

HNO 2004 · 52:837–843
 DOI 10.1007/s00106-004-1097-x
 Online publiziert: 15. Juli 2004
 © Springer-Verlag 2004

Redaktion

M. Ptok, Hannover

B. J. Kröger¹ · P. Hoole² · R. Sader³ · C. Geng⁴ · B. Pompino-Marschall⁵
 C. Neuschaefer-Rube¹

¹ Klinik für Phoniatrie, Pädaudiologie und Kommunikationsstörungen des Universitätsklinikums der RWTH Aachen · ² Institut für Phonetik und Sprachliche Kommunikation, Ludwig-Maximilians-Universität München · ³ Abteilung für Kiefer- und Gesichtschirurgie der Universitätsklinik für Wiederherstellende Chirurgie, Kantonsspital Basel – Universitätskliniken · ⁴ Zentrum für Allgemeine Sprachwissenschaft, Typologie und Universalienforschung (ZAS) Berlin
⁵ Institut für Deutsche Sprache und Linguistik, Humboldt-Universität zu Berlin

MRT-Sequenzen als Datenbasis eines visuellen Artikulationsmodells

Artikulationsmodelle können in der Phoniatrie zur Visualisierung von Sprechfehlern und damit in der Lehre (Aus- und Weiterbildung von Logopäden und Phoniatern), in der Beratung von Patienten bzw. deren Angehörigen sowie in der Therapie von Sprechstörungen [6, 9] genutzt werden. Hierzu wurde auf der Basis von Arbeiten zur artikulatorischen Modellierung des Sprechens [15] ein zweidimensionales mediosagittales Artikulationsmodell zur Visualisierung von lautlichen Zielpositionen und Artikulationsbewegungen realisiert [18]. Im Rahmen dieses Modells können zeitlich konstante lautliche Zielpositionen in Form von mediosagittalen Schnittbildern dargestellt werden (■ **Abb. 1**).

Artikulationsbewegungen und Zielpositionen

Artikulationsbewegungen hingegen werden in Form von Animationen (Videos) realisiert. Letztere zeigen die Bewegungen der Artikulationsorgane, d. h. von Zunge, Unter- und Oberlippe, Unterkiefer, Gaumensegel und Kehlkopf. Sie umfassen die Produktion von Silben, Wörtern oder auch kurzen Sätzen. Das Modell steht via Internet zur Verfügung und kann von

Wissenschaftlern ohne Einschränkung genutzt werden [17].

Die Datenbasis des visuellen Artikulationsmodells war bisher auf statische MRT-Daten beschränkt [16]. Diese Daten zeigen aber nur die Positionierungen der Artikulationsorgane von Lauten, die in Abweichung vom natürlichen Sprechen statisch gehalten – d. h. ohne Artikulationsbewegung und ohne lautlichen Kontext – realisiert werden können [30]. Dies sind z. B. Langvokale [a:], [i:], [u:] (z. B. in „Adel“, „Isar“, „Udo“), Nasallaute [m], [n], [ŋ] (z. B. in „Mann“, „eng“), Frikativlaute [f], [s], [ʃ], [ç], [x] (z. B. in „Vater“, „Ass“, „Asche“, „ich“, „ach“) und der Laterallaut [l] (z. B. in „Land“).

Statische MRT-Daten können aber nur wenig Information über Laute geben, bei denen die *Bewegung der Artikulationsorgane* eine wichtige Rolle spielt. Dies sind insbesondere die Plosivlaute [p], [t], [k], [b], [d] und [g] (z. B. „Peter“, „Tee“, „Kanne“, „Beet“, „Dach“, „Gabe“), da sie nur im Kontext mindestens eines Vokals realisiert werden können. Hier liegt ein Großteil der lautlichen Information zum Artikulationsort (bilabial, alveolar oder velar) gerade in den akustisch-perzeptiven Korrelaten der Artikulationsbewegungen vom konsonantischen Verschluss zum Vokal [25].

Koartikulation

Darüber hinaus ist die Artikulation eines Konsonanten auch vom vokalischen Kontext abhängig bzw. durch den vokalischen Kontext erst vollständig definiert (z. B. [d] in [di:] vs. [d] in [da:]). Dieses Phänomen wird als Koartikulation bezeichnet und ist ein wesentlicher Grund dafür, dass aus MRT-Daten eines statisch gehaltenen Lautes nur bedingt auf die generellen artikulatorischen Eigenschaften dieses Lautes geschlossen werden kann (s. auch die umfassende Sammlung von Arbeiten zur Koartikulation in [11]).

Durch die Verfügbarkeit der im Folgenden beschriebenen MRT-Sequenzen ist es nun möglich, mediosagittale Konturen von Konsonanten im natürlichen Sprechverlauf und damit innerhalb definierter vokalischer Kontexte zu messen. Anhand dieser Daten können die vom visuellen Artikulationsmodell berechneten koartikulatorischen Einflüsse in Hinblick auf die mediosagittalen Konturen von Konsonanten validiert und verbessert werden.

Der Inhalt dieser Arbeit wurde auszugsweise am 14.09.2003 auf der 20. Jahrestagung der Dt. Gesellschaft für Phoniatrie und Pädaudiologie, Fachmedizin für Kommunikationsstörungen, in Rostock vorgestellt.

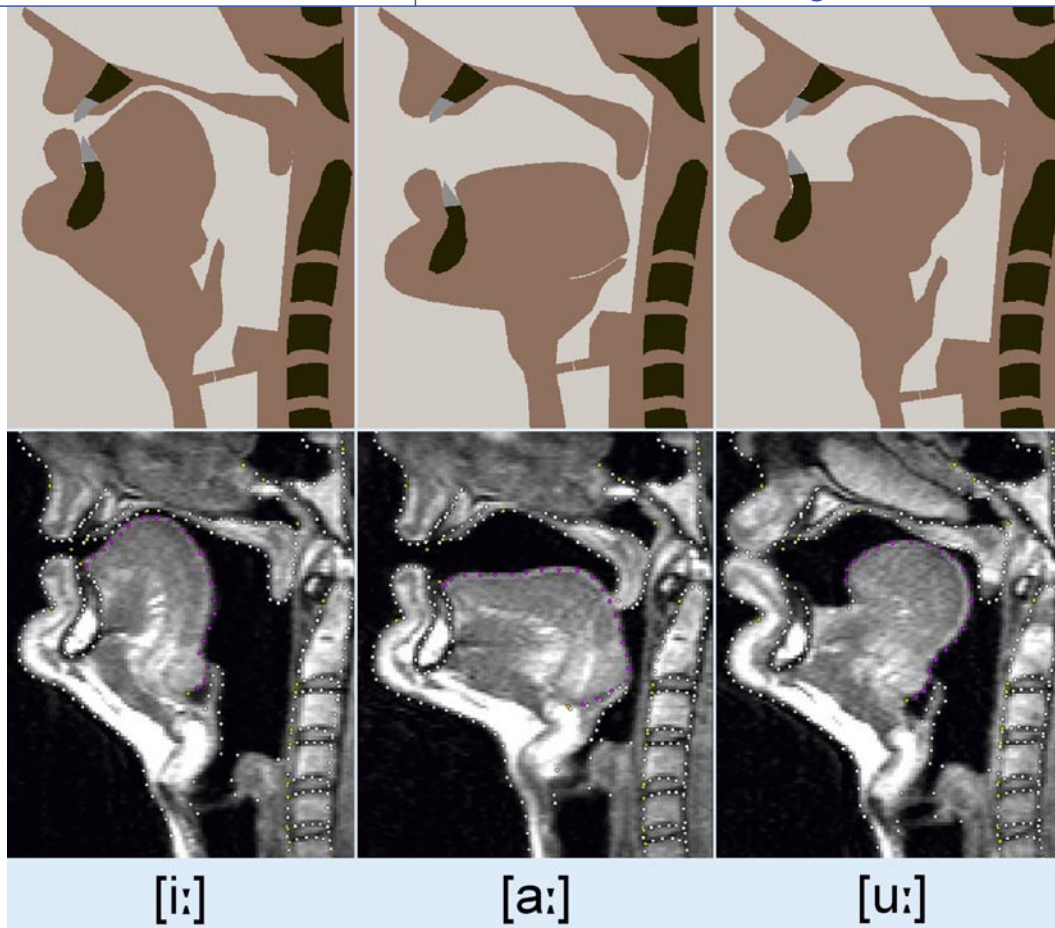


Abb. 1 ◀ **Darstellung der Eckvokale [i:], [a:] und [u:] im visuellen Artikulationsmodell (Zeile 1) und anhand von statischen MRT-Daten eines Sprechers (Zeile 2) des Hochdeutschen („Modellsprecher“).** Die Koordinaten der in den statischen MRT-Daten eingezeichneten Punktmengen stellen die Datenbasis für das visuelle Artikulationsmodell dar. Die Punktmenge des Zungenrückens ist *dunkelgrau* dargestellt. Die Realisierung des [a:] im Artikulationsmodell zeigt eine stärkere Absenkung des Unterkiefers als die MRT-Realisierung dieses Vokals

Methodik

Die MRT-Sequenzen wurden mittels eines Philips-ACS-NT-Gyroscan gewonnen ([3, 20, 21], „T1 fast gradient echo sequence, sensitivity encoding system“, Schichtdicke 10 mm). Es wurden 11 Messungen durchgeführt. Jede Messung entspricht einer Äußerung des Modellsprechers (■ **Tabelle 1**).

Messung von Logatomfolgen

Die Äußerungen 1–10 bestehen aus einer Folge von 3 Logatomen (sinnleere Neologismen mit phonotaktisch erlaubten Reihungen von Lauten, z. B. „bata“), Äußerung 11 besteht aus 5 Logatomen. Der Sprecher wurde angewiesen, die Logatome in normalem Sprechtempo direkt hintereinander zu sprechen und diese Logatomfolge über die gesamte Aufzeichnungszeit fortwährend zu wiederholen. Diese Produktion wurde während der Messung nur durch eine Atempause unterbrochen.

Die gesamte Messdauer für jede Äußerung betrug 15 s mit einer Aufzeichnungsrate von 8 mediosagittalen Schnittbildern pro Sekunde. Dies ergab eine Gesamtmenge von 1320 auszuwertenden MRT-Bildern über alle 11 Äußerungen. Das Korpus war so angelegt, dass insgesamt 12 Konsonanten [b, d, g, t, k, l, n, ŋ, s, ʃ, ç, x] mit der angegebenen Anzahl von Wiederholungen in jeweils 3 lautlichen Kontexten [i:...i:], [a:...a:] und [u:...u:] realisiert wurden (Messung 1–10). Bis auf die Kombinationen zu [b] wurden alle anderen Vokal-Konsonant-Kombinationen nur in jeweils einer Äußerung realisiert. Anhand dieser Messungen 1–10 wurden pro Logatomfolge jeweils 5 Konsonanten zur Analyse ausgewählt (s. Unterstreichung in ■ **Tabelle 1**, Spalte „Logatomfolge“). Diese Konsonanten traten je nach Anzahl der Wiederholungen der Logatomfolgen pro Messung 8- bis 12-mal auf. Messung 11 diente der Ermittlung der artikulatorischen Zielkonturen der Eckvokale [i:], [a:] und [u:].

Datenanalyse

Die Sichtung der Daten (Analyse der Bildfolge) ergab, dass bei der Rate von 8 Bildern/s nicht die artikulatorische Zielkontur jedes Konsonanten (d. h. im Fall von Plosivlauten der Zeitpunkt der maximalen oralen Verschlussbildung; im Fall von Frikativlauten der Zeitpunkt der maximalen oralen Engebildung) in einem MRT-Bild festgehalten („getroffen“) wurde. Grund hierfür ist, dass das Sprechtempo mehr als 8 Laute/s beträgt und dass keine zeitliche Synchronisation zwischen Lautproduktion und Zeitpunkt der Aufnahme eines MRT-Bildes durchgeführt werden kann.

Daher wurden im 1. Schritt der Datenauswertung für jeden Konsonant jeder Messung (1–10) alle MRT-Bilder ausgewählt, bei denen eine konsonantische Enge- oder Verschlussbildung erkennbar war (Analyse der Bildfolge). Entsprechend wurden bei Messung 11 MRT-Bilder der Eckvokale [i:], [a:] und [u:] anhand des Kriteriums einer erkennbaren palatalen, pharyngealen bzw. velaren maximalen vokali-

schen Engebildung ausgewählt. In **■ Tabelle 1** ist die aufgrund dieser Analyse der Bildfolge ermittelte Anzahl der ausgewählten MRT-Bilder pro Laut (Anzahl der MRT-Bilder pro Laut) angegeben.

Im 2. Schritt der Datenauswertung wurden nun die MRT-Bilder der Treffer nach Analyseschritt 1 für jeden Laut jeweils miteinander verglichen, und es wurde eine weitere Auswahl aus diesen Lautbildern vorgenommen (Analyse der Treffer). In diesem Schritt fielen diejenigen MRT-Bilder heraus, bei denen die konsonantische (oder vokalische) Enge- bzw. konsonantische Verschlussbildung nicht annähernd maximal war. Die Anzahl der nach diesem Analyseschritt markierten MRT-Bilder ist für jeden Laut und jede Messung (d. h. für jeden vokalischen Kontext) in **■ Tabelle 2** angegeben.

Summenbilder

Die Abweichung der Konturen der Artikulationsorgane in den so ausgewählten Bildern war so gering, dass eine Mittelung bzw. eine Überlagerung der Bilder zur Erstellung von „Summenbildern“ möglich war (s. **■ Abb. 2** für [s] aus „basa“). Die höhere Helligkeit des Summenbildes resultiert aus einer hier zusätzlich durchgeführten Kontrastverstärkung durch lineares Dehnen der Grauwerte, [1], S. 158. Die Konturerkennung ist mittels des Verfahrens von Canny [4] durchgeführt worden).

Ergebnisse

Die mittels der oben beschriebenen Methode erhaltenen 33 Summenbilder zeigen die mediosagittalen Zielkonturen für die 12 Konsonanten [b, d, g, t, k, l, n, ŋ, s, ʃ, ç, x] im definierten Kontext der 3 Eckvokale [i:, a:, u:] (**■ Abb. 3**; aus Platzgründen sind nur 5 der untersuchten 12 Konsonanten dargestellt; die Laute [ç] und [x] treten jeweils nur in komplementärer vokalischer Umgebung auf: [ç] nach [i:], [x] nach [a:] und [u:]). Aufgrund des hier gegebenen jeweiligen vokalischen Kontextes ist bei diesen mediosagittalen Schnittbildern – im Unterschied zu mediosagittalen Schnittbildern statisch gehaltener Konsonanten ohne vokalischen Kontext – die gesamte mediosagittale Kontur funktio-

Zusammenfassung · Abstract

HNO 2004 · 52:837–843
DOI 10.1007/s00106-004-1097-x
© Springer-Verlag 2004

B. J. Kröger · P. Hoole · R. Sader · C. Geng · B. Pompino-Marschall · C. Neuschaefer-Rube

MRT-Sequenzen als Datenbasis eines visuellen Artikulationsmodells

Zusammenfassung

Artikulationsmodelle können in der Phoniatrie zur Visualisierung von Sprechfehlern und damit in der Lehre, in der Patienten- bzw. Angehörigenberatung sowie in der Therapie genutzt werden. Das hier realisierte Artikulationsmodell basiert auf statischen MRT-Daten gehaltener Laute. Zur Weiterentwicklung des Modells in Hinblick auf Sprechbewegungen sollen nun ergänzend MRT-Sequenzen genutzt werden. Im vorliegenden Korpus wurden mediosagittale MRT-Schnittbilder von 12 Konsonanten im symmetrischen Kontext der 3 Eck-

vokale [i:], [a:] und [u:] mit einer Rate von 8 Bildern/s aufgezeichnet. Die Daten zeigen den starken Einfluss des vokalischen Kontextes auf die artikulatorischen Zielpositionen der Konsonanten. Es wird eine Methode zur Reduzierung der MRT-Daten für nachfolgende qualitative und quantitative Auswertungen vorgestellt.

Schlüsselwörter

MRT · Artikulation · Artikulationsmodell · Koartikulation · Sprechstörungen

MRT sequences as a database for a visual articulatory model

Abstract

Articulatory models can be used in phoniatrics for the visualisation of speech disorders, and can thus be used in teaching, the counselling of patients and their relatives, and in speech therapy. The articulatory model developed here was based on static MRI data of sustained sounds. MRI sequences are now being used to further refine the model with respect to speech movements. Medio-sagittal MRI sections were recorded for 12 consonants in the symmetrical context of the three point

vowels [i:], [a:] and [u:] for this corpus. The recording-rate was eight images/s. The data show a strong influence of the vocalic context on the articulatory target-positions of all consonants. A method for the reduction of the MRI data for subsequent qualitative and quantitative analyses is presented.

Keywords

MRI · Articulation · Articulatory model · Coarticulation · Speech disorders

Tabelle 1

Logatomfolge, Anzahl der Wiederholungen jeder Logatomfolge, ausgewählte Laute pro Messung und Anzahl der MRT-Bilder pro Laut und pro Messung nach Analyseschritt 1 (Analyse der Bildfolge)

Messung	Logatomfolge	Anzahl der Wiederholungen	Ausgewählte Laute	Anzahl der MRT-Bilder pro Laut
01	ba <u>ta</u> ba <u>da</u> ba <u>na</u>	12	[t], [b] ₂ , [d], [b] ₃ , [n]	8, 5, 6, 3, 5
02	bi <u>ti</u> bi <u>di</u> bi <u>ni</u>	10	[t], [b] ₂ , [d], [b] ₃ , [n]	7, 4, 5, 5, 4
03	bu <u>tu</u> bu <u>du</u> bu <u>nu</u>	10	[t], [b] ₂ , [d], [b] ₃ , [n]	7, 5, 5, 5, 6
04	ba <u>la</u> ba <u>sa</u> ba <u>scha</u>	10	[l], [b] ₂ , [s], [b] ₃ , [ʃ]	5, 5, 8, 5, 7
05	bi <u>li</u> bi <u>si</u> bi <u>schi</u>	10	[l], [b] ₂ , [s], [b] ₃ , [ʃ]	6, 5, 7, 5, 6
06	bu <u>lu</u> bu <u>su</u> bu <u>schu</u>	9	[l], [b] ₂ , [s], [b] ₃ , [ʃ]	6, 3, 8, 5, 7
07	ba <u>cha</u> ba <u>ga</u> ba <u>ka</u>	9	[x], [b] ₂ , [g], [b] ₃ , [k]	9, 6, 5, 4, 6
08	bi <u>chi</u> bi <u>gi</u> bi <u>ki</u>	8	[ç], [b] ₂ , [g], [b] ₃ , [k]	3, 6, 3, 6, 5
09	bu <u>chu</u> bu <u>gu</u> bu <u>ku</u>	8	[x], [b] ₂ , [g], [b] ₃ , [k]	8, 7, 5, 5, 6
10	ba <u>nga</u> bu <u>ngu</u> bi <u>ngi</u>	9	[ŋ] ₁ , [b] ₂ , [ŋ] ₂ , [b] ₃ , [ŋ] ₃	5, 5, 5, 5, 4
11	i e a o u	6	[i:], [a:], [u:]	17, 9, 15

Die ausgewählten Laute wurden in der Spalte „Logatomfolge“ jeweils unterstrichen
Bei Messung 11 ist die Anzahl der zur Analyse nutzbaren Bilder größer als die Anzahl der Messungen, da die Vokale länger als 1/8 s ausgehalten wurden und somit mehrere zeitlich benachbarte Bilder zur weiteren Analyse genutzt werden konnten

Tabelle 2

Anzahl der MRT-Bilder pro Laut und pro Messung (bzw. Kontext) nach Analyseschritt 2 (Analyse der Treffer)

Lauttyp und Kontext	Messung	Anzahl der MRT-Bilder pro Laut
[i:], [a:], [u:]	11, 11, 11	14, 9, 15
[i:bi:], [a:ba:], [u:bu:]	08, 07, 09	4, 5, 5
[i:di:], [a:da:], [u:du:]	02, 01, 03	3, 3, 2
[i:gi:], [a:ga:], [u:gu:]	08, 07, 09	3, 2, 4
[i:ti:], [a:ta:], [u:tu:]	02, 01, 03	6, 4, 4
[i:ki:], [a:ka:], [u:ku:]	08, 07, 09	4, 3, 4
[i:ni:], [a:na:], [u:nu:]	02, 01, 03	4, 5, 5
[i:ŋi:], [a:ŋa:], [u:ŋu:]	10, 10, 10	4, 6, 4
[i:li:], [a:la:], [u:lu:]	05, 04, 06	2, 5, 5
[i:si:], [a:sa:], [u:su:]	05, 04, 06	3, 5, 4
[i:ʃi:], [a:ʃa:], [u:ʃu:]	05, 04, 06	6, 6, 7
[i:çi:], [a:xa:], [u:xu:]	08, 07, 09	7, 6, 8

Der lautliche Kontext entfällt im Fall der Vokalanalyse (Messung 11)
Für den Laut [b] (Zeile 2) wurde [b]₂ aus Messung 7, 8 und 9 analysiert (s. Tabelle 1)
Die Lautverbindungen wurden nach Lauttyp des Konsonanten und vokalischem Kontext geordnet (vgl. auch Anordnung in Abb. 3). In der 2. Spalte wird die zugehörige Messung angegeben (vgl. Tabelle 1)

nal determiniert. Somit kann anhand dieser Datenbasis der koartikulatorische Einfluss der Vokale auf die Konsonantartikulation studiert werden. Die Ergebnisse stehen in Einklang mit den in der Literatur beschriebenen koartikulatorischen Effekten [5, 12, 13, 14, 22, 23, 26]:

1. Gut erkennbar ist die labiale Koartikulation hinsichtlich der Lippenrundung aufgrund des koartikulatorischen Einflusses des [u:] im Vergleich zu [i:] und [a:] bei allen Konsonanten.
2. Hinsichtlich der Lage und Formung des Zungenrückens ist der koartikulatorische Einfluss des vokalischen Kontextes für den labialen Konsonant [b] am stärksten: Die gesamte Form des Zungenrückens entspricht hier weitgehend der des vorhergehenden bzw. nachfolgenden Vokals. Dieser kontextuelle Einfluss der Vokale wird vom labialen Konsonant [b] über die koronalen Konsonanten [d, t, n, s, ʃ,] bis hin zu den dorsalen Konsonanten [g, k, ŋ, ç, x] zunehmend geringer.
3. Die Lage der Artikulationsstelle (d. h. der konsonantischen Enge- bzw. Verschlussbildung) variiert beim labialen Konsonanten und bei den koronalen Konsonanten aufgrund des vokalischen Kontextes kaum. Somit wird beispielsweise die Artikulationsstelle des [t] aufgrund des hinteren Vokals [u:] im Vergleich zum vorderen Vokal [i:] nicht sichtbar nach hinten verlagert (vgl. aber [26]). Bei den dorsalen Konsonanten [g, k, ŋ, ç, x] kann hingegen ein sichtbarer Einfluss des vokalischen Kontextes auf die Lage der konsonantischen Artikulationsstelle erkannt werden. Hier wird die konsonantische Enge- bzw. Verschlussbildung im [i:]-Kontext wesentlich weiter vorne gebildet als im Fall des [a:] bzw. [u:]-Kontextes. Beim velaren Frikativ [x] ([a:]- und [u:] -Kontext) wird die konsonantische Enge sehr weit hinten (annähernd uvular) realisiert. Der palatale Frikativ [ç] ([i:]-Kontext) zeigt hingegen eine annähernd [i:] -ähnliche Formung des Zungenrückens. Der Bereich der Engebildung liegt hier somit sehr weit vorne und deckt sich weitgehend mit dem des Eckvokals [i:].
4. Es ist erkennbar, dass das Gaumensegel beim tiefen Vokal [a:] gegenüber den hohen Vokalen ([i:] und [u:]) stärker abge-

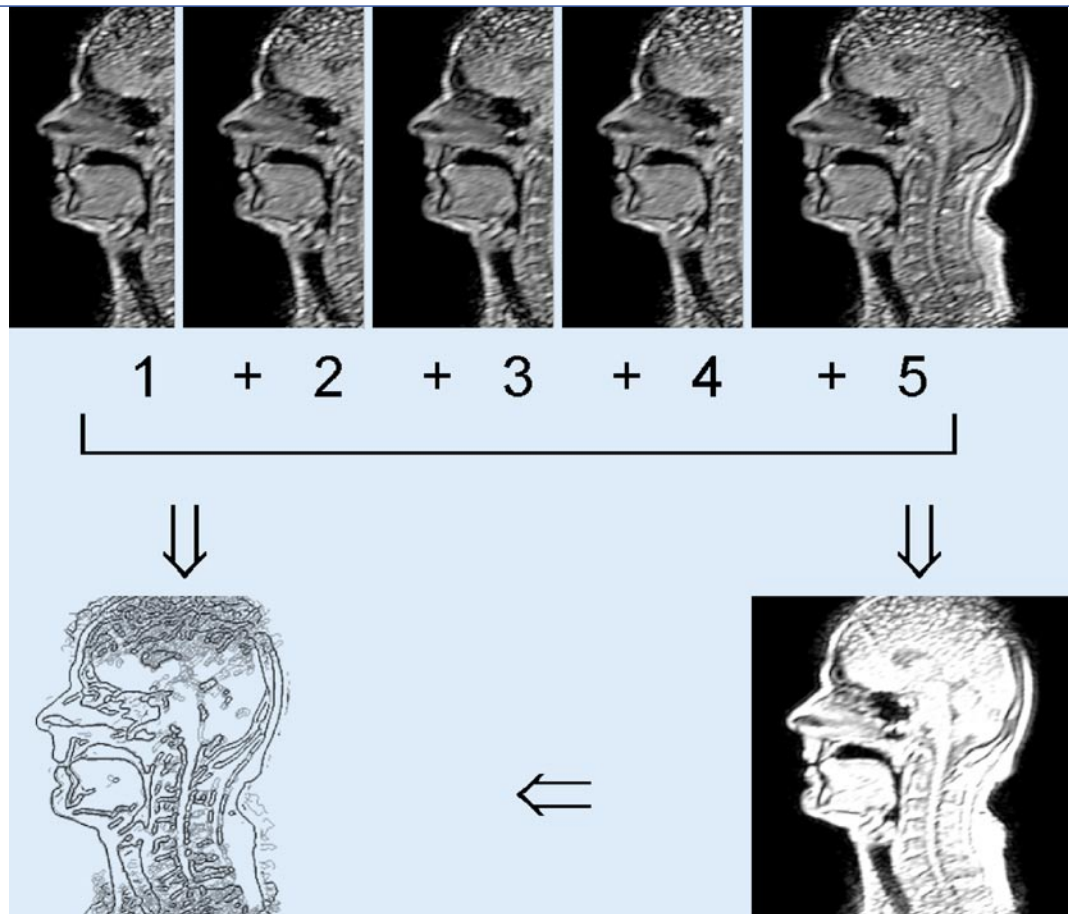


Abb. 2 ► Nach beiden Analyseschritten ausgewählte Einzelbilder (Zeile 1); das Summenbild (Zeile 2, rechts) und das Konturbild (Zeile 2, links) für [s] aus „basa“. Das Konturbild zeigt eine Überlagerung der Konturen der 5 Einzelbilder (grau) und der Kontur des Summenbildes (schwarz)

senkt ist (► Abb. 1 und Abb. 3 für [b] und [d]). Dies trägt der Tatsache Rechnung, dass der akustische Effekt der Nasalisierung bei tiefen Vokalen erst bei größerer Öffnung der velopharyngealen Pforte entsteht, sodass ein starkes Anheben des Gaumensegels hier nicht erforderlich ist (vgl. [19]).

Daten in kompakter Form

Über die auf der qualitativen Ebene auch aus der Literatur bereits zum Großteil bekannten koartikulatorischen Effekte hinaus ist ein wichtiges Resultat dieser Arbeit, dass artikulatorische und koartikulatorische Daten zur Konsonantartikulation nun in kompakter und übersichtlicher Form anhand von wenigen mediosagittalen Schnittbildern – die in 2 Dimensionen (Konsonant und Vokalkontext) geordnet werden können (s. ► Abb. 3) – verfügbar sind.

Weitergehende qualitative und quantitative artikulatorische Analysen können anhand von Literaturdaten allein nicht geleistet werden, da diese auf der Basis von

Material unterschiedlicher Sprachen und Sprecher und auch mittels verschiedener artikulatorischer Analysemethoden (z. B. Röntgen-Microbeam-Technik [13], Elektropalatographie [10], Zungensonographie [28], elektromagnetische Artikulographie [7, 8, 24, 27] erhoben wurden und damit nur schwer vergleichbar bzw. generalisierbar sind.

Diskussion und Ausblick

Die Messung mediosagittaler *Konturen* von Sprechbewegungen ist bisher mittels der MRT-Technik noch nicht mit ausreichender zeitlicher Auflösung möglich. Messmethoden mit ausreichender Zeitauflösung (40 Bilder/s und mehr), liefern aber zumeist lediglich die Lage bzw. den Bewegungsverlauf weniger ausgezeichnete Oberflächenpunkte von Artikulatoren (siehe [2]). Die Magnetresonanztomographie ist hingegen in der Lage, die *gesamte mediosagittale Kontur* mit ausreichender räumlicher Auflösung wiederzugeben.

Konsonantische Zielkonturen

Die hier eingesetzte MRT-Methode kann zwar im derzeitigen Stadium (Aufzeichnung von 8 Bildern/s) noch keine Sprechbewegungsverläufe auflösen, erlaubt aber mittels der hier vorgestellten Analysemethode bereits die Erfassung *artikulatorischer Zielkonturen von Lauten bei normalem fließendem Sprechen*. Diese Daten ermöglichen damit die Kontrolle des lautlichen Kontextes und geben damit insbesondere Aufschluss über die Zielkonturen von Konsonanten in unterschiedlichen vokalischen Kontexten.

Anhand dieser Daten kann quantifiziert werden, welche Bereiche konsonantischer Zielkonturen kontextuell annähernd invariant bleiben (primäre Artikulation, [18], S. 405) und welche Bereiche kontextuell variieren (sekundäre Artikulation ([18], S. 405)). Diese Analyse ist in Hinblick auf primäre und sekundäre Artikulation für die Entwicklung quantitativer Artikulationsmodelle ([18], S. 405) von großer Bedeutung und kann anhand

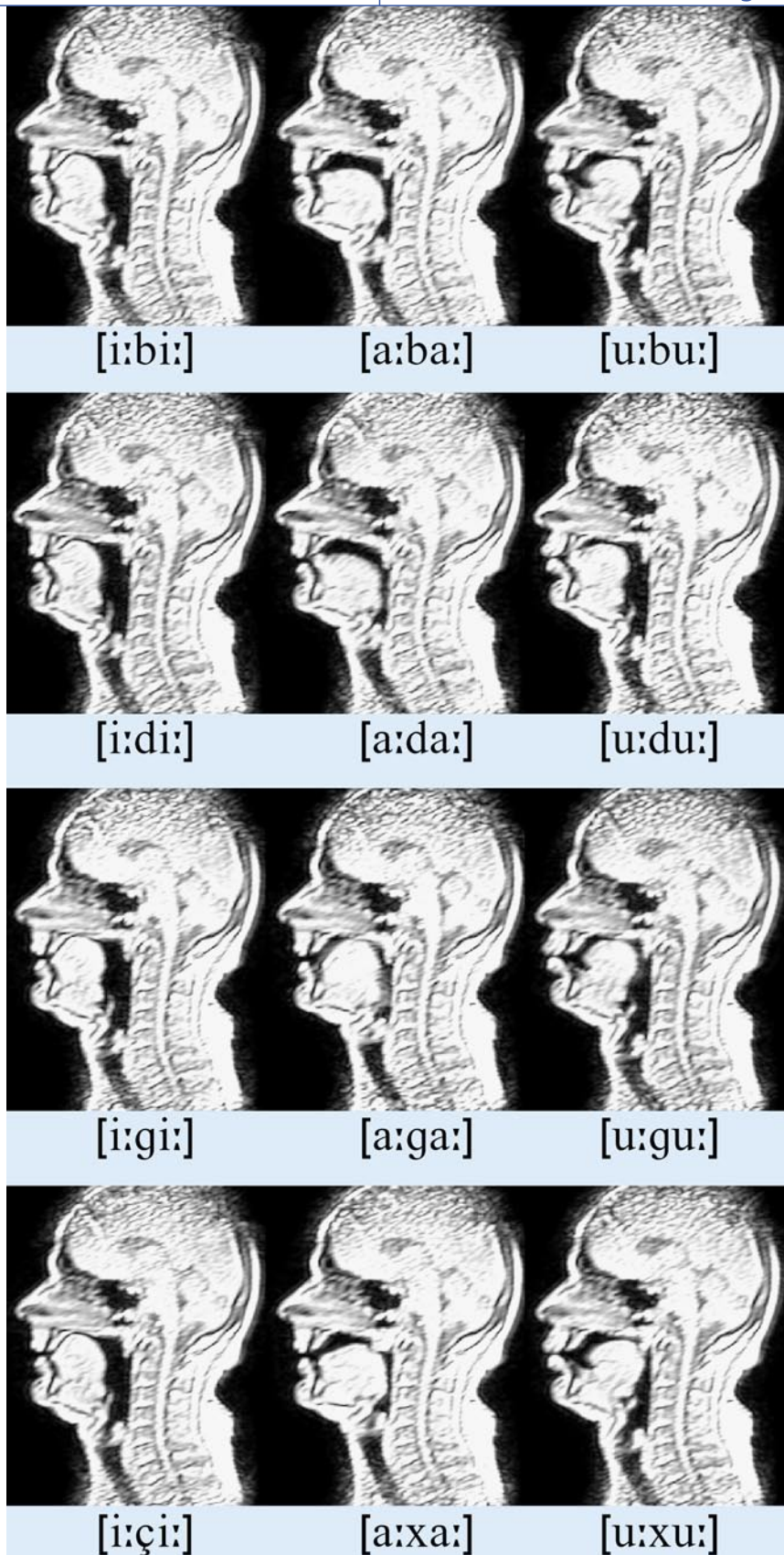


Abb. 3 ▲ Zusammenstellung der Summenbilder der Konsonanten [b], [d], [g] und [ç]/[x] (Zeile 1–4) im jeweiligen Kontext der Eckvokale [i:], [a:] und [u:] (Spalte 1–3)

von den in der Literatur üblichen mediosagittalen Schnittbildern statisch gehaltener Konsonanten ohne vokalischen Kontext (für das Deutsche z. B. [29]) nicht durchgeführt werden.

Weiterentwicklung

Insbesondere wurde in dieser Arbeit eine Methode zur Reduzierung eines MRT-Datensatzes auf wenige aussagekräftige Summenbilder vorgestellt. Bei dem hier genutzten Korpus zur Analyse von 12 Konsonanten [b, d, g, t, k, l, n, ŋ, s, ʃ, ç, x] im Kontext von jeweils 3 Vokalen [i:, a:, u:] führt dies zu 33 konsonantischen Summenbildern ([ç] und [x] sind komplementär verteilt).

Interessant ist nun die weitere Nutzung der Summenbilder in Hinblick auf die quantitative Validierung und Weiterentwicklung des visuellen Artikulationsmodells [18]. In weiteren Arbeitsschritten können nun erstens die räumlichen Ausdehnungen der im Artikulationsmodell definierten primären und sekundären Artikulation und zweitens die Effekte der Vokal-Konsonant-Koartikulation ([18], S. 406) quantitativ erfasst und in das Artikulationsmodell übertragen werden.

Korrespondierender Autor

Prof. Dr. phil. Dipl.-Phys. B. J. Kröger

Klinik für Phoniatrie,
Pädaudiologie und Kommunikationsstörungen,
Universitätsklinikum Aachen,
Pauwelsstraße 30, 52074 Aachen
E-Mail: bkroeger@ukaachen.de

Danksagung

Unser besonderer Dank gilt A. Zimmermann und K. Mady (Klinik und Poliklinik für Mund-Kiefer-Gesichts-Chirurgie, Klinikum r. d. Isar, TU München) und A. Beer und C. Hannig (Institut für Röntgendiagnostik, Klinikum r. d. Isar, TU München) für die Erstellung der MRT-Sequenzen.

Interessenkonflikt: Der korrespondierende Autor versichert, dass keine Verbindungen mit einer Firma, deren Produkt in dem Artikel genannt ist, oder einer Firma, die ein Konkurrenzprodukt vertreibt, bestehen.

Literatur

1. Abmayr W (1994) Einführung in die digitale Bildverarbeitung. Teubner, Stuttgart
2. Barlow SM, Finan DS, Andreatta RD, Paseman LA (1997) Kinematic measurement of the human vocal tract. In: MR McNeil (ed) Clinical management of sensorimotor speech disorders. Thieme, New York, pp 107–148
3. Beer AJ, Hellerhoff P, Zimmermann A, Mady K, Sader R, Rummeny EJ, Hannig C (2004) Real-time magnetic-resonance imaging (MRI) for analysis of velopharyngeal closure in comparison to videofluoroscopy. J Magn Reson Imaging (eingereicht)
4. Canny J (1986) A computational approach to edge detection. IEEE Transactions on Pattern Analysis and Machine Intelligence, vol PAMI-8: 679–698
5. Carney P, Moll K (1971) A cinefluorographic investigation of fricative consonant-vowel-coarticulation. Phonetica 23: 193–202
6. Diem A, Kröger BJ (2003) Einsatz eines Programms zur Visualisierung von Sprechbewegungen in der Therapie eines Kindes mit motorisch bedingter Artikulationsstörung. Abstractband zum 32. dbl-Jahreskongress 19.–21. Juni 2003 in Karlsruhe, S 40
7. Engelke W, Bruns T, Striebeck M, Hoch G (1996) Mid-sagittal velar kinematics during production of VCV sequences. Cleft Palate Craniofac J 33: 236–244
8. Goozee JV, Lapointe LL, Murdoch BE (2003) Effects of speaking rate on EMA-derived lingual kinematics: a preliminary investigation. Clin Linguist Phonet 17: 375–381
9. Gotto J (in Vorb.) PC-gestützte Therapie der Sprechapraxie – Eine Einzelfallstudie. Diplomarbeit, Studiengang Lehr- und Forschungslogopädie, RWTH Aachen
10. Hardcastle WJ, Gibbon F, Nicolaidis K (1991) EPG data reduction methods and their implications for studies of lingual coarticulation. J Phonet 19: 251–266
11. Hardcastle WJ, Hewlett N (1999) Coarticulation. Cambridge University Press, Cambridge
12. Heike G (1987) Coarticulation' in an articulatory model of German. Proc 11th Int Congress Phonetic Sci 1: 214–216
13. Kiritani S, Itoh K, Fujimura O (1975) Tongue-pellet tracking by a computer-controlled X-ray microbeam system. J Acoust Soc Am 57: 1516–1520
14. Kiritani S, Sawashima M (1987) The temporal relationship between articulations of consonants and adjacent vowels. In: R Channon, L Shockley: In Honor of Ilse Lehiste. Fortis, Dordrecht, pp 139–149
15. Kröger BJ (1998) Ein phonetisches Modell der Sprachproduktion. Niemeyer, Tübingen
16. Kröger BJ, Winkler R, Mooshammer C, Pompino-Marschall B (2000) Estimation of vocal tract area function from magnetic resonance imaging: Preliminary results. Proc 5th Semin Speech Production: Models and Data. Kloster Seeon, Bavaria, pp 333–336
17. Kröger BJ (2002) <http://www.phoniatrie.ukaachen.de > Lehre > SpeechTrainer>
18. Kröger BJ (2003) Ein visuelles Modell der Artikulation. Laryngorhinootologie 82: 402–407
19. Lubker JF (1968) An electromyographic-cineradiographic investigation of velar function during normal speech production. Cleft Palate J 5: 1–8
20. Mady K, Sader A, Zimmermann A, Hoole P, Beer A, Zeilhofer HF, Hannig C (2001) Use of real-time MRI in assessment of consonant articulation before and after tongue surgery and tongue reconstruction. In: Maassen B, Hulstijn W, Kent R, Peters H, Lieshout P van (eds) Speech motor control in normal and disorderd speech. 4th Int Speech Motor Conference, Nijmegen, pp 142–145
21. Mady K, Sader A, Zimmermann A, Hoole P, Beer A, Zeilhofer HF, Hannig C (2002) Assessment of consonant articulation in glossectomy speech by dynamic MRI. Proc ICSLP 2002, Denver, pp 961–964
22. Menzerath P, Lacerda A (1933) Koartikulation, Steuerung und Lautabgrenzung. Ferdinand Dümmlers, Berlin
23. Metoui M (2001) Strategien der Artikulation. Über die Steuerungsprozesse des Sprechens. Shaker, Aachen
24. Perkell JS, Cohen MH, Svirsky MA, Mattheis ML, Garabeta I, Jackson MTT (1992) Electromagnetic midsagittal articulometer systems for transducing speech articulatory movements. J Acoust Soc Am 92: 3078–3096
25. Pompino-Marschall B (1995) Einführung in die Phonetik. De Gruyter, Berlin
26. Recasens D (1999) Lingual coarticulation. In: Hardcastle WJ, Hewlett N: Coarticulation – theory, data and techniques. Cambridge University press, Cambridge, pp 80–104
27. Schönle PW, Grabe K, Wenig P, Hohne J, Schrader J, Conrad B (1987) Electromagnetic articulography: use of alternating magnetic fields for tracking movements of multiple points inside and outside the vocal tract. Brain Lang 31: 26–35
28. Stone M, Davis EP (1995) A head and transducer support system for making ultrasound images of tongue/jaw movement. J Acoust Soc Am 98: 3107–3112
29. Wängler HH (1976) Atlas deutscher Sprachlaute. Akademie-Verlag, Berlin
30. Wein B, Drobnitzky M, Klajman S, Angerstein W (1991) Evaluation of functional positions of tongue and soft palate with MR imaging: initial clinical results. J Magn Reson Imaging 1:381–383

Wie wir lernen, unseren Weg zu finden

Kernspintomographie-Studie zeigt, wie das Gehirn automatisch und unbewusst wichtige Wegmarkierungen in einer bestimmten Hirnregion abspeichert

Damit wir den richtigen Weg durch unsere Umgebung finden, müssen wir uns wichtige Informationen über die Wegstrecke merken. Bisher war nicht bekannt, wie das Gehirn dies bewerkstelligt. Gabriele Janzen und Miranda van Turenhout vom Max-Planck-Institut für Psycholinguistik und dem niederländischen FC Donders Centrum für kognitives Neuroimaging in Nijmegen haben jetzt mittels funktioneller Magnetresonanztomographie nachgewiesen, dass das Gehirn selektiv nur jene Markierungen im so genannten parahippocampalen Gyrus abspeichert, die an navigationsrelevanten Positionen entlang einer Route platziert sind (Nature Neuroscience 7:673–7). Diese Abspeicherung von Schlüsselinformationen erfolgt automatisch und häufig unbewusst und bildet offensichtlich die neuronale Grundlage für effizientes und erfolgreiches Navigieren durch bekannte oder unbekannte Umgebungen.

Quelle: Max-Planck-Gesellschaft zur Förderung der Wissenschaften e.V.