

Das Phonem

Einleitung

In unserer bisherigen Behandlung von Sprachlauten spielen ausschließlich Fragen nach deren Produktion und der darauf basierenden Klassifikation eine Rolle. Wir verstehen Sprachlaute dabei als relativ konkrete und präzise zu beschreibende, physische Instanzen. Damit liegt die bisherige Betrachtungsweise klar im Gebiet der Phonetik, die sich im Falle der artikulatorischen Phonetik mit dem menschlichen Lautbildungspotential und den physiologischen Mechanismen der Lauterzeugung beschäftigt und auf diese Weise anstrebt, ein umfassendes Bild aller in den Sprachen der Welt vertretenen Laute zu erstellen – sozusagen eine Übersicht über das globale Sprachlautinventar. Fragen danach, wie die Einzelsprachen dieses globale Inventar individuell ausschöpfen, haben wir dabei nicht weiter berücksichtigt, und um genau diese wird es in den nächsten Abschnitten gehen. Damit verlassen wir den Bereich der Phonetik und betreten das Gebiet der Phonologie.¹ Dabei werden wir eine Reihe von Konzepten einführen, die für die phonologische Analyse von Sprachdaten eine zentrale Rolle spielen.

Wir beginnen dabei mit einem Phänomen, dass uns allen bekannt ist – dem fremdsprachlichen Akzent:



Abbildung 1: Deutscher Akzent in den Katzenjammer-Kids

Der hier dargestellte Cartoonausschnitt stammt aus einem der ältesten modernen Comicstrips, den *Katzenjammer Kids*, die ab 1897 regelmäßig in einer Beilage des *New York Journals* erschienen sind.

Für uns interessant ist der in diesem Ausschnitt zutage tretende deutsche Akzent, der hier orthographisch wiedergegeben ist in Wörtern wie *ve* und *vot* (statt *we* und *what*), *chust* (statt *just*), *dey're* und *dey* (statt *they're* und *they*). Offensichtlich werden Sprachlaute, die nicht im muttersprachlichen Lautinventar vorhanden sind (wie die englischen [w] oder [ð]), durch ähnliche Laute des eigenen Lautinventars ersetzt – ein wesentlicher Faktor eines fremdsprachlichen Akzents. Im Kontrast Deutsch-Englisch sind u.a. folgende Ersetzung üblich:

[ð] als [d] oder [z], [w] als [v], [dʒ] als [tʃ], [θ] als [t] oder [s], [ɑ] als [a], [ɜ] als [œ], [eɪ] als [e:] usw.

Sprachen schöpfen das universale Lautinventar unterschiedlich aus: [θ], [w], [ɑ] usw. kommen im Englischen vor, aber nicht im Deutschen. Umgekehrt kommen [ç], [x] oder [ø] im Deutschen vor, nicht aber im Englischen.

¹ Dazu Otto Jespersen (1860-1943), einer der führenden Phonetiker und Grammatiker des 20. Jahrhunderts: *It would, perhaps, be advisable to restrict the word "phonetics" to universal or general phonetics and to use the word "phonology" of the phenomena peculiar to a particular language (e.g. "English Phonology")*. (Jespersen 1924:35). Kenneth Pike (Pike 1971¹²: 57) drückt es salopper so aus: *Phonetics gathers raw material. Phonemics [=phonology] cooks it*.

Betrachten wir dazu einige weitere Beispiele:²



Abbildung 2: Eric Clapton und Harry Potter

Auf diesen Bildern tritt ein bekanntes und häufig unangemessen belächeltes Phänomen auf: die Schwierigkeiten, bestimmte Liquide (r-Laut und l-Laut) auseinanderzuhalten, die bei Sprechern einiger asiatischer Sprachen – im vorliegenden Fall Koreanisch und Japanisch – zu beobachten ist. Können wir die Argumentation aus dem Kontext »deutscher Akzent« übertragen und aus diesem Beispiel schließen, dass in den fraglichen Sprachen entweder kein r- oder kein l-Laut vertreten ist? Die Antwort lautet »nein«, da sowohl im Japanischen als auch im Koreanischen zumindest vergleichbare Laute vertreten sind.

Um dieses Phänomen gründlich zu beschreiben, müssen wir uns das Konzept »Sprachlaut« genauer ansehen und präzisieren. Erst dann können wir erklären, was ursächlich für die r-l-Problematik anzusehen ist.

Phone und Phoneme

Im phonologischen System einer Einzelsprache gibt es zahlreiche Laute, die sich ähneln, deren Artikulation sich aber voneinander unterscheidet. Trotzdem werden diese Laute von den Sprechern der Sprache als ein- und derselbe Laut wahrgenommen: Sprecher einer Einzelsprache ignorieren in vielen Fällen substantielle Eigenschaften von Sprachlauten und betrachten verschiedene Laute als den gleichen Laut. Sehen wir uns dazu ein Beispiel aus dem Deutschen an.

1. Skat, kahl, Oktober, Kiel

Diese Wörter haben eine Gemeinsamkeit: in jedem dieser Wörter kommt der gleiche Laut vor, nämlich der k-Laut. Wenn Sie diese Wörter bewusst aussprechen und dabei genau auf die Produktion achten, können Sie allerdings Unterschiede feststellen:

- Der k-Laut in *Kiel* wird im Unterschied zu den anderen k-Lauten weiter vorne, also eher am harten, nicht am weichen Gaumen gebildet. Dieses Merkmal, fachsprachlich *fronting* genannt, ist durch ein kleines Pluszeichen unter dem frontierten Laut angezeigt: [k̟].
- In *Oktober* wird der Verschluss des k-Lautes im Unterschied zu den anderen k-Lauten nicht gelöst, sondern erfolgt erst bei der Produktion des darauffolgenden [t]. Das diakritische Zeichen dafür ist ein kleines Häkchen: [k̟̚]
- In *Kiel* und *kahl* wird der k-Laut aspiriert, nicht aber in *Skat* und *Oktober*.

Zu diesem Punkt gibt es ein kleines Video namens *Oktober* auf der Webseite.

Mit den diakritischen Zeichen für Aspiration (^h) und keine Verschlusslösung (engl. *unreleased*) ([̚]) sind die Unterschiede zu erfassen und die enge phonetische Transkription liefert uns die folgenden Formen und die folgenden Versionen des k-Lautes:

2. [ska:t^h] — [k^ha:l] — [ʔɔk̟̚'t^ho:be] — [k̟̚^hi:l]
3. [k] — [k^h] — [k̟̚] — [k̟̚^h]

Tatsächlich und physisch betrachtet haben wir hier also vier etwas unterschiedliche Laute. In unserer Wahrnehmung aber handelt es sich um denselben Laut. Betrachten wir auf dieser Grundlage nun die folgenden Phantasiewörter:

² www.english.com

4. skup, kait, aktal, kiebe

Wenn Sie diese Wörter wieder laut aussprechen, stellen Sie fest, dass Sie die k-Laute jeweils genau so artikulieren, wie sie in (3) transkribiert sind – und das tun Sie, ohne die Wörter überhaupt zu kennen und ohne darüber nachzudenken.

Offensichtlich haben Sie also die folgenden phonologischen Regeln des k-Lauts im Deutschen internalisiert:

- er wird vor einem Vorderzungenvokal frontiert, spricht am harten Gaumen produziert (wie in *Kiel*, *kiebe*),
- er wird vor einem weiteren Konsonanten ohne Verschlusslösung produziert (wie in *Oktober*, *aktal*),
- er wird am Beginn einer betonten Silbe aspiriert (wie in *Kiel*, *kiebe*, *kahl* und *kait*).

Diese Regeln hat Ihnen niemand beigebracht, Sie haben sie beim Erwerb ihrer Muttersprache automatisch verinnerlicht.³ Wir können hier erkennen, dass die Frage danach, welchen Laut wir verwenden, offensichtlich stark davon abhängt, was dem Laut unmittelbar folgt oder vorausgeht.

Eine zentrale Aufgabe der Phonologie besteht nun darin, solche phonologischen Regeln einer Einzelsprache zu erfassen und präzise zu beschreiben. Zu diesem Zweck führen wir eines der zentralen Konzepte der Phonologie ein: das Phonem.

Um terminologisch zwischen einem Sprachlaut als tatsächlicher, physischer Einheit und einem Sprachlaut als eher abstrakter, wahrgenommener Einheit zu differenzieren, unterscheiden wir zwischen Phonem auf der einen Seite und Phonemen auf der anderen Seite. Während das phonetische Alphabet universal ist, variiert das phonemische Alphabet von Sprache zu Sprache. Was die Notation angeht, werden Phone in eckigen Klammern geführt ([p], [ɔ], [ç] usw.), Phoneme zwischen Schrägstrichen (/p/, /ɔ/, /ç/ usw.).

Um auf die Einleitung zurückzukommen: die Laute [k], [kʰ], [kʷ] und [k̟] sind vier verschiedene Phone, die – im Deutschen – jeweils dasselbe Phonem /k/ ("den k-Laut") repräsentieren.

Die Differenzierung zwischen Phon und Phonem geht u.a. zurück auf Daniel Jones (1881-1967), einem der bekanntesten und einflussreichsten Phonetiker des 20. Jahrhunderts, dem wir auch das System der Kardinalvokale zu verdanken haben. Die erste Verwendung des Terminus »Phonem«⁴ im heutigen Sinn findet sich bei Jones in einem Artikel zum Setswana, den er 1917 vor der britischen *Philological Society* in London vortrug:

The Sechuana language appears to contain twenty-eight phonemes, i.e. twenty-eight sounds or small families of sounds which are capable of distinguishing one word from another. (Jones 1919)

In diesem Abschnitt sind für uns zwei Dinge relevant: zum einen erhalten wir einen Hinweis darauf, dass ein Phonem einer Sprache X eine *small family of sounds*, also mehr als nur einen einzigen Laut umfassen kann. Das haben wir im Falle des /k/ im Deutschen ja gerade an einem Beispiel gesehen. Zum anderen wird die Funktion des Phonems benannt: es dient dazu, ein Wort von einem anderen zu unterscheiden. Dieser Punkt findet sich auch in zahlreichen modernen Beschreibungen, in denen das Phonem als »kleinste bedeutungsunterscheidende Einheit der Sprache« beschrieben wird. Diese Funktion ist überhaupt essentiell dafür, zu ermitteln, ob bestimmte Laute in einer Sprache ein Phonem repräsentieren oder nicht. Ein Verfahren, das zu diesem Zweck eingesetzt wird, ist die sogenannte Minimalpaaranalyse, auf die wir gleich noch etwas eingehen.

Einige Jahre nach seinem Artikel über das Setswana liefert Jones eine etwas detailliertere Beschreibung für das Phonem (1932, reprint 1957:49):

³ Eine der großen Diskussionen in der theoretischen Linguistik dreht sich um die Frage, welche Menge grammatischer Regeln von Kindern im Erstspracherwerb tatsächlich 'erworben' werden. Gibt es eine Teilmenge von Regeln, für die die Spezies Mensch bereits bei der Geburt eine kognitive Voraussetzung mitbringt, sozusagen ganz grundlegende, universelle Strukturprinzipien, die im menschlichen Gehirn verankert sind? Auf diese Diskussion gehen wir in unserem Modul nicht näher ein.

⁴ Die erste Verwendung des Begriffs findet sich im 19. Jh. bei Baudouin de Courtenay, allerdings in einer etwas anderen Bedeutung.

In describing the sound-system of any language it is necessary to distinguish between speech-sounds and what are called phonemes. A speech-sound is a sound of definite organic formation and definite quality which is incapable of variation. A phoneme may be described roughly as a family of sounds consisting of an important sound of the language (i.e. the most frequently used member of that family) together with other related sounds which 'take its place' in particular sound-sequences or under particular conditions of length or stress.

In dieser Definition wird ein entscheidender Aspekt für die Ermittlung der Phoneme einer Sprache benannt, nämlich der sprachliche Kontext, die Umgebung, in dem die fraglichen Laute auftreten. Auch auf die Funktion des Phonems geht Jones ausführlicher ein (1932, reprint 1957:51 leicht modifiziert):

Phonemes are capable of distinguishing one word of a language from other words of the same language. There is an English word *sin* (sɪn) and another English word *sing* (sɪŋ). [...] The existence of such words is a proof that the n and ŋ sounds belong to separate phonemes in English.

Wir halten fest: ein Phonem ist eine relativ abstrakte Einheit, die in spezifischen sprachlichen Kontexten (*particular sound-sequences*) durch ein oder mehrere verwandte Phone realisiert wird. Die Funktion eines Phonems besteht darin, Bedeutungen zu unterscheiden. Diese Aussage zieht zwei Fragen nach sich, die es im Einzelnen zu klären gilt:

- was ist mit »verwandt« gemeint?
- was ist mit »spezifischen sprachlichen Umgebungen« gemeint?

Diese Fragen müssen wir klären, damit wir zu einer präzisen (und operationalisierbaren) Definition des Konzepts »Phonem« kommen.

Phonetische Ähnlichkeit

Zu diesem Punkt, also der Verwandtschaft zweier oder mehrerer Laute, können wir knapp sagen, dass diese Verwandtschaft, sprich die phonetische Ähnlichkeit zweier Laute umso größer ist, je mehr artikulatorische Merkmale sie teilen.

Danach wären sich beispielsweise [p] und [b] relativ ähnlich, da sie den Artikulationsort als auch die Artikulationsart teilen und sich nur im Merkmal der Stimmhaftigkeit unterscheiden. [p] und [ʒ] wären sich relativ unähnlich, da sie sich sowohl bei Artikulationsort als auch -art und in der Stimmhaftigkeit unterscheiden – hier steht ein stimmloser bilabialer Plosiv einem stimmhaften postalveolaren Frikativ gegenüber. [p] und [b] sind sich also phonetisch ähnlicher als [p] und [ʒ].

Was die Konsonanten betrifft, sind sich in aller Regel Laute ähnlicher, deren Artikulationsstellen näher benachbart sind als solche, die der Artikulationsstellen weit auseinander liegen.

Sprachliche Umgebung

Was die Frage nach den spezifischen sprachlichen Umgebungen betrifft, so geht es hier um die Distribution, also um die Verteilung von Lauten in einem Sprachdatensatz. Das Konzept »Distribution« ist nicht nur für die Phonologie, sondern ganz allgemein für die strukturelle Grammatik, also auch für die Morphologie, Syntax und Semantik von größter Bedeutung.

Zum Thema »Distribution« finden Sie ein Video auf der Webseite. Kenntnis über verschiedene Typen von Distribution wird in den kommenden Sitzungen vorausgesetzt.

Wir beginnen zunächst mit zwei Definitionen:

Distribution

Die Distribution eines sprachlichen Segmentes X ist die Menge der Umgebungen, in der X in einem Datensatz vorkommt.

Umgebung

Die Umgebung eines sprachlichen Segmentes X ist definiert über diejenigen Segmente, die X unmittelbar vorausgehen und X unmittelbar folgen.

Im nächsten Beispiel sehen wir vier Ketten von Phonen, die zusammen den Mini-Datensatz 1 bilden:

5. (1) bast (2) lan (3) alt (4) 'tu:ba

In jeder dieser Ketten kommt der Laut [a] vor, und zwar in folgenden Umgebungen: in (1) nach einem [b] und vor einem [s], in (2) nach einem [l] und vor einem [ŋ], in (3) steht das [a] hinter einer Wortgrenze, für die wir die Raute '#' als Symbol einführen, und vor einem [l], in (4) steht das [a] schließlich nach einem [b] und vor der Wortgrenze #. Zusammengenommen ergibt diese Menge von Umgebungen die Distribution von [a] in unserem Minidatensatz. Die Notation dieses Sachverhaltes sieht so aus:

$$[a] / \left\{ \begin{array}{l} [b] __[s] \\ [l] __[ŋ] \\ \# __[l] \\ [b] __ \# \end{array} \right\}$$

Abbildung 3: Distribution von [a] im Mini-Datensatz 1

Der waagerechte Strich in den Umgebungen ist ein Platzhalter für das Segment, das am Anfang des Ausdrucks steht, hier also das [a], der Schrägstrich bedeutet 'kommt vor in Umgebung' und die Schweifklammer zeigt an, dass wir es hier mit einer Menge zu tun haben.

Besonders relevant ist die Untersuchung der Distribution eines einzelnen Segmentes dann, wenn wir sie in Beziehung setzen zur Distribution eines anderen Segments, wenn wir also die Distribution von X vergleichen mit der Distribution von Y. Dabei können wir zwei wichtige Untertypen ausmachen, nämlich die äquivalente Distribution einerseits, die komplementäre Distribution andererseits. Auch diese beiden Konzepte werden zunächst definiert und dann an einem Beispiel illustriert.

Äquivalente Distribution

Zwei Segmente X und Y sind äquivalent verteilt, wenn überall dort, wo X vorkommen kann, auch Y vorkommen kann und umgekehrt.

In Minidatensatz 2 sehen wir wieder Ketten von Phonem, achten Sie hier auf die Distribution von [m] einerseits und [p] andererseits:

6. 'mapə amt 'o:ma 'hæm 'papə apt 'o:pa hæp

Wir sehen, dass die Umgebungen für [m] und [p] identisch sind, wir haben es mit äquivalenter Distribution zu tun:

$$\left\{ \begin{array}{l} [m] \\ [p] \end{array} \right\} / \left\{ \begin{array}{l} \# __[a] \\ [a] __[t] \\ [o:] __[a] \\ [aɪ] __ \# \end{array} \right\}$$

Abbildung 4: äquivalente Distribution von [m] und [p] im Mini-Datensatz 2

Wenn Sie sich die Wörter in (6) ansehen, stellen Sie fest, dass wir hier Wortpaare bilden können, die sich in nur einem Segment in identischer Position voneinander unterscheiden:

7. Mappe–Pappe Amt–Abt Oma–Opa Heim–Hype

Der Austausch der beiden Laute [m] und [p] in diesen Umgebungen führt dazu, dass wir es hier mit jeweils zwei verschiedenen Wörtern zu tun haben: Mappe bedeutet etwas anderes als Pappe, Amt etwas anderes als Abt usw. Hier sehen wir die Kernfunktion sprachlicher Laute: sie dienen dazu, den Kontrast zwischen sprachlichen Zeichen auszudrücken. In den Wortpaaren in (7) kontrastieren also die Laute [m] und [p].

Für derartige Paare von Wörtern, die sich nur in einem Segment in identischer Umgebung unterscheiden, gibt es einen eigenen Fachbegriff:

Minimalpaar

Ein Minimalpaar ist ein Paar von Wörtern, das sich in genau einem Laut in identischer Position unterscheidet. Die Laute weisen äquivalente Distribution auf und kontrastieren.

Zur Minimalpaaranalyse gibt es ein Video auf der Webseite. Kenntnis über Minimalpaare wird in den nächsten Sitzungen vorausgesetzt.

Danach haben wir es in den folgenden Beispielen jeweils mit Minimalpaaren zu tun, in denen die Laute [ç]-[g], [ʁ]-[ʁ̥], [aʊ]-[i:], [m]-[x] und [d]-[l] kontrastieren:

8. 'tsaɪçən—'tsaɪgən 'ʁɪnə—'zɪnə faʊl—fi:l bəʊm—bəʊx gə'dɛŋkə—gə'lɛŋkə

Minimalpaare liefern einen wichtigen Hinweis auf die Frage, ob zwei Laute X und Y verschiedene oder dasselbe Phonem repräsentieren. Wenn wir für zwei Laute X und Y ein Minimalpaar finden, haben wir einen Beleg dafür, dass sie verschiedene Phoneme repräsentieren. Wenn Sie zum letzten Exzerpt von Daniel Jones auf Seite (4) zurückblättern, sehen Sie, dass auch Jones Minimalpaare einsetzt, konkret das Minimalpaar sɪn—sɪŋ.

Sehen wir uns nun das nächste Beispiel an:

9. 'rɔ:zə—'ʁɔ:zə brɑ:t—bʁɑ:t

In (9) sehen wir zwei weitere Paare, die sich in nur einem Laut in identischer Position unterscheiden. In diesen Paaren sind die Laute [r] und [ʁ] äquivalent verteilt. Sind die Paare in (9) also Minimalpaare? Die Antwort lautet 'nein', denn der Austausch von [r] und [ʁ] führt nicht zu einem jeweils anderen Zeichen. Sowohl 'rɔ:zə als auch 'ʁɔ:zə haben dieselbe Bedeutung, beide Ketten stehen für das Wort *Rose*. Also kontrastieren [r] und [ʁ] nicht. Das gleiche sehen wir auch Paaren wie brɑ:t—bʁɑ:t, e:ɐ'trɑ:ŋ—e:ɐ'tʁɑ:ŋ usw. Um diesen Sachverhalt verallgemeinert zu beschreiben, dient der Terminus »freie Variation«:

Freie Variation

Zwei Segmente X und Y stehen in freier Variation, wenn X und Y äquivalent verteilt sind, ohne zu kontrastieren.

Laute wie [r] und [ʁ] im Deutschen, also phonetisch ähnliche Laute in freier Variation, repräsentieren stets dasselbe Phonem.

Hier können wir bereits eine grobe Faustregel formulieren: zwei Laute, die phonetisch ähnlich und äquivalent verteilt sind und kontrastieren, repräsentieren verschiedene Phoneme. Zwei Laute, die phonetisch ähnlich sind und in freier Variation auftreten, repräsentieren dasselbe Phonem.

Zum Ende dieses Abschnittes betrachten wir einen weiteren wichtig Untertyp, die komplementäre Distribution.

Komplementäre Distribution

Zwei Segmente X und Y sind komplementär verteilt, wenn überall dort, wo X vorkommen kann, Y nicht vorkommen kann und umgekehrt.

Hierfür betrachten wir Minidatensatz 3 mit folgenden Lautketten, achten Sie auf die Umgebungen von [h] und [ŋ]:

10. hɛs hals ho:f 'hɔ:tə sɔŋ dʊŋ gɔŋ gə'ʁɪŋ

Die Umgebungen für [h] und [ŋ] sehen so aus:

$$[h] / \left\{ \begin{array}{l} \#_ [aɪ] \\ \#_ [l] \\ \#_ [o:] \\ \#_ [ɔɪ] \end{array} \right\} \quad [ŋ] / \left\{ \begin{array}{l} [a]_\# \\ [ʊ]_\# \\ [ɔ]_\# \\ [ɪ]_\# \end{array} \right\}$$

Abbildung 5: komplementäre Distributionen von [h] und [ŋ] im Mini-Datensatz 3

Hier springt sofort ins Auge, dass das [h] in einer Position vorkommt, in der das [ŋ] nicht vorkommt, nämlich in der Umgebung #_X (das X ist ein Variable für verschiedene Laute), das [ŋ] hingegen steht in der Umgebung X_#, in der wiederum das [h] nicht auftreten kann. Diese Beobachtung können für die ganze Sprache generalisieren: im Deutschen kann das [h] am Beginn eines Wortes vorkommen, das [ŋ] nicht, dagegen kann das [ŋ] am Ende eines Wortes auftreten, das [h] nicht. Damit sind die beide Laute [h] und [ŋ] im Deutschen komplementär verteilt.

Mit diesen Erkenntnissen im Gepäck haben wir alles beisammen, was wir für die Definition des Phonems benötigen:

Phonemdefinition

Ein Phonem ist eine Menge von phonetisch ähnlichen Phonemen (diese Menge kann auch aus nur einem Element bestehen), die entweder in freier Variation auftreten oder komplementär verteilt sind. Die Phone, die ein Phonem repräsentieren, nennt man Allophone dieses Phonems.

Diese Definition liefert eine präzise Erklärung für das Phonemkonzept. Außerdem ist sie gut auf die Phonembeschreibung von Daniel Jones (s.o.) anwendbar:

Wo Jones von Lauten als *speech-sound of definite organic formation and definite quality which is incapable of variation* spricht, sprechen wir vom Phon. Wo Jones sagt, dass es beim Phonem um eine Menge von Lauten geht, die *related* sind, sprechen wir von phonetischer Ähnlichkeit. Wo Jones sagt, dass die Allophone den Platz des Phonems in *particular sound-sequences or under particular conditions of length or stress* einnehmen, sprechen wir von Distribution, Umgebung, freier Variation und komplementärer Verteilung.⁵

Aus dieser Definition ist beispielsweise ableitbar, dass [ʀ], [r] und [ʁ] (siehe Beispiel (9)) dasselbe Phonem im Deutschen repräsentieren: sie sind phonetisch ähnlich und stehen in freier Variation. Ebenso ist ableitbar, dass [h] und [ŋ] nicht dasselbe Phonem im Deutschen repräsentieren: die Laute sind zwar komplementär verteilt, aber nicht phonetisch ähnlich.

Analyse und phonologische Regel

In diesem Teil geht es darum, auf Basis des bisher Gesagten ein paar konkrete Analysen durchzuführen und etwas über phonologische Regeln und deren Notation zu erfahren.

Phoneme im Deutschen

Wir haben bereits mit zwei Phonemen des Deutschen Bekanntschaft geschlossen, dem /k/ mit seinen Allophonen {[k], [kʰ], [kʷ], [cʰ]} und dem /r/ mit seinen Allophonen {[ʀ], [ʁ], [r]}. Beim /r/ kommt noch ein weiterer Laut hinzu, nämlich das Zungenspitzen-R, also der alveolare Vibrant [r]:

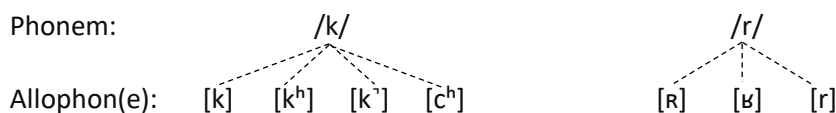


Abbildung 6: /k/ und /r/ im Deutschen

Die Allophone des /k/ sind komplementär verteilt, die des /r/ stehen in freier Variation.

Als nächstes betrachten wir die Laute [ç] und [x] im Standarddeutschen etwas genauer. Um herauszufinden, ob diese phonetisch ähnlichen Laute dasselbe oder verschiedene Phoneme repräsentieren, wenden wir die folgenden Schritte an:

- wir schauen zunächst, ob wir Minimalpaare mit [ç] und [x] finden, also Paare, in denen die Laute äquivalent verteilt sind und kontrastieren.
- wir analysieren die genaue Distribution von [ç] und [x]. Wir versuchen dabei herauszufinden, ob die Laute komplementär verteilt sind oder nicht. Wenn ja, handelt es sich um Allophone desselben Phonems.

Was den ersten Punkt angeht, lautet die Antwort 'nein'.

Was den zweiten Punkt angeht, betrachten wir den folgenden Datensatz, in dem es um den Laut [ç] geht:

11. 'kœçɪn 'viçtə 'zi:çən 'bɛçnə 'by:çə miç

Hier sehen wir die folgende Distribution:

⁵ Diese Definition geht insofern über die von Daniel Jones hinaus, als dessen Versionen die freie Variation nicht direkt umfassen. Das heißt nicht, dass er diesen Punkt nicht beachtet hätte – er spricht dann aber von *Variphones* (z.B. in Jones 1962²:205, siehe weiter unten).

$$12. [\zeta] / \left\{ \begin{array}{l} [\text{œ}] _ [\text{i}] \\ [\text{i}] _ [\text{t}] \\ [\text{i:}] _ [\text{ə}] \\ [\text{ɛ}] _ [\text{n}] \\ [\text{y:}] _ [\text{e}] \\ [\text{ɪ}] _ \# \end{array} \right\}$$

An dieser Stelle erfolgt ein wichtiger Schritt: wir wollen herausfinden, ob in der Distribution bestimmte Generalisierungen möglich sind, ob wir also eine Systematik in der Umgebung ausmachen können: Haben die Laute, die jeweils vor bzw. hinter $[\zeta]$ auftreten, eine Gemeinsamkeit?

Bei den Lauten, die auf $[\zeta]$ folgen, können wir keine solche Systematik ausmachen: hier liegt eine sehr heterogene Gruppe vor von Konsonanten wie bei *Wichte* oder *rechne*, Vokalen wie bei *Köchin*, *siechen* oder *Bücher* oder auch der Wortgrenze wie bei *mich*: $_ \{[\text{t}, \text{n}, \text{ɪ}, \text{ə}, \text{e}, \#]\}$. Wir konzentrieren uns folglich auf das, was dem $[\zeta]$ vorausgeht.

Hier sehen wir folgendes: in allen Fällen geht $[\zeta]$ ein Vokal voraus. Die Klasse 'Vokal' kürzen wir mit 'V' ab, damit gilt $[\zeta] / V _$. Diese Vokale können wir noch genauer spezifizieren, denn sie bilden keine arbiträre Menge: wir haben es in jeder einzelnen Umgebung mit einem Vorderzungenvokal zu tun, also um Vokale mit dem Merkmal [vorne]. Damit können wir das Entscheidende für die Distribution von $[\zeta]$ wie folgt notieren, zu lesen als ' $[\zeta]$ folgt auf einen Vorderzungenvokal':

$$13. [\zeta] / V_{[\text{vorne}]} _$$

Kommen wir als nächstes zum $[x]$ in den folgenden Beispielen:

$$14. \text{bu:x} \quad \text{dɔ:x} \quad \text{'bo:xʊm} \quad \text{zʊxt}$$

Hier sehen wir diese Distribution:

$$15. [x] / \left\{ \begin{array}{l} [\text{u:}] _ \# \\ [\text{ɔ}] _ \# \\ [\text{o:}] _ [\text{ʊ}] \\ [\text{ʊ}] _ [\text{t}] \end{array} \right\}$$

Hier haben wir eine analoge Situation zum $[\zeta]$, denn auch hier ist nur der vordere Teil der Umgebung relevant (der hintere Teil ist wieder heterogen, wir finden hier mal nichts wie in *Buch* und *doch*, mal einen Vokal wie in *Bochum* und mal einen Konsonanten wie in *Sucht*). Vor dem $[x]$ steht aber immer ein Vokal. Wenn wir uns diese Vokale genau ansehen, stellen wir auch hier eine Gemeinsamkeit fest: es sind allesamt Hinterzungenvokale, sprich Vokale mit dem Merkmal [hinten]. Es gilt also: $[x]$ folgt auf Hinterzungenvokale:

$$16. [x] / V_{[\text{hinten}]} _$$

Zusammengenommen ergibt sich folgendes Bild: Hinter einem Vorderzungenvokal kann das $[\zeta]$ auftreten, aber nicht das $[x]$. Hinter einem Hinterzungenvokal kann das $[x]$ auftreten, aber nicht das $[\zeta]$.

Damit gilt folgendes: die beiden Laute $[\zeta]$ und $[x]$ sind phonetisch ähnlich und im Deutschen komplementär verteilt. Damit repräsentieren sie dasselbe Phonem bzw. sind Allophone desselben Phonems $/\zeta/$:

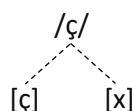


Abbildung 7: $/\zeta/$ im Deutschen

Wenn ein Phonem über mehr als nur ein Allophon verfügt, stellt sich die Frage, welches dieser Allophone für die Bezeichnung des Phonems gewählt, dem Phonem also zugrundeliegt. Wir wählen hier das $[\zeta]$, manche Autoren verwenden das $[x]$. Es gibt Argumente für beide Optionen. Der Grund für unsere Wahl liegt u.a. darin, dass die Ableitung von Phonem aus der Phonemen sehr viel ökonomischer ist, wenn dasjenige Allophon, das die größte Distribution aufweist, als Phonem zugrunde gelegt wird. Dazu später noch mehr. In unserem Beispiel hat das $[\zeta]$ die größte Distribution, da dieses, anders als das $[x]$, auch hinter Approximanten wie $[\text{l}]$ oder $[\text{ʁ}]$ auftreten kann (z.B. in

Milch oder *Furche*) oder auch hinter einem Nasal (z.B. in *manche*). In diesen Umgebungen kommt [x] nicht vor. Sehen wir uns dazu nochmal die vollständigen Distributionen von [ç] und [x] an:

- [ç] tritt in folgenden Umgebungen auf: hinter einem Vorderzungenvokal, hinter [l], [r] und [n] und hinter den Diphthongen [ai] und [ɔy]. In einigen Lehnwörtern wie ['çi:na], [çɛ'mi:] oder [çi'ni:n] steht [ç] auch am Beginn einer Silbe, dies berücksichtigen wir hier nicht weiter.
- [x] tritt in folgenden Umgebungen auf: hinter einem Zentral- oder Hinterzungenvokal und hinter dem Diphthong [aʊ].

Als Sprecher des Deutschen haben wir diese Realisierungen internalisiert. Um einen Argumentationsschritt aus dem ersten Abschnitt aufzugreifen: wenn Sie die folgenden Phantasiewörter aussprechen, folgen Sie genau den Regelmäßigkeiten der o.a. Distribution:

17. [x] *kucht*: [kʊxt], nicht [kʊçt], *spauche*: ['ʃpaʊxə], nicht [ʃpaʊçə], *knocht*: [knɔ:xt], nicht [knɔ:çt] usw.

18. [ç] *palch*: [palç], nicht [palx], *ziecht*: [tsi:çt], nicht [tsi:xt], *flurche*: ['flʊɐçə], nicht ['flʊɐxə] usw.

Phonologische Regeln

Wie weiter oben bereits gesagt, besteht eine Aufgabe der Phonologie darin, diese eben beobachteten Regelmäßigkeiten transparent zu machen. Zu diesem Zweck führen wir als nächstes phonologische Regeln ein. Was das Phonem /ç/ angeht, betrachten wir zunächst dessen folgende, explizite Regel:

$$\begin{array}{l} /ç/ \rightarrow \left\{ \begin{array}{l} V[\text{hinten}] __ \\ V[\text{zentral}] __ \\ [aʊ] __ \end{array} \right\} \\ \left\{ \begin{array}{l} V[\text{vorne}] __ \\ [ai] __ \\ [ɔy] __ \\ [l] __ \\ [ʁ] __ \\ [n] __ \end{array} \right\} \end{array}$$

Abbildung 8: phonologische Regel für /ç/ im Deutschen, Version A

Eine weitverbreitete Möglichkeit, die Regel in Abbildung 8: phonologische Regel für /ç/ im Deutschen ökonomischer zu gestalten, besteht darin, nur diejenigen Fälle aufzuzählen, in denen das /ç/ als [x] realisiert wird und dann anzugeben, dass /ç/ in allen anderen Fällen als [ç] realisiert wird. Das sieht dann so aus:

$$\begin{array}{l} /ç/ \rightarrow \left\{ \begin{array}{l} V[\text{hinten}] __ \\ V[\text{zentral}] __ \\ [aʊ] __ \end{array} \right\} \\ [ç] / \text{sonst} \end{array}$$

Abbildung 9: phonologische Regel für /ç/ im Deutschen, Version B

Beide dieser Regeln entsprechen Grundmuster für phonologische Regeln:

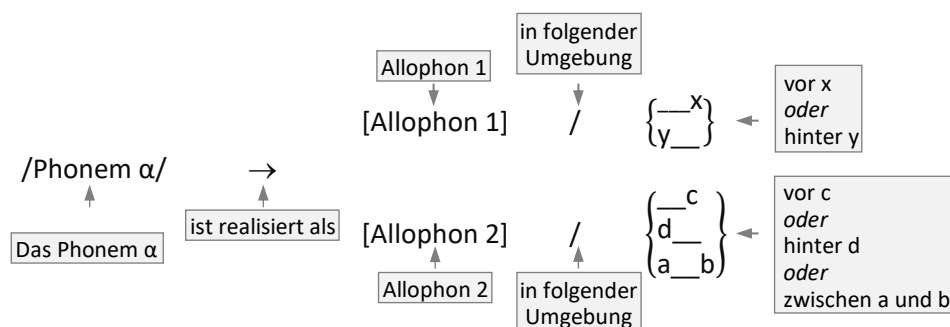


Abbildung 10: Annotiertes Grundmuster für phonologische Regeln

Zu lesen ist diese Regel wie folgt: das Phonem /ç/ wird realisiert ('ausgesprochen') als [x], wenn es hinter einem Hinterzungen- oder Zentralvokal auftritt oder hinter dem Diphthong [aʊ]. In allen anderen Fällen wird es als [ç] realisiert ('ausgesprochen'). Die phonologischen Regeln A und B (s.o.) sind gleichwertig, wenn es darum geht, die stellungsbedingten Varianten des Phonems darzustellen, Regel B ist aber ökonomischer.

Das Phonem als sprachspezifisches Konstrukt

Wenn Sie den Text aufmerksam gelesen haben, ist Ihnen vielleicht aufgefallen, dass immer vom »Phonem /ç/ im Deutschen« oder dem »Phonem /r/ im Deutschen« die Rede gewesen ist. Dies ist ein Hinweis darauf, dass es sich beim Phonem stets um ein sprachspezifisches Konzept handelt. Nachstehend sehen Sie eine Reihe von Beispielen dafür, dass sich das Verhältnis Phon–Phonem im phonologischen System verschiedener Sprachen ganz unterschiedlich gestalten kann.

/k/ (Deutsch, Thai)

Wir beginnen mit dem /k/, das wie gesehen u.a. die Allophone [k] und [kʰ] umfasst.

Wenn wir uns das Lautpaar [k]–[kʰ] in anderen Sprachen ansehen, ist die Sachlage eine andere. Im nächsten Beispiel sehen wir zwei Wortpaare aus dem Thailändischen:

19. กด: [kòtʰ] 'drücken' ขด: [kʰòtʰ] 'Spule'
 กวด: [kù:atʰ] 'verfolgen' ขวด: [kʰù:atʰ] 'Flasche'

Wir haben es bei [kòtʰ]–[kʰòtʰ] und [kù:atʰ]–[kʰù:atʰ] mit Minimalpaaren zu tun: die Laute [k] und [kʰ] sind äquivalent verteilt und kontrastieren. Damit repräsentieren Sie zwei verschiedene Phoneme:

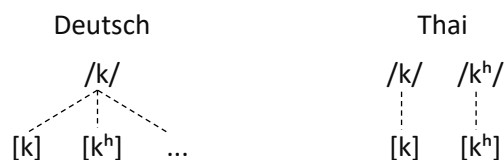


Abbildung 11: /k/ im Deutschen und Thai

/l/ (Deutsch, Englisch, Albanisch)

Betrachten wir als nächstes das /l/. Dieses hat im Deutschen nur ein Allophon. Schauen wir uns nun vier Wörter aus dem Englischen an:

20. leaf [li:f] silence [saɪləns]
 21. mild [maɪld] feel [fi:l]

Wir sehen hier, dass im Englischen zwei l-Laute vorkommen, das sog. clear-l [l] und das sog. dark-l [lʲ].⁶ Bei letzterem handelt es sich um eine velarisierte Form des [l]: bei der Artikulation wird der hintere Zungenrücken in Richtung Velum angehoben. Die Distribution von clear- und dark-l ist im Englischen regelhaft komplementär: das clear-l wird innerhalb einer Silbe vor Vokalen, das dark-l vor Konsonanten und am Wortende verwendet. [l] und [lʲ] sind phonetisch ähnlich und komplementär verteilt, sind im Englischen also Allophone desselben Phonems /l/.

Nehmen wir nun eine weitere Sprache hinzu, das Albanische. Das albanische Lautinventar umfasst, wie das englische auch, sowohl clear- wie auch dark-l, unterscheidet sich aber insofern vom Englischen, als die Laute hier nicht dasselbe, sondern verschiedene Phoneme repräsentieren. Dies kann über die folgenden Minimalpaare gezeigt werden:

22. mal: [mal] 'Berg' mall: [malʲ] 'Waren'
 23. lak: [lak] 'Schleife' llak: [lʲak] 'Spray, Lack'

Hier sind [l] und [lʲ] äquivalent verteilt und kontrastieren, also haben wir es mit Repräsentanten zweier verschiedener Phoneme zu tun.

⁶ Dieser Laut wird auch häufig durch [ɫ] transkribiert.

Wenn wir von Phänomenen wie Lippenrundung absehen, können wir also feststellen, dass das deutsche /l/-Phonem durch ein Allophon [l] repräsentiert ist, das englische dagegen durch zwei Allophone [l] und [ɫ], die ihrerseits im Albanischen zwei verschiedene Phoneme repräsentieren:

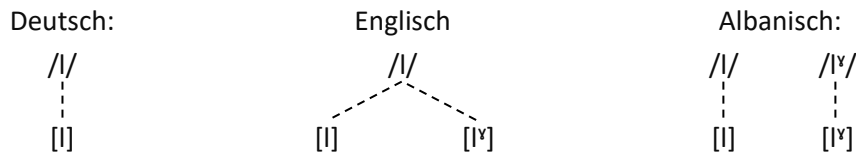


Abbildung 12: /l/ im Deutschen, Englischen und Albanischen

/f/ (Altenglisch, modernes Englisch)

Im Gang ihrer diachronen Entwicklung ist es auch durchaus möglich, dass Sprachen ihre phonologische Struktur verändern und spezifische Laute ihren Status im Phonemsystem ändern. Dazu ein paar Beispiele aus dem Altenglischen, achten Sie auf die Verteilung von [f] und [v]:

24. oft [ɔft] 'oft', rof [ro:f] 'tapfer', friþ [friθ] 'Friede', lif [li:f] 'Leben', æfter [æftər] 'danach/hinter'
 25. gifu ['=jivu] 'Geschenk', alyfan [a:ly:van] 'erlauben', æfen ['æ:vən] 'Abend', lifan ['livan] 'leben'

Die Distribution des [f] ist heterogen, die Distribution des [v] dagegen hat eine klare Systematik: es tritt immer zwischen zwei Vokalen auf.⁷ [f] und [v] sind phonetisch ähnlich und komplementär verteilt, damit repräsentieren sie dasselbe Phonem. Im modernen Englisch dagegen repräsentieren [v] und [f] verschiedene Phoneme, wie die folgenden Minimalpaare untermauern:

26. fat–vat, ferry–very, leave–leaf usw.

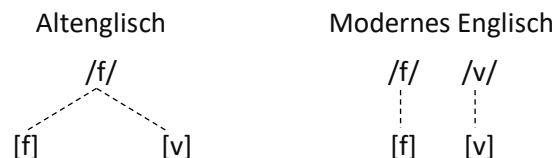


Abbildung 13: /k/ im Alt- und im modernen Englisch

Ein Echo des Umstandes, dass [f] und [v] im Altenglischen dasselbe Phonem repräsentieren, hat übrigens bis dato einen Reflex im modernen Englisch und ist z.B. in folgenden Paaren zu erkennen:

27. life–lives, wife–wives, wolf–wolves usw.

/r/ (Englisch, Koreanisch, Japanisch)

Mit diesen Erkenntnissen im Hinterkopf können wir nun zurückkehren zu dem eingangs diskutierten Phänomen der l-r-Verwechslung, die bei Sprechern des Japanischen und Koreanischen zu beobachten ist. Beginnen wir mit dem Koreanischen und betrachten dazu die folgenden Daten. Links steht jeweils die deutsche Übersetzung, in der Mitte die Verschriftlichung im Hangul, dem koreanischen Alphabet und rechts die für uns interessante Aussprache im IPA-Code. Es geht darum, zu untersuchen, wie [l] und [r]-Laut⁸ in dieser Datensammlung jeweils verteilt ist. Gehen Sie davon aus, dass die Daten repräsentativ sind, d.h. dass alle für die Generalisierung nötigen Fälle erfasst sind. Die Tabellen sind insofern vorsortiert, als Sie links Belege für [r], rechts für [l] finden:

	Hangul	IPA		Hangul	IPA
Wind	바람	[param]	Alkohol	알코올	[alko:l]
Name	이름	[irum]	Wasser	물	[mul]
Speise	요리	[jori]	Sieben	일곱	[ilɡop]
Kopf	머리	[mɾi]	Draht	철사	[tɕʌlsa]
Wolke	구름	[kurum]	Wespe	말벌	[malbʌl]
Glas	유리	[juri]	Biene	벌	[pəl]

⁷ Genauer gesagt: zwischen zwei stimmhaften Lauten, wie z.B. auch in hæfde ['hævdə] 'haben'.

⁸ [r] ist der alveolare Tap, ein Allophon des engl. /r/.

Afrika	아프리카	[apwrika]	Eisenbahn	철도	[tɕʌlto]
Bein	다리	[tari]	Gasse	골목	[kolmog]
Straße	도로	[to:ro]	Tafel	칠판	[tɕilpan]

Abbildung 14: Daten aus dem Koreanischen

Wir stellen zunächst fest, dass wir mit [r] und [l] keine Minimalpaare bilden können.

Wenn wir uns die Umgebungen ansehen, in denen [r] und [l] auftreten, sehen diese wie folgt aus:

28. [r] / {a__a, i__w, o__i, ʌ__i, u__w, u__i, w__i, a__i, o:__i}

29. [l] / {a__k, u__#, i__g, ʌ__s, a__b, ə__#, ʌ__t, o__m, i__p, . }

Die Umgebungen in (28) lassen sich leicht generalisieren: [r] tritt immer zwischen zwei Vokalen (V) auf. An dieser Position taucht das [l] nie auf. Auch die Verteilung des [l] in (29) ist eindeutig: es tritt immer zwischen einem Vokal und einem Konsonanten auf, außerdem vor der Wortgrenze. In diesen Umgebungen tritt dagegen das [r] nie auf. Damit haben wir etabliert, dass es sich bei den Liquiden [l] und [r] im Koreanischen um zwei stellungsbedingte Varianten handelt: sie sind phonetisch ähnlich und komplementär verteilt, repräsentieren also dasselbe Phonem.

$$/l/ \rightarrow \begin{cases} [r] / V_V \\ [l] / \text{sonst} \end{cases}$$

Abbildung 15: Distributionsregel für /l/ im Koreanischen

Im Japanischen dagegen ist die Sachlage anders und keinesfalls so klar. Jones schreibt dazu folgendes:

One of the most noteworthy cases of a variphone is "the Japanese r". In the pronunciation of many if not most Japanese this sound is very variable; they sometimes use a sound resembling an English fricative r (ɹ), sometimes a lingual flap r, sometimes a kind of retroflex ɻ, sometimes a kind of l and sometimes sounds intermediate between these. (Jones 1962:205-206)

Wie weiter oben in Fußnote 5 bereits gesagt wurde, benutzt Jones den Terminus »variphone« im Sinne von »Allophon in freier Variation«, und das scheint auch die bis heute akzeptierte Analyse für die Reihe von Phonem zu sein, die den japanischen r-Laut repräsentieren. Um welche Laute es sich dabei genau handelt, und ob es möglicherweise nicht doch regelhafte Stellungsvariationen gibt, wird weiterhin untersucht, aber der aktuelle Stand der Dinge ist, dass es sich bei [l], [ɹ] und [r] im Japanischen um Allophone desselben Phonems handelt, die in freier Variation auftreten.

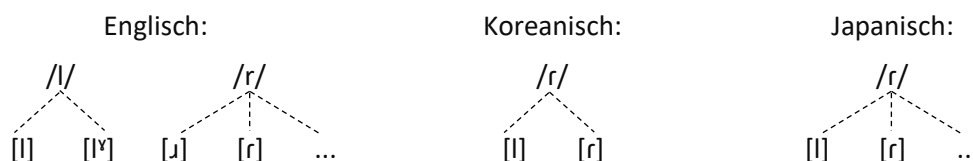


Abbildung 16: /l/ im Deutschen, Englischen und Albanischen

Auf dieser Basis ist gut nachvollziehbar, wieso Sprecher des Japanischen oder Koreanischen Probleme mit der /r/-/l/-Distinktion haben: In dem Maße, in denen die Sprecher einer Sprache das eigene phonologische System verinnerlichen, kann die Wahrnehmung für andersartige Systeme leiden: man hört die Unterschiede gar nicht mehr richtig, die für andere Sprachen phonemisch, in der eigenen Sprache aber allophonisch sind. Um diesen Sachverhalt vernünftig zu erklären, sind die Konzepte »Phonem« und »Allophon« unabdingbar.

Was die Beispiele in diesem Text zeigen, ist, dass sich Sprachen in ihrer phonologischen Struktur auf verschiedenen Ebenen unterscheiden können:

- in ihrem jeweiligen Phoneminventar (bestimmte Phoneme sind in Sprache X vertreten, nicht aber in Sprache Y)
Beispiel: engl. [w], dt. [ç] usw.
- in Anzahl und Art der Allophone, durch die ein Phonem repräsentiert sein kann,
Beispiel: engl. /l/ — {[l], [ɫ]}, dt. /l/ — {[l]}

- im Status einzelner Sprachlaute hinsichtlich der Frage, ob diese Phonemstatus haben oder nicht.
Beispiel: engl. /l/ — {[l], [lʲ]}, albanisch /l/, /lʲ/

Das Phänomen »fremdsprachlicher Akzent« kann also nicht erschöpfend mit dem Umstand erklärt werden, dass spezifische Laute aus Sprache X in Sprache Y fehlen. Es müssen auch die phonologischen Systeme genau untersucht werden, also die individuelle Organisation dieser Systeme bezüglich der Einteilung Phonem–Allophon. Hat man mit der Muttersprache einmal ein solches System internalisiert, also unbewusst verinnerlicht, ist es in einigen Fällen sehr schwierig, davon abweichende phonologische Strukturen zu reproduzieren oder überhaupt hören zu können. Zum fremdsprachlichen Akzent gehört allerdings noch mehr, da sich Sprachen nicht nur in ihren Phonemstrukturen unterscheiden, sondern auch in ihrer Phonotaktik und den in ihnen auftretenden phonologischen Prozessen. Dazu später mehr.

Phonematische Orthographie

Zum Abschluss dieses Textes können wir jetzt noch kurz einen Punkt ansprechen, der sich auf die Verschriftlichung der Sprache bezieht. Im Zusammenhang mit Alphabetschriften sehen wir, dass die Grapheme einer Sprache Phoneme repräsentieren, nicht Allophone. Basierend auf den Beispielen in diesem Text können wir diese Aussage wie folgt präzisieren: es sind in aller Regel die Phoneme einer Sprache, die auf Buchstaben bzw. Buchstabenkombinationen abgebildet werden, und nicht die Allophone. Damit ist die Orthographie phonematisch orientiert, nicht phonetisch.

Gehen wir dafür nochmal die bisherigen Beispiele durch:

- Das deutsche /ç/ ist durch zwei Allophone [ç] und [x] repräsentiert, aber nur durch ein Graphem <ch> verschriftlicht.
- Das deutsche /k/ ist u.a. durch die Allophone [k] und [kʰ] repräsentiert, aber nur durch ein Graphem <k> verschriftlicht. Im Thai, in dem /k/ und /kʰ/ jeweils Phonemstatus haben, gibt es entsprechend zwei Grapheme: ก für /k/ und ข für /kʰ/.
- Das englische /l/ ist durch zwei Allophone [l] und [lʲ] repräsentiert, aber nur durch ein Graphem verschriftlicht: <l>. Im Albanischen repräsentieren [l] und [lʲ] verschiedene Phoneme, also gibt es auch zwei Grapheme: <l> für /l/ und <ll> für /lʲ/.
- Das altenglische /f/ ist durch die Allophone [f] und [v] repräsentiert, aber nur durch ein Graphem verschriftlicht. In den wenigen Texten, in denen das Altenglische mit Runen geschrieben wurde, steht für beide Allophone nur eine solche Rune zur Verfügung: <ƿ>. In der Mehrheit der altenglischen Texte wurde das lateinische Alphabet verwendet, in dem für beide Allophone nur das <f> verwendet wird (vgl. *wifel* ['wɪvəl] 'Rüsselkäfer' (heute *weevil*), *gifan* ['gɪvən] 'geben' (heute *give*) usw.). Im modernen Englisch, in dem /f/ und /v/ unterschiedliche Phoneme repräsentieren, gibt es zwei verschiedene Grapheme: <f> und <v>.
- Das koreanische /l/ ist durch zwei Allophone [l] und [ɾ] repräsentiert, aber nur durch ein Graphem <ㄹ> verschriftlicht. Im Englischen, in dem /l/ und /r/ Phonemstatus haben, gibt es entsprechend die beiden Grapheme <l> und <r>.

Literatur

Jones, Daniel:

- 1919: The phonetic structure of the Sechuana language. In: *Transactions of the Philological Society* 28.1 (99–106).
- 1957: *An outline of English Phonetics*. Cambridge: Heffer & Sons Ltd.
- 1962²: *The Phoneme. Its Nature and Use*. Cambridge: Heffer & Sons Ltd.

Jespersen, Otto:

- 1924: *The philosophy of grammar*. New York: Henry Holt and Company.

Pike, Kenneth:

- 1971¹²: *Phonemics: a technique for reducing languages to writing*. Ann Arbor: University of Michigan Press.