier eine Struktur, die in allen Sprachen der Welt möglich ist und in manchen Sprachen der Welt sogar ausdrücklich verlangt wird. Eine unmarkierte Silbe ist also somit eine wohlgeformte Silbe in allen Sprachen der Welt und manche Sprachen erlauben nur Silben dieser Art.

Wir fassen die beiden Beobachtungen zur unmarkierten Silbenstruktur in der Onset-Bedingung zusammen (s. auch das Silbenanlautgesetz in Hall 2000: 213):

- (42) Onset-Bedingung, Teil I: Unmarkierte Silben haben einen Onset  
 Onset-Bedingung, Teil II: Unmarkierte Silben haben nur einen Konsonanten im Onset

Wie verhält es sich nun mit der Coda? Welche universelle Tendenzen können wir hier beobachten? Schauen wir uns wieder ein paar Beispiele an (s. Hall 2000: 214 und die dort angegebenen Quellen):

- |      |                                |                  |                    |
|------|--------------------------------|------------------|--------------------|
| (43) | [o:psts]                       | 'Obsts, Genitiv' | Deutsch            |
|      | [t̪sɛhɛjələwtx <sup>w</sup> s] | 'ihre Kirche'    | Upriver Halkomelem |
|      | [χosgd̪ʒ]                      | 'ich befahl'     | Svan               |
| (44) | [a.lo.ha]                      | 'Liebe'          | Hawaiianisch       |
|      | [wa.hi.ne]                     | 'Frau'           | Samoanisch         |
|      | [a.pá]                         | 'Dach'           | Ngiti              |

Die ersten drei Beispiele stammen aus Sprachen, in denen ziemlich lange Coda möglich sind. So haben wir im Beispiel aus dem Deutschen eine Coda, die aus vier Konsonanten besteht.<sup>11</sup> Im Svan gibt es sogar Silben, die fünf Konsonanten in der Coda haben können (die letzten beiden Konsonanten sind allerdings eigentlich nur ein Segment, eine Affrikate). Dann gibt es aber auch Sprachen, in denen überhaupt keine Silbencodas erlaubt sind. Die letzten drei Beispiele stammen aus Sprachen dieser Art: keine Silbe hat einen Konsonanten in der Coda.

Wir können also hier wieder eine universelle Tendenz feststellen. In allen Sprachen der Welt sind Silben ohne Coda erlaubt. In manchen Sprachen der Welt (wie im Hawaiianischen, Samoanischen und Ngiti), sind nur Silben ohne Coda erlaubt. Es gibt aber keine Sprache der Welt, in der Silben ohne Coda verboten sind.

Die unmarkierte Silbe scheint also eine Silbe zu sein, die nicht auf einen Konsonanten endet. Markierte Silben hingegen haben einen Konsonanten in der Coda.

Es gibt außerdem Sprachen, wie z.B. das Italienische, in denen nur ein Konsonant in der Coda stehen darf. Wir können also wie beim Onset-Prinzip beobachten, dass komplexe Coda markierter als einfache Coda sind.

Wir können diese Beobachtungen wiederum in zwei Bedingungen formulieren:<sup>12</sup>

- (45) Coda- Bedingung, Teil I: Unmarkierte Silben haben keine Coda  
 Coda- Bedingung, Teil II: Coda mit nur einem Konsonanten sind  
 unmarkierter als Coda mit mehreren Konsonanten

Wenn wir nun Onset- und Coda-Bedingung zusammen betrachten, dann müssen wir den Schluss ziehen, dass die universell unmarkierte Silbe eine CV-Silbe ist, also eine

<sup>11</sup> Wenn wir die Sprachen, aus denen diese Beispiele stammen, allerdings etwas genauer untersuchen, dann sehen wir, dass extrem komplexe Coda oft nur am Wortende auftauchen können. Das gilt auch für das Deutsche. Eine Coda, die aus der Sequenz *psts* besteht, finden wir zwar in wortfinalen Coda, nicht aber im Inneren eines Wortes. Dasselbe gilt oft für komplexe Onsets: am Anfang eines Wortes finden wir oft sehr komplexe Onsets, die im Inneren des Wortes nicht erlaubt sind.

<sup>12</sup> Für eine ausführliche Diskussion dieser universellen Prinzipien s. auch Jakobson (1962) und Prince&Smolensky (1993)

Silbe, die nur aus einem Konsonanten und einem Vokal besteht. In der Tat ist diese Silbe die einzige Silbe, die in allen Sprachen der Welt vorkommt!

Es gibt eine weitere Beobachtung, die darauf hinweist, dass es sich bei der CV-Silbe um die universell am wenigsten markierte Silbe handelt. Wenn Kinder anfangen, die ersten Wörter von sich zu geben, dann verkürzen sie diese Wörter gerne auf eine CV-Silbe. *Buch* wird dann zu *bu* und *Elefant* wird zu *fa*. Erst nach diesem CV-Stadium kommt dann ein weiteres Stadium, in dem Kinder auch Silbencodas zulassen. Und erst nach diesem zweiten Stadium kommt ein drittes Stadium, in dem Kinder schließlich auch komplexe Onsets und Codas aussprechen:

(46) Stadien im Spracherwerb:

- |             |       |             |
|-------------|-------|-------------|
| 1. Stadium: | CV    | <i>fa</i>   |
| 2. Stadium: | CVC   | <i>fan</i>  |
| 3. Stadium: | CCVCC | <i>fant</i> |

Wenn man es sich genau überlegt, dann ist dieser Spracherwerb in Stadien ziemlich überraschend, denn eigentlich können die Kinder die meisten Buchstaben schon aussprechen. Warum sagt das Kind *fa*, wenn es doch das *n* und das *t* der Coda schon aussprechen kann? Der Grund liegt darin, dass das Kind in seinen phonologischen Entwicklungsstadien generell mit unmarkierten Strukturen beginnt und erst nach und nach die markierteren Strukturen erwirbt.

Ein letzter Hinweis darauf, dass eine wohlgeformte Silbe einen Onset aber keine Coda besitzen sollte, bildet der Mechanismus der Silbifizierung. Wenn wir eine Sequenz von Lauten vor uns haben, dann gibt es im Prinzip mehrere Möglichkeiten, wie diese Sequenz silbifiziert werden könnte. Nehmen wir z.B.

- (47) Amerika: a. A.me.ri.ka  
b. \*Am.er.ik.a

Es ist klar, dass ein Wort wie *Amerika* immer wie in a. silbifiziert werden wird und nie wie in b. Das heißt, die Silbifizierung wird immer so vor sich gehen, dass jede Silbe, wenn möglich, einen Onset hat, aber keine Coda.

Der Mechanismus der Silbifizierung ist allerdings noch etwas komplizierter. Generell wird eine Lautsequenz so silbifiziert, dass der Onset maximiert wird. Das bedeutet, dass die Konsonanten, soweit möglich, alle einem Onset zugeordnet werden. Nur wenn ein Konsonant im Onset keine wohlgeformte Struktur mehr ergeben würde, wird er einer Coda zugeordnet. Wir sehen das im folgenden Beispiel:

- (48) Andreas: a. \*A.ndre.as  
b. An.dre.as  
c. \*And.re.as  
d. \*Andr.e.as

In a. wurde versucht, die Konsonanten /n, d/ und /r/ alle in den Onset der zweiten Silbe zu stecken. Das ist nicht möglich, denn die Sequenz /ndr/ ergibt keinen wohlgeformten Onset, da die Sonoritätshierarchie verletzt werden würde. Deshalb muss *n* in der Coda bleiben. Die anderen beiden Konsonanten können jedoch dem Onset zugeordnet werden, wie Beispiel b. zeigt. Beispiel c. und d. sind nicht wohlgeformt, da hier der Onset nicht maximiert wurde. In Beispiel c. wäre die erste Silbe *And* durchaus eine wohlgeformte Silbe, doch diese Silbifizierung wird nicht gewählt, da das *d* auch im Onset der darauffolgenden Silbe stehen kann (und also dort stehen muss).

#### 4.1.4 Die deutsche und die italienische Silbe und ihre Restriktionen

In diesem Kapitel wollen wir uns die Silbenstruktur des Deutschen und die Silbenstruktur des Italienischen anschauen und sie miteinander vergleichen. Wie sieht der deutsche Onset aus, wie sieht der italienische Onset aus? Gibt es Ähnlichkeiten oder Unterschiede was die Segmente betrifft, die in den Untereinheiten der Silbe auftreten können? Am Ende des Kapitels wollen wir uns einem Problem zuwenden, das jedem, der Deutsch und Italienisch spricht, sofort auffällt: deutsche Wörter sind kürzer als italienische Wörter.<sup>13</sup> Wir werden sehen, dass unsere kontrastiven Beobachtungen zur Silbenstruktur uns helfen können, diesen Unterschied zu erklären.

Beginnen wir mit dem Silbenonset im Deutschen. Der Onset der Silbe des Deutschen besteht aus höchstens drei, fast immer aber aus mindestens einem Konsonanten.

Wenn wir im Onset drei Konsonanten haben, dann ist der erste immer ein [s] oder ein [ʃ]. Solche maximale Onsets finden wir in der ersten Silbe der folgenden Beispiele:

|                      |            |
|----------------------|------------|
| (49) [ʃ]trumpf       | <Strumpf>  |
| [ʃ]prache            | <Sprache>  |
| [ʃ]plitter           | <Splitter> |
| [ʃ]pray oder [s]pray | <Spray>    |

Das Graphem <s> wird am Wortanfang vor einem Konsonanten meistens als [ʃ] ausgesprochen. Nur in einigen Fremdwörtern wie *Spray*, *Skat*, *Slalom*, *Smoking* wird es als [s] ausgesprochen.

Da der dritte Konsonant in einem Onset immer nur ein [s] oder ein [ʃ] sein kann, haben einige Phonologen angenommen, dass der deutsche Onset eigentlich höchstens aus zwei Konsonanten bestehen kann, und dass es sich bei dem [s] oder [ʃ] um ein sogenanntes **extrasilbisches Segment** handelt (s. Wiese 1996). Es gibt auch noch einen weiteren Grund, um das anzunehmen: wenn das [s] oder [ʃ] wirklich Teil des Onset wäre, dann würde es die Sonoritätshierarchie verletzen, da es als Frikativ vor einem Plosiv stehen kann.

Es gibt noch einige andere Beispiele, in denen es so aussieht, als gäbe es drei Konsonanten im Onset:

|                |            |
|----------------|------------|
| (50) [pfl]aume | <Pflaume>  |
| [pfr]opfen     | <Pfropfen> |
| [tsv]eig       | <Zweig>    |
| [tsv]anzig     | <zwanzig>  |

<sup>13</sup> Ich meine hier morphologisch einfache Wörter. Natürlich hat das Deutsche auch ziemlich lange Wörter, wenn man morphologisch komplexe Wörter in Betracht zieht, da das Deutsche sehr lange und komplexe Komposita zulässt.

Wenn man sich diese Wörter allerdings genauer ansieht, dann sieht man, dass es sich in allen Fällen um Sequenzen handelt, in denen am Wortanfang eine Affrikate [pf] oder [ts] steht. Die Tatsache, dass Affrikaten hier in einem Kontext stehen, in dem normalerweise nur ein einzelner Konsonant vorkommt (also im Onset vor einem anderen Konsonanten) ist ein gutes Argument dafür, dass es sich bei Affrikaten im Grunde um ein Einzelsegment handelt, nicht um zwei unterschiedliche Segmente.

Wenn es in der zugrundeliegenden Form am Anfang einer Silbe keinen Konsonanten gibt, dann wird ein glottaler Plosiv [ʔ] eingesetzt. So bekommt auch eine Silbe, die sonst mit einem Vokal beginnen würde einen Onset. Wir hören den glottalen Plosiv besonders gut in einem Satz, in dem alle Wörter mit einem Vokal beginnen:

(51) [ʔ]Anton [ʔ]aß [ʔ]am [ʔ]Abend [ʔ]immer [ʔ]Auflauf

Aber nicht nur am Wortanfang finden wir den glottalen Plosiv. Er wird von den meisten Sprechern auch vor eine vokalinitiale Wurzel gesetzt:

(52) ver-[ʔ]arbeiten  
ver-[ʔ]achten  
be-[ʔ]eilen  
[ʔ]auf-[ʔ]essen

Sprecher des Italienischen müssen bei der Aussprache dieser Wörter besonders aufpassen. Das Italienische hat keinen glottalen Plosiv, deshalb neigen Sprecher des Italienischen dazu, in diesen Wörtern die Coda des Präfixes als Onset für die Wurzelsilbe zu verwenden. Diese Aussprache klingt aber für ein deutsches Ohr ziemlich seltsam:

(53) Italienische Aussprache vokalinitialer Wurzeln:  
\*ve.rar.bei.ten  
\*ve.rach.ten  
\*au.fes.sen

Es gibt noch einen Kontext, in dem wir den glottalen Plosiv finden. Auch im Inneren eines Morphems kann es vorkommen, dass eine Silbe mit einem Vokal beginnt. Das geschieht dann, wenn zwei Vokale aufeinander treffen und einen sogenannten **Hiatus** bilden:

(54) The.[ʔ]á.ter  
cha.[ʔ]ó.tisch  
Be.[ʔ]á.te  
in.di.vi.du.[ʔ]éll

Süddeutsche Sprecher produzieren in diesem Kontext eigentlich selten einen glottalen Plosiv, aber wer nicht aus dem Süden kommt, wird in diesen Wörtern wahrscheinlich diesen epenthetischen Konsonanten einsetzen.

Beachtet, dass in diesen Beispielen der Vokal, vor dem der glottale Verschluss steht, immer betont ist. Es gibt nämlich einen Kontext, in dem morphemintern normalerweise kein glottaler Plosiv produziert wird, obwohl die Silbe mit einem Vokal beginnt. Das geschieht, wenn der silbeninitiale Vokal nicht betont ist:

|      |          |           |                       |
|------|----------|-----------|-----------------------|
| (55) | [zé:.ən] | *[zé:ʔən] | <sehen>               |
|      | [ʔé:.ə]  | *[ʔé:ʔə]  | <Ehe>                 |
|      | [bé:.a]  | *[bé:ʔa]  | <Bea>                 |
|      | [ká:ɔs]  | *[ká:ʔɔs] | <Chaos> <sup>14</sup> |

Der glottale Plosiv steht also bei vokalinitialen Silben in folgenden Kontexten:

- am Wortanfang
- am Anfang einer Wurzel
- morphemintern, wenn der darauffolgende Vokal betont ist

Wir sehen, auch wenn es so aussieht, als gäbe es im Deutschen Silben ohne Onset, so müssen wir doch schließen, dass das in den meisten Fällen nicht so ist, da in diesen Fällen im Onset ein glottaler Plosiv steht.

Der Onset des Deutschen richtet sich im Großen und Ganzen nach der Sonoritätshierarchie. Allerdings finden wir nicht alle Sequenzen, die nach der Sonoritätshierarchie erlaubt wären. Die folgenden Sequenzen müssten eigentlich möglich sein, es gibt aber keine Onsets dieser Art im Deutschen:

- (56) \*kfa  
 \*mra  
 \*tla

Wir haben schon weiter oben gesehen, dass der Silbennukleus im Deutschen aus einem kurzen oder langen Vokal, einem Diphthong oder einem silbischen Konsonanten bestehen kann:

|      |                          |          |          |
|------|--------------------------|----------|----------|
| (57) | a. Kurzer Vokal:         | W[a].che | <Wache>  |
|      | b. Langer Vokal:         | L[i:].be | <Liebe>  |
|      | c. Diphthong:            | L[au].be | <Laube>  |
|      | d. Silbischer Konsonant: | war.[tɪ] | <warten> |
|      |                          | ne.[bm]  | <neben>  |
|      |                          | Fle.[gl] | <Flegel> |

In der Coda können im Deutschen im Allgemeinen maximal zwei Konsonanten stehen, wobei sich ihre Abfolge wieder nach den Prinzipien der Sonoritätshierarchie richtet:

- (58) Werk  
 halb  
 Hand  
 Helm

Es gibt allerdings Wörter, in denen stehen nach dem Nukleus mehr als zwei Konsonanten (s. Wiese 1996: 48):

<sup>14</sup> Es gibt Sprecher, die auch in *Chaos* einen glottalen Verschluss verwenden. Es könnte sein, dass auf der zweiten Silbe dieses Wortes, -os-, ein (schwächerer) Nebenakzent liegt, und dass einige Sprecher glottale Plosive auch vor Vokale setzen, die einen Nebenakzent tragen.

- (59) Obst  
Markt  
Herbst  
Werft

Es ist interessant, dass der dritte und vierte Konsonant in diesen Fällen immer ein [s] oder ein [t] ist. Es können also nur alveolare Konsonanten als dritte und vierte Konsonanten in der Coda stehen. Da die Art der Segmente, die nach dem zweiten Konsonanten in der Coda stehen können so beschränkt ist, wurde auch hier vorgeschlagen, dass diese alveolaren Segmente als extrasilbisch zu werten sind.

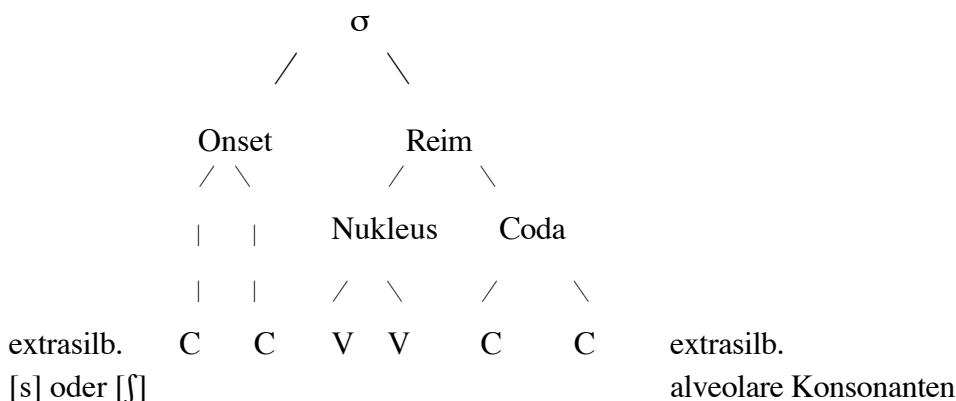
Was die Affrikaten betrifft, so können wir auch in der Coda eine interessante Beobachtung machen. Schauen wir uns die folgenden Beispiele an:

- |              |          |
|--------------|----------|
| (60) Pro[pf] | <Pfropf> |
| Ko[pf]       | <Kopf>   |
| Wi[ts]       | <Witz>   |
| Sa[ts]       | <Satz>   |

Die Affrikaten [pf] und [ts] bestehen aus einem Plosiv-Teil ([p], bzw. [t]) und einem Frikativ-Teil ([f], bzw. [s]). Wenn diese beiden Teile als je ein Segment zählen würden, dann würden [pf] und [ts] in der Coda die Sonoritätshierarchie verletzen. Da wir jedoch diese Affrikaten, wie in den obigen Beispielen, in der Silbencoda finden, müssen wir annehmen, dass Affrikaten nur als ein Segment zählen. Wir haben also ein weiteres Argument für den monosegmentalen Status der Affrikaten gefunden.

Fassen wir noch einmal die gerade genannten Generalisierungen zur deutschen Silbe zusammen. Die maximale deutsche Silbe kann man schematisch folgendermaßen darstellen:

- (61) Die maximale deutsche Silbe:



Die maximale Silbe im Deutschen besteht also aus zwei Onset-Konsonanten, einem langen Vokal oder Diphthong im Nukleus und zwei Coda-Konsonanten. Dazu können noch extrasilbische Sibilanten am Silbenanfang oder alveolare Konsonanten am Silbenende dazukommen.

Die minimale Silbe im Deutschen hat mindestens einen Nukleus und (außer in einigen wenigen Fällen wie die zweite Silbe in *sehen*) eigentlich auch immer einen Onset, der mindestens aus einem epenthetischen glottalen Plosiv besteht.

Wenden wir uns nun der italienischen Silbe zu. Wir wollen versuchen, die Charakteristiken der italienischen Silbenstruktur anhand einer Übung selbst zu erarbeiten. Wir können uns dabei an der Beschreibung der deutschen Silbenstruktur orientieren. Die italienischen Daten und Generalisierungen stammen weitgehend aus Nespor (1993).

### Übung 8

1. Als erstes wollen wir herausfinden, wie der Onset der italienischen Silbe aussieht. Sehen wir uns dazu die folgenden Daten an:

- (62) [pr]anzo            [str]ano  
       [kl]ima            [spr]uzzo  
       [tr]eno            [skr]ivere  
       [gl]icino          [zgr]adevole  
       a.[tl]e.ta

Beantwortet nun die folgenden Fragen:

- Wieviele Konsonanten können im Italienischen maximal im Onset stehen?
- Muss mindestens ein Konsonant im Onset stehen?
- Welche Konsonanten können bei einem maximalen Onset in der ersten Position stehen, welche in der zweiten, welche in der dritten?
- Welche Gemeinsamkeiten bzw. Unterschiede gibt es zum Onset des Deutschen?
- Befolgen die Onsets des Italienischen die Sonoritätshierarchie?

2. Wenden wir uns nun dem Nukleus zu:

- (63) [a].ra.bo  
       [au].la  
       p[ie].no

- Welche Segmente können im Italienischen die Nukleusposition besetzen?
- Welche Unterschiede zum Deutschen lassen sich hier feststellen?

3. Und nun zur Coda:

- (64) pa[r].te            fa[t].to  
       a[l].to            le[g].go  
       ba[n].da          ma[s].sa  
       ba[m].bi.no        a[v].vi.so

pa[s].ta oder pa.[s]ta?

- Wieviele Konsonanten können im Italienischen maximal in der Coda stehen?
- Welche Art von Konsonanten kann in der Coda stehen?
- Welche Unterschiede gibt es zwischen der italienischen und der deutschen Silbencoda?

Wenn ihr diese Übung in allen ihren Teilen ausgearbeitet habt, dann solltet ihr die folgenden Dinge beobachtet haben. Der italienische Onset gleicht in großen Zügen dem deutschen Onset. Auch im Italienischen haben wir maximal zwei Konsonanten, die die Onset-Position besetzen können, plus einen dritten Konsonanten, der aber

dann ein /s/ oder /z/ sein muss. Diese Sibilanten könnten genauso wie im Deutschen als extrasilbisch gewertet werden. Das erste Element des maximalen Onset ist also immer ein /s/ oder /z/. Das zweite Element kann ein beliebiger Konsonant sein. Das dritte Element – und hier unterscheidet sich das Italienische wieder vom Deutschen – ist immer ein Liquid, also ein /r/ oder ein /l/. Es gibt noch einen Unterschied zum Deutschen. Während wir im Deutschen vor vokalinitiale Silben meistens einen glottalen Plosiv setzen, tun wir das im Italienischen nicht. Im Italienischen gibt es also häufiger Silben ohne Onset als im Deutschen.

Schließlich können wir noch beobachten, dass es im Italienischen Onsets gibt, die es im Deutschen nicht gibt, wie z.B. den Onset [tl] im Wort *a.tle.ta*.

Der Nukleus des Italienischen ist dem des Deutschen ziemlich ähnlich. In beiden Sprachen kann er aus einem einfachen Vokal oder auch aus einem Diphthong bestehen. Im Unterschied zum Deutschen können allerdings im Nukleus einer italienischen Silbe keine silbischen Konsonanten stehen.

In der Coda der beiden Sprachen finden wir wahrscheinlich die größten Unterschiede. Im Italienischen kommt in der Coda immer nur höchstens ein Konsonant vor. Dieser Konsonant kann außerdem nur ein Sonorant (also ein Nasal oder ein Liquid), die erste Hälfte einer Geminale oder ein [s] sein. Im Deutschen hingegen kann die Coda aus mehreren Konsonanten bestehen.

Einen etwas problematischen Fall stellen wortinterne Konsonantensequenzen wie in den folgenden Beispielen dar:

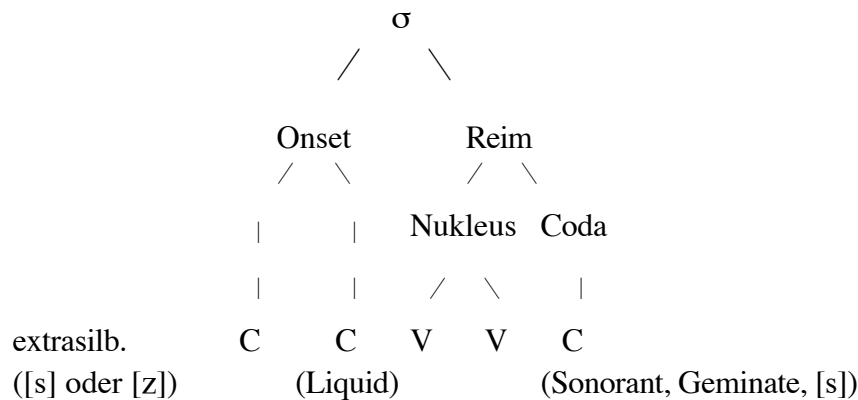
- (65) As.co.li    oder    A.sco.li?  
       pas.ta    oder    pa.sta?  
       a.spro    oder    as.pro?

Es ist in diesen Fällen nicht ganz klar, ob das [s] zur Coda der vorhergehenden Silbe oder zum Onset der nächsten Silbe gezählt werden soll. In der Schule lernt man, dass man diese Wörter als *A.sco.li* bzw. *pa.sta*, *a.spro* trennen soll. Aber es gibt gute phonologische Argumente dafür, dass die richtige Silbentrennung eigentlich *As.co.li* bzw. *pas.ta*, *as.pro* ist (s. Nespor 1993: 176). Wir können nämlich beobachten, dass der Vokal vor dem [s] immer kurz ist. Wenn die Trennung jedoch *A.sco.li*, *pa.sta*, *a.spro* wäre, dann müsste der Vokal verlängert werden, denn Vokale in offenen, betonten Silben werden im Italienischen immer verlängert. Wir können das in Beispielen wie den folgenden sehen, in denen die betonten Vokale alle in offenen Silben stehen und deshalb verlängert werden:

- (66) cá.sa  
       á.ra.bo  
       prá.to

Wir können also die Struktur der italienischen maximalen Silbe folgendermaßen darstellen:

(67) Die maximale italienische Silbe:



Wir wenden uns hier zum Abschluss noch einem interessanten Problem zu, das die Länge der Wörter des Italienischen und des Deutschen betrifft. Wer die beiden Sprachen auch nur oberflächlich kennt, bemerkt bald, dass das Italienische eigentlich ziemlich lange Wörter hat, während deutsche Wörter meistens aus einer einzigen Silbe bestehen. Wir sprechen hier nur von morphologisch einfachen Wörtern, denn im Deutschen kann man natürlich sehr lange morphologisch komplexe Wörter bilden, wenn man z.B. an Komposita wie *Donaudampfschiffahrt* denkt. Doch es ist gar nicht so einfach, mehrsilbige deutsche Wörter zu finden, die keine Komposita sind und keine Suffixe oder Präfixe enthalten. Im Bereich der Fremdwörter gibt es natürlich auch mehrsilbige Wörter, wie z.B. *Universität*, *individuell* u.s.w., aber die mehrsilbigen Wörter, die nicht eindeutige Fremdwörter sind, sind eher selten. Es handelt sich um so Wörter wie *Arbeit*, *Ameise*, oder auch Wörter, die auf *-er* und *-el* enden, wie *Engel*, *Muster* oder ähnliches. Die meisten typisch deutschen Wörter bestehen jedoch aus nicht mehr als einer Silbe. Im Italienischen ist es umgekehrt: das typische italienische Wort besteht aus mehr als einer Silbe. Wie kann man sich den Unterschied zwischen den beiden Sprachen erklären?

Ein Grund für die Kürze der deutschen Wörter könnte in der Silbenstruktur liegen. Nehmen wir einmal an, dass es, um einen Begriff in einem Wort auszudrücken, im Durchschnitt sieben Segmente braucht. Im Deutschen können diese sieben Segmente ohne weiteres in eine einzige Silbe gezwängt werden, wenn wir z.B. eine Silbe nehmen, die ein oder mehrere extrasilbische Segmente hat. Im Italienischen brauchen wir für die gleichen sieben Segmente hingegen mindestens zwei Silben, da in einer italienischen Silbe maximal fünf Segmente Platz haben.

Ein weiterer Grund für die "Mehrsilbigkeit" des Italienischen liegt sicherlich darin, dass italienische (lexikalische) Wörter immer mit einem Vokal enden müssen. Ein einsilbiges Wort kann diese Bedingung nur erfüllen, wenn die Silbe eine (C)CV-Struktur hat. In einer Struktur dieser Art kann man aber sehr wenige Segmente unterbringen (höchstens vier, wenn man die extrasilbischen Elemente dazuzählt). Es handelt sich hier also um eine nicht sehr "informative" Silbe.

#### 4.1.5. Verletzbare Bedingungen - Deutsch und Italienische kontrastiv

Wir haben nun gesehen, wie die deutsche und die italienische Silbe aussehen. In diesem Kapitel wollen wir uns fragen, wie sich die Silben dieser Sprachen in Bezug auf die universellen Onset- und Coda-Bedingungen verhalten.

Auf den ersten Blick scheint die Antwort auf diese Frage banal zu sein: beide Sprachen scheinen diese Bedingungen nicht zu befolgen.

Im Deutschen haben wir Beispiele wie das folgende:

(68) [ze:.ən] <sehen>

In diesem Beispiel hat die zweite Silbe keinen Onset, sie verstößt also gegen die Onset-Bedingung (1. Teil) und sie hat zugleich eine Coda, verstößt also gegen die Coda-Bedingung (1. Teil).

Im Italienischen haben wir Beispiele wie:

(69) am.ba.scia.ta

wo die erste Silbe keinen Onset, wohl aber eine Coda hat.

Doch auch wenn die beiden universellen Bedingungen in vielen Wörtern verletzt werden, so können wir doch sehen, dass sie auch in diesen beiden Sprachen eine gewisse Rolle spielen.

Wir werden sehen, dass beide Sprachen, wann immer möglich, die beiden Bedingungen befolgen, und dass sie dazu bestimmte Strategien verfolgen.

Beginnen wir mit dem Deutschen. Wir haben im Deutschen Wörter, die mit einem Vokal beginnen. Wenn dieser Vokal am Wortanfang steht, dann wird vor ihm ein glottaler Plosiv eingesetzt. Das Wort *Esel* wird also folgendermaßen realisiert:

(70) /e:səl/ → [ʔe:səl]

Der glottale Plosiv sorgt dafür, dass die erste Silbe einen Onset bekommt. Wortintern kann der glottale Plosiv nur vor einen Vokal gesetzt werden, wenn dieser betont ist.

Im Beispiel [ze:.ən] kann also kein glottaler Plosiv vor das schwa gesetzt werden, da dieses nicht betont ist.

Prinzipiell könnte man sich für das Wort /e:səl/ verschiedene Strategien der Realisierung vorstellen:

(71) /e:səl/ →      a. [ʔe:səl]  
                              b. \*[səl]  
                              c. \*[e:səl]  
      /pak e:səl/ →    d. \*[pa.k e:.səl]      <Packesel>

Wir können, wie in a., einen Konsonanten für die vokalinitiale erste Silbe einfügen. Dies ist die Strategie, die schließlich angewandt wird. Wir könnten uns aber auch vorstellen, dass man vielleicht den ersten Vokal streichen könnte, wie in b. Durch diese Streichung erreichen wir, dass nur eine Silbe mit Onset übrigbleibt. Wir könnten auch die zugrundeliegende Form so lassen, wie sie ist, und damit die Onset-Bedingung verletzen, wie in c. Schließlich könnten wir, wenn dem Wort *Esel* noch ein anderes Wort vorangeht, den Weg der Resilbifizierung wählen, wie in d. Das bedeutet, dass wir den Konsonanten des ersten Wortes nehmen und ihn als Onset für



## Übung 9

1. Wir verwenden in dieser Übung die Beispiele *amore* und *con amore*. Erstellt nun mit diesen Beispielen eine Liste von möglichen outputs, wie wir sie im letzten Beispiel (*Esel*) für das Deutsche gesehen haben. Wie könnten *amore* und *con amore* prinzipiell realisiert werden? Ihr solltet alle vier Möglichkeiten anführen: Einfügen eines Segmentes, Streichen eines Segmentes, keine Änderung und Resilbifizierung.

- a. \_\_\_\_\_
- b. \_\_\_\_\_
- c. \_\_\_\_\_
- d. \_\_\_\_\_

Welche dieser Realisierungen gibt es nun im Italienischen wirklich, welche sind nicht möglich?

2. Gebt nun an, welche Bedingungen in diesen vier Realisierungen verletzt werden:

- a. \_\_\_\_\_
- b. \_\_\_\_\_
- c. \_\_\_\_\_
- d. \_\_\_\_\_

Welche Bedingungen können also im Italienischen verletzt werden, welche nicht?

3. Versucht nun, die Hierarchie der Bedingungen für das Italienische zu bauen. In dieser Hierarchie werden die Bedingungen, die verletzt werden können, von den Bedingungen, die nicht verletzt werden können, dominiert. Stellt diese Hierarchie auf die gleiche Art und Weise wie im Deutschen dar.

4. Die Hierarchie des Italienischen ist allerdings etwas komplizierter als die im Deutschen. Schaut euch noch mal die beiden Strategien an, die im Italienischen möglich sind, und die Bedingungen, die durch sie verletzt werden. Lässt sich zwischen diesen beiden Bedingungen noch eine Subhierarchie feststellen? Welche Strategie wird in *con amore* bevorzugt? Welche von den beiden relevanten Bedingungen wird also eher verletzt?

Wenn ihr diese Übung in all ihren Schritten gelöst habt, dann solltet ihr zu dem Schluss gelangt sein, dass im Italienischen die beiden Bedingungen *Streiche nichts!* und *Füge nichts hinzu!* am höchsten stehen. Die Onset-Bedingung und die Bedingung *Keine Resilbifizierung!* hingegen müssen von diesen beiden Bedingungen dominiert werden, denn sie werden in den outputs *amore* und *co.n a.mo.re* verletzt. Wenn wir versuchen, eine Hierarchie zwischen der Onset-Bedingung und *Keine Resilbifizierung!* herzustellen, dann müssen wir uns den output *co.n a.mo.re* anschauen. In diesem output wird die Bedingung *Keine Resilbifizierung!* verletzt, während die Onset-Bedingung befolgt wird. *Keine Resilbifizierung!* muss also von der Onset-Bedingung dominiert werden. Wir erhalten also für das Italienische die folgenden Hierarchie:

(74) *Streiche nichts!*      *Füge nichts hinzu!*

|  
Onset-Bedingung

|  
*Keine Resilbifizierung!*

Die beiden Bedingungen *Streiche nichts!* und *Füge nichts hinzu!* stehen am höchsten in der Hierarchie, denn sie werden immer befolgt. Die Onset-Bedingung und *Keine Resilbifizierung!* stehen tiefer, denn sie werden in einigen Fällen verletzt. Wenn man aber die Wahl zwischen diesen beiden Bedingungen hat, wie im Fall von *con amore*, dann entscheiden wir uns dafür, die Onset-Bedingung auf Kosten von *Keine Resilbifizierung!* zu befolgen. Wir resilbifizieren also die beiden Wörter als *co.n a.mo.re*, anstatt das zweite Wort ohne Onset zu lassen (wie in *con amore*).

Was haben wir nun eigentlich in diesem Vergleich zwischen italienischen und deutschen Onset-Bedingungen gesehen? Wir haben gesehen, dass die Bedingungen, die die Silbenstruktur bestimmen, als universell angesehen werden können. Was bedeutet universell? Es bedeutet, dass wir ihren Einfluss in den meisten, vielleicht in allen Sprachen der Welt finden. Es bedeutet aber nicht, dass diese Bedingungen immer und überall gelten. Sie können auch verletzt werden, aber nur, wenn es dadurch möglich ist, wichtigere Bedingungen zu erfüllen.

Die Bedingungen, die die Silbenstruktur bestimmen, sind universell, aber die Hierarchie ihrer Wichtigkeit ist sprachspezifisch. So ist z.B. die Onset-Bedingung sowohl im Deutschen als auch im Italienischen ein wichtiges Prinzip. Aber im Deutschen steht sie höher in der Hierarchie als im Italienischen, sie wird im Deutschen also eher befolgt werden, als im Italienischen. Die **Optimalitätshierarchie** (Prince&Smolensky 1993) hat diese Ideen in den letzten 10 Jahren ausgearbeitet:

- die Idee, dass die Wohlgeformtheit von sprachlichen Strukturen durch universelle Bedingungen (constraints) bestimmt wird;
- dass diese universellen Bedingungen verletzbar sind;
- dass diese Bedingungen in jeder Sprache in einer sprachspezifischen Hierarchie angeordnet sind.

### Übung 10

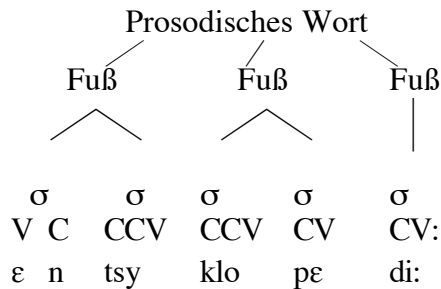
Kennt ihr die Romane von Andrea Camilleri, in denen der commissario Montalbano die Hauptrolle spielt? Im commissariato von Montalbano gibt es den Polizisten Catarella, der manchmal eine etwas eigenartige Sprache spricht. So sagt er z.B. "Dottore! È arrivato un facchisi!" Er meint damit, dass ein *Fax* angekommen ist.

1. Transkribiert die Wörter *fax* und *facchisi*
2. Was hat Catarella mit dem Wort *fax* gemacht? Wie sieht die Silbenstruktur des Wortes *fax* aus, wie sieht die Silbenstruktur von *facchisi* aus?
3. Welche Silbenstrukturbedingungen verletzen die Wörter *fax* und *facchisi*, welche befolgen sie?
4. Wie sieht also die Grammatik Catarellas aus, welche Bedingungen stehen hoch in der Hierarchie, welche stehen tief?
5. Ist Catarellas Grammatik mit der des Standarditalienischen identisch?



Sie sind sich aber noch nicht sehr einig, wie genau die Füße des Deutschen auszusehen haben. Bilden immer zwei Silben einen Fuß oder manchmal nur eine?

Segmente verbinden sich zu Silben, Silben verbinden sich zu Füßen und Füße verbinden sich am Ende zu einem prosodischen Wort, das ungefähr dem grammatikalischen Wort gleichkommt (die schwierigen Fälle, wo das nicht so ist, interessieren uns hier nicht). Wir haben also am Ende für ein Wort die folgende **prosodische Struktur**:



#### 4.2.2 Kurzer Kommentar zum deutschen Akzent

**Fremdwörter im Deutschen:** haben ähnlich wie im Italienischen eine alternierende Akzentstruktur, wobei der Hauptakzent auf eine der letzten drei Silben fällt:

- (77)    s`y.ste.má.tisch      sì.ste.má.ti.co  
          Èn.zy.klò.pä.díe      èn.ci.clò.pe.dí.a  
          ù.ni.ver.sál          ù.ni.ver.sá.le

##### **native Wörter:**

- **morphologisch einfache Wörter** sind im Deutschen sehr kurz. Meistens besteht die Wurzel eines Wortes nur aus einer einzigen Silbe. Folglich fällt der Akzent auf diese Silbe. Wenn das Wort allerdings aus mehreren Silben besteht, dann fällt der Hauptakzent auf die erste dieser Silben:

- (78)    Árbeit            nicht: \*Arbéit  
          Ántwort        nicht: \*Antwórt  
          Úrlaub           nicht: \*Urláub

##### - **morphologisch zusammengesetzte Wörter:**

- (79)    untrennbare Präfixe: Hauptakzent auf der Wurzel  
          ver-tréiben  
          zer-bréchen  
          er-láuben  
          über-sétzen (tradurre)
- (80)    trennbare Präfixe: verhalten sich wie Komposita: Hauptakzent fällt auf das Präfix, Nebenakzent auf die Wurzel  
          áuf-hálten  
          úm-dénken  
          ´über-sétzen (traghettare)
- (81)    Komposita: Akzent fällt auf den ersten Teil:  
          Léhrstuhl  
          Hóchschùle  
          Kíndergärten

- (82) Arbiträr:  
 frohlócken                      aber:                      réchtfèrtigen  
 vollbríngen                      aber:                      fr´úhst`ücken

- (83) schwa [ə] und vokalisiertes /r/ [ʁ] tragen nie einen Akzent:  
 die Heitereren =                      die Heit[ə]r[ə]r[ə]n

## 5. Laute und Schrift

Im Deutschen gibt die Orthographie die Laute der Sprache ziemlich genau, aber doch nicht perfekt wieder. So gibt es z.B. Sequenzen von Graphemen (Zeichen der Schrift), denen nur ein Phonem entspricht:

- <sch> [ʃ]                      'schlafen'                      3 Grapheme, ein Phonem  
 <ch> [x] oder [ç]                      'ich, ach'                      2 Grapheme, ein Phonem, 2 Allophone

Ein weiterer Fall ähnlicher Art sind die Grapheme <f> und <v>: sie drücken im Deutschen immer denselben Laut [f] aus:

- <für> ['fy:ʁ]                      2 Grapheme, ein Phonem  
 <viel> ['fi:l]

Umgekehrt gibt es Fälle, in denen wir zwei Laute haben, aber dieser Unterschied wird in der Orthographie nicht ausgedrückt. So gibt es z.B. kein orthographisches Zeichen für das stimmhaft [z]. Wir schreiben sowohl für [z] als auch für [s] ein <s>. Das "scharfe" <ß> (auch ess-zett genannt) entspricht allerdings immer einem stimmlosen [s].

- <See>                      ['ze:]                      zwei Phoneme, ein Graphem  
 <das>                      ['das]

Auch die Grapheme <z> und <x> drücken in Wahrheit zwei Laute aus:

- <zahm>                      ['tsa:m]                      zwei Laute, ein Graphem  
 <Hexe>                      ['hɛksə]

Ein weiteres Beispiel für die Nicht-Entsprechung zwischen Phonemen und Graphemen sind die Doppelkonsonanten. Im Deutschen gibt es keine Geminaten, d.h., es gibt keine Konsonanten, die phonetisch gesehen lang sind. In der Schrift gibt es aber sehr viele Doppelkonsonanten. Diese orthographischen Doppelkonsonanten drücken nicht die Länge der Konsonanten aus, sondern die Kürze des vorhergehenden Vokals:

- <Mitte>                      ['mɪtə]  
 <Mappe>                      ['mapə]  
 <Ecke>                      ['ɛkə]

<Spitze>      [ʃpɪtsə]<sup>16</sup>

Auch bei den Diphthongen ist die Orthographie etwas eigenartig. So schreibt man <eu>, wie z.B. in <euch>, aber dieser Diphthong setzt sich nicht aus den Lauten [e] und [u] zusammen, sondern aus [ɔ] und [ɪ].

~~~~~

ÜBUNG zur Orthographie:

Wie wird in der deutschen Orthographie die Vokallänge wiedergegeben? Untersucht dazu folgende Daten:

dir
viel
ihr
See

~~~~~

Die deutsche Orthographie wurde von Mönchen im Mittelalter entwickelt, auf der Basis der lateinischen Schrift. Ende des 19. Jahrhunderts und Anfang des 20. Jahrhunderts gab es dann einige Reformen in der Orthographie und deren Festlegung in einer Norm. Einige orthographische Eigenheiten wurden eliminiert, aber andere (siehe oben) sind geblieben. Die Orthographie ist eben auch das Resultat einer historischen Entwicklung. Wir können sie nicht zu einer perfekten Übereinstimmung mit den Phonemen bringen, ohne die Sprachgemeinschaft zu verstören. Stellt euch nur vor, wie das wäre, wenn es plötzlich ein Gesetz gäbe, das befiehlt, dass man alle Texte in phonetischer Transkription schreiben muss. Allerdings sind die deutsche und die italienische Rechtschreibung doch noch relativ nahe an der Lautstruktur der Wörter, im Vergleich z.B. zu der englischen Orthographie.

---

<sup>16</sup> Komischerweise gibt es in der deutschen Orthographie keine Doppelkonsonanten <kk> und <zz>. Stattdessen wird <ck> und <tz> geschrieben.

## Appendix

### Kommentar zur Transkription des Deutschen:

**Die beste Grundlage für gutes Transkribieren ist gutes Hinhören. Wenn aber eine Stelle kommt, an der man den eigenen Ohren nicht recht traut, dann können vielleicht diese Hinweise weiterhelfen:**

- **phonetische Transkription kommt immer zwischen eckige Klammern** ( [ ... ] )! Phoneme kommen zwischen Schrägstriche ( / .... / ) und orthographische Zeichen kommen zwischen spitze Klammern ( < ... > ). Der Hauptakzent wird mit einem Apostroph vor der betonten Silbe gekennzeichnet (z.B. [t<sup>h</sup>a:l]), oder mit einem akuten Akzent "´" (z.B. [t<sup>h</sup>á:l]).

- **Achtung! Passt auf die Klein- und Großschreibung auf!** Große Buchstaben bedeuten im phonetischen Alphabet etwas anderes als Kleinbuchstaben! z.B.

[R] = uvularer Vibrant

[r] = apikaler Vibrant

- **Vokalqualität und Vokalquantität:** manchen Studenten fällt es schwer, die Vollvokale korrekt zu transkribieren. Vielleicht hilft euch dabei die Generalisierung, über die wir schon im Unterricht gesprochen haben:

außer in Fremdwörtern besteht im Deutschen folgende Korrelation: gespannte Vokale sind lang, ungespannte Vokale sind kurz. Wenn ich also ein Wort wie <Düne> vor mir habe, und ich weiß nicht, wie ich den Wurzelvokal transkribieren soll, dann versuche ich zuerst einmal festzustellen, ob dieser Vokal kurz oder lang ist. Er ist eindeutig lang (sonst hätte ich ja das Wort <dünne>). Deshalb transkribiere ich

[y:] (und in <dünne> transkribiere ich <ü> als [ʏ]). Und vergesst nicht den Doppelpunkt für die Vokallänge! Es gibt eine einzige Ausnahme zu der oben genannten Generalisierung: das ungespannte [ɛ] kann manchmal auch lang sein (also [ɛ:]), wie z.B. in "Ähre")

- **das schwa (= [ə]) in den Flexionsendungen:** in den typischen deutschen Flexionsendungen wird immer ein schwa gesprochen, nicht ein [ɛ] oder gar ein [e]. (Oder aber es wird überhaupt kein Vokal, sondern ein silbisches [ŋ], [m] oder [l] gesprochen). Ein schwa taucht auch auf bei Nomen, die auf <e> enden, auch wenn es sich hier nicht um ein Flexionsmorphem handelt. Also <Mühe> = [my:ə], nicht [my:e] oder [my:ɛ]

- **ein typischer Fehler: Verwechseln von Orthographie und phonetischen Zeichen.** Im Deutschen *schreiben* wir <v>, aber dieser *Buchstabe* entspricht meist der *Aussprache* [f], wird also mit dem *phonetischen Zeichen* [f] transkribiert. z.B.

<ver-> in <veranlassen> = [fɛə]

Wer sich hier nicht sicher fühlt, sollte das IPA noch einmal durchgehen und die Zeichen heraussuchen, die nicht den orthographischen Zeichen des Deutschen entsprechen.

- **vergesst die glottalen Plosive (= [ʔ]) nicht!** glottale Plosive werden mindestens in folgenden Kontexten gesetzt:

- vor Vokalen am Wortanfang
- vor Vokalen am Anfang einer Wurzel  
(z.B. [ʔ]er[ʔ]arbeiten)

Die meisten Leute machen außerdem noch einen im Wortinnern, wenn die darauffolgende Silbe den Hauptakzent trägt:

z.B. The[ʔ]áter

Einige wenige (?) machen auch einen glottalen Plosiv, wenn die darauffolgende Silbe einen Nebenakzent trägt:

z.B. Oze[ʔ]ànographíe

N.B.: glottale Plosive werden im Aussprache-Duden nicht transkribiert. Da das aber eine wichtige Charakteristik des Deutschen ist, will ich sie in eurer Transkription sehen

#### - **r-Laute am Silbenanfang:**

- [r] (= apikaler Vibrant = "Zungenspitzen-r"): gibt es im Deutschen so gut wie gar nicht, außer in einigen Dialekten und in der Bühnenaussprache der dreißiger Jahre (s. Hitler).
- [R] (= uvularer Vibrant): kommt selten vor, meist auch eher bei Dialektprechern, oder höchstens am Wortanfang.
- [ʀ] (= uvularer Frikativ): ist wahrscheinlich das gebräuchlichste "r" im Deutschen (bitte das [ʀ] mit Bauch nach rechts!)

**am Silbenende:** am Silbenende hat das Phonem /r/ ein vokalisiertes Allophon [ɐ], d.h., am Silbenende wird jedes /r/ zu [ɐ]. <ver> wird also als [fɛɐ] transkribiert, nicht als [fɛʀ], [fɛR], oder gar [fɛr]! Ich bin bereit, etwas toleranter zu sein, wenn nach dem /r/ vor Silbenende noch ein Konsonant kommt<sup>17</sup>, wie z.B. in

<Antarktis> = [ʔant.ʔaʀk.tɪs] oder [ʔant.ʔaʀk.tɪs]

d.h., in dieser Position benutzen manche Sprecher auch einen uvularen Frikativ. N.B.: der Duden ist ziemlich schlampig bei der Transkription der <r's>, in diesem Fall kann man sich nicht auf ihn verlassen.

- **das "stumme h" heißt deshalb so, weil es nicht ausgesprochen wird.** Die typische Situation für jemanden, der Deutsch lernen will, sieht so aus: er weiß nicht, ob man die <h's>, die da immer zwischen Vokalen *geschrieben* stehen, ausspricht, oder nicht. Also fragt er einen Muttersprachler. Der sagt ihm das Wort schön langsam vor, und macht dabei ein ganz deutlich zu hörendes [h]. Also macht das der

<sup>17</sup> Silbengrenzen gebe ich mit einem Punkt "." an

Deutschlernende auch, und fällt deshalb durch seine lächerliche Aussprache auf.

Wenn man Deutsch lernt, sollte man wissen:

- beim ganz langsamen Sprechen (z.B. Diktieren) wird in Wörtern wie <Mühe> oft ein [h] gesprochen
- bei der normalen Aussprache fällt das [h] vollkommen weg. Also:

<Mühe> = [my:ə] nicht [my:hə]

### - Aspiration:

Plosive werden mindestens im folgenden Kontext aspiriert:

- wenn sie alleine am Wortanfang stehen, vor betonter Silbe  
z.B. in [t<sup>h</sup>]ip

Der einzige Fall, wo man das vielleicht nicht so gut hört, ist bei [k], wie in  
[k<sup>h</sup>]int

Am Wortende ist Aspiration optional, also

Ti[p<sup>(h)</sup>]  
Kin[t<sup>(h)</sup>]

**- wie transkribiert man Diphthonge?** Das Deutsche hat drei Diphthonge, die als <ei> (auch <ai>), <au> und <eu> *geschrieben* werden. Aber wie werden sie *ausgesprochen*? Da scheiden sich die Geister. Die Qualität des ersten Vokals ist ziemlich klar: [a] für die ersten beiden Diphthonge, und [ɔ] (nicht [o]!) für den dritten. Den zweiten Teil des Diphthongs transkribiert jeder anders. Ich führe hier ein paar Möglichkeiten an:

|                          |                  |
|--------------------------|------------------|
| <ei>, <ai> (z.B. <ein>): | [aɪ], [ai], [aj] |
| <au> (z.B. <Haus>):      | [aʊ], [au]       |
| <eu> (z.B. <euch>):      | [ɔɪ], [ɔi], [ɔj] |

### Transkriptionsbeispiele:

|           |                                      |
|-----------|--------------------------------------|
| <Tat>     | [t <sup>h</sup> a:t <sup>(h)</sup> ] |
| <sieht>   | [ʰzi:t <sup>(h)</sup> ]              |
| <wo>      | [ʰʷo:]                               |
| <müsst>   | [mʏst <sup>(h)</sup> ]               |
| <Helm>    | [hɛlm]                               |
| <könnt>   | [k <sup>h</sup> œnt <sup>(h)</sup> ] |
| <Zimmer>  | [tsɪmɐ]                              |
| <Sonne>   | [zɔnə]                               |
| <böse>    | [bø:zə]                              |
| <Höhe>    | [hø:ə]                               |
| <welcher> | [vɛlçɐ]                              |
| <Honig>   | [ho:nɪç], [ho:nɪk] (süddt.)          |
| <Tochter> | [t <sup>h</sup> ɔxtə]                |
| <Töchter> | [t <sup>h</sup> œçtə]                |

|            |                                            |
|------------|--------------------------------------------|
| <Paar>     | [p <sup>h</sup> a:ɐ]                       |
| <Bar>      | [b̥a:ɐ]                                    |
| <mehrere>  | [ˈme:RƏRə], [ˈme:ʁƏʁə]                     |
| <mehr>     | [ˈme:ɐ]                                    |
| <Dächer>   | [ˈdɛçɐ]                                    |
| <sterben>  | [ˈʃtɛʁbən], [ˈʃtɛʁbm̩]                     |
| <König>    | [ˈk <sup>h</sup> ø:nɪç]                    |
| <Sänger>   | [ˈzɛŋɐ]                                    |
| <Tier>     | [ˈt <sup>h</sup> i:ɐ]                      |
| <Söhne>    | [ˈzø:nə]                                   |
| <sinken>   | [ˈzɪŋkən], [ˈzɪŋkŋ]                        |
| <Hase>     | [ˈha:zə]                                   |
| <Wasser>   | [ˈʏasɐ]                                    |
| <Strich>   | [ˈʃtʁɪç], [ˈʃtʁɪç]                         |
| <Wiese>    | [ˈʏi:zə]                                   |
| <Wissen>   | [ˈʏɪsən], [ˈʏɪsŋ]                          |
| <Ding>     | [ˈdɪŋ]                                     |
| <verehren> | [ˈfɛʁʔe: Rən], [ˈfɛʁʔe: ʁən], [ˈfɛʁʔe: ʁŋ] |
| <erste>    | [ˈʔɛʁstə]                                  |
| <Wind>     | [ˈʏɪnt <sup>(h)</sup> ]                    |
| <Häuser>   | [ˈhɔɪzɐ]                                   |
| <Häuschen> | [ˈhɔɪʃçən]                                 |
| <Arbeit>   | [ˈʔaɐbaɪt <sup>(h)</sup> ]                 |
| <Tuch>     | [ˈt <sup>h</sup> u:x]                      |
| <sehr>     | [ˈze:ɐ]                                    |
| <brav>     | [ˈb̥ʁa:f], [ˈb̥ʁa:f]                       |
| <weitere>  | [ˈʏaɪtəʁə], [ˈʏaɪtəRə]                     |
| <süchtig>  | [ˈzʏçtɪç]                                  |
| <Möhre>    | [ˈmø:ʁə], [ˈmø: Rə]                        |
| <doch>     | [ˈdɔx]                                     |
| <Hunger>   | [ˈhʊŋɐ]                                    |
| <zum>      | [ˈtsum]                                    |

### Bibliographie:

- Canepari, Luciano (1979). *Introduzione alla fonetica*. Einaudi, Torino.  
 Di Meola, Claudio (2003). *La linguistica tedesca. Un' introduzione con esercizi e bibliografia ragionata*. Bulzoni, Roma.  
 DUDEN, *Aussprachewörterbuch: Wörterbuch der deutschen Standardaussprache*. 2004. Bearbeitet von Max Mangold in Zusammenarbeit mit der Dudenredaktion. (Der Duden 6). Dudenverlag, Mannheim/Wien/Zürich.

- Eisenberg, Peter (1998). *Grundriß der deutschen Grammatik. Band 1: Das Wort*. Verlag J.B. Metzler, Stuttgart, Weimar.
- Fromkin, Victoria (2000) (Hrsg.), *Linguistics. An Introduction to Linguistic Theory*. Blackwell, Oxford.
- Goldsmith, John A. (1995) (Hrsg.). *The Handbook of Phonological Theory*. Blackwell, Oxford
- Graffi, Giorgio & Sergio Scalise (2002). *Le lingue e il linguaggio. Introduzione alla linguistica*. Bologna, Il Mulino.
- Großes Wörterbuch der deutschen Aussprache* (1982). Herausgegeben von dem Kollektiv Eva-Maria Krech, Eduard Kurka, Helmut Stelzig, Eberhard Stock, Ursula Stötzer und Rudi Teske. VEB Bibliographisches Institut, Leipzig.
- Gussenhoven, Carlos & Haike Jacobs (1998). *Understanding Phonology*. Arnold, London.
- Hall, Tracy Alan (1992). *Syllable structure and syllable-related processes in German* (= Linguistische Arbeiten 276). Max Niemeyer, Tübingen.
- Hall, Tracy Alan (2000). *Phonologie. Eine Einführung*. Berlin/New York, de Gruyter.
- Jakobson, Roman (1962). *Selected writings I: phonological studies*. Mouton, The Hague.
- Kenstowicz, Michael (1994). *Phonology in Generative Grammar*. Blackwell, Oxford.
- Kohler, K.J. (1977). *Einführung in die Phonetik des Deutschen*. Schmidt, Berlin.
- Ladefoged, Peter (1993). *A Course In Phonetics*. Harcourt Brace College Publishers, Fort Worth et al.
- Nespor, Marina (1999). *Fonologia*. Bologna, Il Mulino.
- Pinker, Steven (1994). *The Language Instinct*. William Morrow and Company, New York. [Dt. Übersetzung (1996): *Der Sprachinstinkt*. Kindler, München.]
- Prince, Alan & Paul Smolensky (1993). *Optimality Theory: Constraint Interaction in Generative Grammar*. Ms. Rutgers University, New Brunswick, and University of Colorado, Boulder. Veröffentlicht (2004) als *Optimality Theory. Constraint Interaction in Generative Grammar*, Blackwell, Oxford.
- Ramers, Karl-Heinz & Heinz Vater (1995). *Einführung in die Phonologie*. (Klage 16), Gabel-Verlag, Hürth.
- Ramers, Karl-Heinz (1998). *Einführung in die Phonologie*. W. Fink Verlag, München.
- Spencer, Andrew (1996). *Phonology*. Blackwell, Oxford.
- Trask, R. L. (1996). *A Dictionary of Phonetics and Phonology*. Routledge, London and New York.
- Vater, Heinz (1996). *Einführung in die Sprachwissenschaft*. 2. Auflage, UTB, Fink-Verlag, München.
- Wiese, Richard. 1988. *Silbische und Lexikalische Phonologie: Studien zum Chinesischen und Deutschen* (= Linguistische Arbeiten 211). Niemeyer, Tübingen.
- Wiese, Richard (1996). *The Phonology of German*. Oxford University Press, Oxford.