

Akkommodation i jysk regionalsprog – Akustiske og perceptuelle perspektiver på lydlig tilpasning i interaktion

Anna Kai Jørgensen
Københavns Universitet

1. Introduktion

Den internationale forskning i akkommodation har længe undersøgt hvordan talere tilpasser sig hinanden i både imitation og interaktion. Forskningen tager primært udgangspunkt i engelsk, mens akkommodation i dansk er sjældent undersøgt (cf. Rathje, 2008; Rathje et al., 2021).

Akkommodationsforskningen beskæftiger sig bredt set med sproglig tilpasning på alle niveauer, fra det segmentelt fonetiske til det pragmatiske (cf. Pardo et al., 2017; Rathje et al., 2021). Denne undersøgelse placerer sig et sted midt imellem disse yderpunkter og undersøger fonetisk tilpasning på ordniveau. Her henvises der ikke til leksikalsk tilpasning, men snarere til en mere holistisk lydlig tilpasning som vil blive beskrevet nærmere nedenfor.

Akustiske akkommodationsanalyser peger i mange retninger med den generelle fællesnævner at akkommodation er udbredt og at én eller anden form for sproglig tilpasning kan betragtes som *default* i menneskelig interaktion (MacLeod, 2013). Fokus er ofte på de psykologiske og kognitive mekanismer der får folk til at tilpasse sig deres samtalepartner, for eksempel ved at inddrage talernes personlighedstræk eller specifikke kompetencer (cf. Aguilar et al., 2016; Cohen Priva & Sanker, 2020; Yu et al., 2013), mens der ikke er tradition for at beskæftige sig med mere sociolingvistiske spørgsmål af variationistisk eller sprogforandringsmæssig karakter.

Pardo (2006: 3) peger på at det interaktionelle niveau kan være *the missing link* mellem individet og større forandringsprocesser i sprogsamfundet. I forlængelse heraf foreslår jeg akkommodationsundersøgelsen som et værktøj til at undersøge lingvistisk variation på niveauet mellem individ og samfund, nemlig i interaktionen. En større viden om hvordan tilpasningsprocesser udspiller sig i interaktioner, kan hjælpe os til at forstå de generelle mekanismer for sproglig tilpasning i større og mindre skala i det danske sprogsamfund.

Denne undersøgelses fokus er på systematisk variation i samtaler mellem to personer og søger at besvare følgende spørgsmål: Tilpasser de to talere sig hinanden? Er der forskel i talestile alt efter om der læses op eller interageres med en samtalepartner? Konvergerer nogle talere mere end andre? Er der regionale forskelle i akkommodationen? Disse spørgsmål undersøges med et akustisk og et perceptuelt mål for lydlig tilpasning ud fra den antagelse at produktion og perception er adskilte størrelser med forskellige mekanismer. Undersøgelsen er en del af et større projekt hvor også det segmentelle niveau undersøges, mens fokus her er på ordniveau.

2. Dataindsamling

Undersøgelsens datagrundlag er optagelser af interaktioner mellem 14 talere, alle førsteårsstuderende på læreruddannelserne, de såkaldte *University Colleges*, i Hjørring og Jelling. Optagelserne er foretaget på uddannelsesinstitutionerne som led i projektet *Speaking Up* der undersøger sprogets rolle for social mobilitet i uddannelse. Social mobilitet er ikke i fokus her. Deltagerne er sat sammen i par efter hvem de tidligere har oplyst at tilbringe mest tid med i skoletiden.

For at opnå et ensartet og sammenligneligt datamateriale tager interaktionerne mellem talerne udgangspunkt i en *map task* udviklet specielt til formålet, herefter kaldet en kortøvelse. Kortøvelsen er en videreudvikling af den der blev brugt i indsamlingen af DanPASS-korpusset (Grønnum, 2022), som igen er baseret på den oprindelige HCRC Map Task (Anderson et al., 1991). Kort fortalt går øvelsen ud på at hver deltager får tildelt et kort med en række kendemærker (*landmarks*) hvoraf nogle er ens for begge kort, mens andre er forskellige. Deltagerne følger en rute der, i den oprindelige version,

Ud over øvelsesdesignet har det oprindelige kortdesign fået en visuel overhaling. Dette er gjort både for bedre at kunne styre de elicerede ord og lyde, øge læsbarheden, som efterhånden var udfordret af gentagen kopiering, og til sidst for at kortene skulle virke upåfaldende på unge deltagere i 2020'erne. Der er desuden den fordel ved det opdaterede design at det bruger Microsoft Words ikoner som frit kan bruges og videredistribueres (en redigerbar version af kortøvelsen kan fås ved at kontakte forfatteren). Et eksempel på et af de to kortsæt i kortøvelsen kan ses i Figur 1.

På hvert kortsæt i kortøvelsen er der 16 kendemærker hvilket giver 32 unikke kendemærker for hele øvelsen. Den tilhørende ordliste indeholder de 32 kendemærker som i oplæsningen indsættes i bæresætningen: ”Nummer _ er en/et _____” efter forlæg fra Pardo et al. (2019). Datagrundlaget for undersøgelsen er alle ord sagt af begge talere i et par, dvs. 459 tokens i alt.

3. Akustisk og perceptuel akkomodation – to mål

Undersøgelsen tager to metoder i brug til at måle akkommodation på ordniveau. Det første mål er akustisk, det andet perceptuelt. Hvor det første mål er et forsøg på at kvantificere den sproglige tilpasning på ordniveau, er formålet med det andet mål at afprøve det akustiske mål perceptuelt; hører lyttere i perceptionsundersøgelsen det som indfanges i den akustiske undersøgelse? Udenlandske studier anvender og kombinerer disse to tilgange, men der findes, så vidt vides, ingen studier af dansk som anvender nogle af metoderne.

I det følgende vil jeg kort gennemgå de to metoder: *amplitude envelopes* og AXB-studiet.

3.1 Akustisk lighed: *Amplitude envelopes*

Amplitude envelopes blev først introduceret til akkommodationsforskningen i Lewandowski (2012) i en undersøgelse af tyske L2 lørnere af engelsk. Undersøgelsens fokus var om mennesker med mere fonetisk talent (*phonetic talent*) tilpasser sig mere til deres samtalepartner (det gør de). Til undersøgelsen bruges *amplitude envelopes* til at kvantificere en mere holistisk tilpasning, en nuancering af de segmentelle mål der hidtil havde domineret feltet (cf. Babel, 2010, 2012; Tilsen, 2009; Yu et al., 2013).

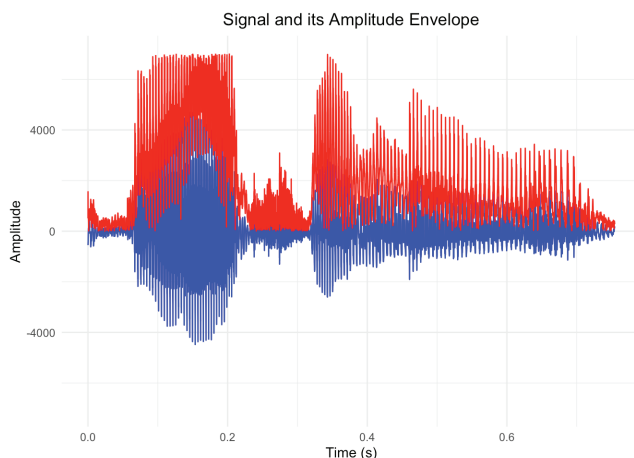
En *amplitude envelope* – på dansk herfra *amplitudekonvolut* – er en repræsentation af de amplitudinale udsving i en lyd, hvad Lewandowski & Jilka beskriver som ”a smoothed global picture of the energy present in the signal” (Lewandowski & Jilka, 2019: 9). Eftersom lyde udfolder sig i tid, er dette samtidig en måde at få en tidslig repræsentation af en lyd, fremfor blot et øjebliksbillede, sådan som det kan siges at være tilfældet i målinger af formantfrekvenser eller *Voice Onset Time* hvor frekvenserne udtrækkes på et, eller ganske få, tidspunkter i lyden. Som både Lewandowski (2012) og Lewandowski & Jilka (2019) påpeger, kan der være store individuelle forskelle på hvad lyttere og talere reagerer på i det akustiske signal, og derfor giver en sammenligning af amplitudekonvolutter et mere holistisk billede af ligheden mellem to ord eller fraser:

Taking a more global measurement, which comprises several or even all features present in the speech signal, should make it possible to capture convergence happening literally “everywhere” in the acoustic signal (Lewandowski, 2012: 134).

Man kan tænke på amplitudekonvolutten som lydens akustiske profil, på samme måde som et menneskes profil er et omrids af ansigtets former.

Til at få fat i en lyds amplitudekonvolut anvendes en *Hilbert transform* (Hilbert, 1912). En sådan transformation efterligner processerne involveret i den menneskelige hørelse, og outputtet består af en *temporal fine structure* og en *amplitude* (eller *temporal envelope*) (He & Dellwo, 2016; Moore, 2008). Det er sidstnævnte der er interessant i denne sammenhæng.

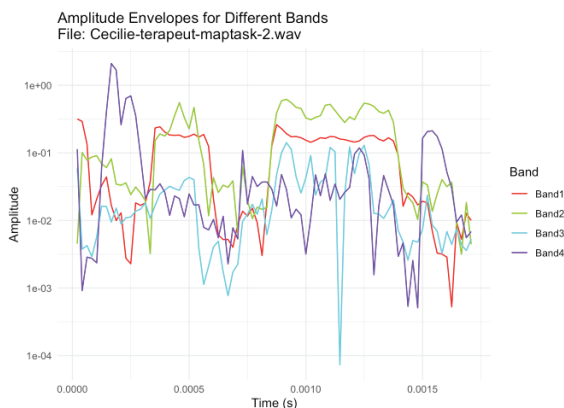
Figur 2 viser et eksempel på en lyd og dets amplitudekonvolut (markeret med rød). I blå ses oscillogrammet for lyden af en deltager der siger ordet *datalager*. Som man næsten kan fornemme ved at se på figuren, er der mest akustisk energi i /a/’et i første stavelse (begyndende omkring tidspunkt 0.1) hvorefter energien aftager efter det andet /a/ i *data*. Dette afspejler at der er tryk på første stavelse i ordet, /’datalager/.



Figur 2: Et eksempel på en lyd (blå) og dens amplitudekonvolut (rød). Lyden er en deltager fra Hjørring der siger *data*lager.

Som det fremgår af Figur 2, er der masser af bittesmå udsving i amplitudekonvolutten i sin rå form. Dette gør den svær at sammenligne med en tilsvarende lyd fra en anden taler, og derfor er det en del af processen at dele lyden op i logaritmiske bånd og herefter udglatte konvolutterne for de enkelte bånd. Proceduren der er brugt i det følgende, er inspireret af Lewandowski (2012), Lewandowski & Jilka (2019) og Abel & Babel (2017). Jeg vil henvise til de enkelte artikler for flere detaljer.

Første skridt i proceduren er for hver lyd at udtrække fire bånd med logaritmisk afstand mellem 80 og 7800 Hz (Abel & Babel, 2017; Wade et al., 2010). Herefter køres hvert bånd gennem en Hilbert transform, så der udtrækkes en amplitudekonvolut for hvert af de fire bånd. Til sidst filtreres og downsamples de fire amplitudekonvolutter for at udjævne udsving, som forklaret ovenfor. Denne proces resulterer i noget der ligner Figur 3.



Figur 3: Et eksempel på resultatet af den beskrevne procedure. På figuren ses de udjævnede amplitudekonvolutter for de fire logaritmiske bånd i ordet *terapeut* sagt af en taler fra Hjørring. X-aksen repræsenterer tid, mens y-aksen repræsenterer amplitudeudsving.

Efter udtrækning og udglatning af de fire amplitudekonvolutter for hver lyd kan disse sammenlignes med en tilsvarende lyd for samtalepartneren. Vær opmærksom på at formålet med undersøgelsen er at undersøge akustisk lighed mellem to talere i en interaktion, og sammenligningen af de to udtaler sker ved en krydskorrelation af de fire bånd for hvert ord sagt af hver taler. Krydskorrelation kan sammenligne forskelle mellem tidsserier, og derfor er målet oplagt til at sammenligne lyd (cf. Wade et al., 2010).

I eksemplet ovenfor med ordet *terapeut* (Figur 3) vil hvert af de fire bånd i den ene talers udtale af lyden blive krydskorreleret med samtalepartnerens tilsvarende. Den ene talers bånd 1 korreleres altså med partnerens bånd 1, bånd 2 med partnerens bånd 2 osv. Til sidst beregnes en samlet *match value* for hvert ordpar, så krydskorrelationsværdierne for hver sammenligning af de fire bånd samles til én enkelt værdi. To ords *match value* repræsenterer altså et gennemsnit af den samlede lighed mellem lydenes fire bånd og afspejler dermed ikke eventuelle forskelle i lighed mellem båndene. Alt dette er gjort vha. et script i databehandlingsprogrammet R, som kan fås ved at kontakte forfatteren. Efter endt krydskorrelation ender

hvert ordpar med en værdi mellem 0 og 1: jo tættere på 1, jo mere ens er de to udtaler af det samme ord.

3.2 Perceptuel lighed: AXB-studiet

Et AXB-studie er et studie af forskelle mellem tre stimuli: A, X og B. I sin grundform består hvert spørgsmål i et AXB-studie af tre stimuli, hvoraf to af dem (A og B) vurderes op mod det tredje (X). Disse stimuli præsenteres for en testdeltager som bliver bedt om at vurdere om A eller B er mest lig med eller forskellig fra X, alt efter undersøgelsesspørgsmålet (Greenaway, 2017). Rækkefølgen af de tre stimuli kan varieres, og der findes studier der anvender rækkefølgen XAB (Kim et al., 2011) eller ABX (Greenaway, 2017), hvor X repræsenterer den stimulus de to andre stimuli skal vurderes i forhold til. I den lingvistiske akkommodationsforskning anvendes oftest rækkefølgen AXB. I disse studier er stimuli X i midten, mens rækkefølgen på de andre stimuli varieres, så testdeltageren bliver præsenteret for stimuli i både rækkefølgen AXB og BXA (Goldinger, 1998; Namy et al., 2002; Pardo et al., 2013, 2017). Hovedpointen her er at X er i midten da det modvirker at nogle af de kendte rækkefølgeeffekter kommer til at skævvride data; det er også den rækkefølge der er anvendt i nærværende studie. Kombinationen af stimuli fremgår af Figur 4.

Ligesom den ovenfor beskrevne metode giver et mere holistisk billede af akkommodation end segmentelle analyser, giver en perceptuel vurdering af lydlig lighed ved hjælp af AXB-testen et mål for hvad lyttere rent faktisk hører når de lytter til de to talere. Der er ingen grund til at antage at lyttere hører det samme som de akustiske mål indfanger. Tværtimod er der evidens for at perceptuelle mål fanger andre aspekter ved akkommodation end de akustiske (cf. Babel & Bulatov, 2012; Pardo et al., 2017).

I dette undersøgelsesdesign, hvor målet altså er at vurdere lydlig lighed mellem to talere, repræsenterer stimulus X den ene taler i talerparrets udtale af et ord, mens denne er flankeret af to forekomster af samme ord sagt af samtalepartneren. Det er altså den ene parts udtale der holdes op mod den andens. X er altid et ord sagt i løbet af kortøvelsen fordi det antages at det er i løbet af interaktionen der foregår en sproglig påvirkning, hvis det er tilfældet, og at det er denne udtale

en samtalepartner vil tilpasse sig. Samtalepartnerens forekomster (A og B) består af udtaler sagt under oplæsning af ordlisterne før og efter kortøvelsen, eller af en forekomst af samme ord under kortøvelsen. Som det fremgår af Figur 4, kombineres disse stimuli, så der til sidst er seks mulige kombinationer for hvert ordpar.



Figur 4: En illustration af rækkefølgen af de præsenterede stimuli. Hver ordkombination optræder seks gange.

Stimuli består af lydbidder klippet ud af optagelserne af kortøvelsen som deltagerne har gennemført i par, og de individuelt indtalte ordlister. Alle forekomster af kendemærker fra kortene er klippet ud af optagelserne, og forekomster med fejl der gør dem dårligt sammenlignelige med andre forekomster, er sorteret fra. Det kan være hvis en deltager bøjer ordet i en form som ikke svarer til den der fremgår af ordlisterne, eller hvis der mumles eller ordet på anden måde modificeres. Hvis én deltager fx siger *datalager*, og den anden siger *datalageret* ved den ene forekomst, men *datalager* ved den anden, kan vi ikke vide om en eventuelt vurderet lighed mellem de to ubestemte former skyldes akkommodation i udtale eller bøjning. Brug af kendeord er ikke medtaget i lydklippene.

Efter udtrækning af alle ord, som beskrevet ovenfor, er stimuli kombineret efter mønstret der fremgår af Figur 4. Hver kombination af tre stimuli varer et par sekunder. Disse lægges ind i et spørgeskema designet og distribueret via onlinetjenesten <http://www.soscisurvey.de/>. Stimuli er fordelt på 15 mindre spørgeskemaer med ca. 80 spørgsmål i hvert for at beholde en overskuelig længde (cf. Pardo et al., 2017: 644). Det tager ca. 10-15 minutter at gennemføre et spørgeskema. Ved hjælp af en *randomizer* tildeles deltagerne tilfældigt og automatisk én af de 15 spørgeskemaer når de klikker på linket til undersøgelsen. Figur 5 viser et eksempel på spørgeskemaets design.



KØBENHAVNS
UNIVERSITET

2 af 79

Hvilken udtale lyder mest som den midterste lyd?

Første

Sidste

Figur 5: Spørgeskemaets design. Lydfilen med stimuli spiller automatisk ved start af spørgsmålet.

Efter en forsøgt bred distribution af undersøgelsen blev der indsamlet 68 besvarelser. Selvom undersøgelsen nåede ud til mange egne af landet, er størstedelen af deltagerne opvokset i Københavnsområdet ($n = 47$) med en medianalder på 28. Svartider over 10.000 millisekunder blev fjernet for at undgå en skævvridning af resultaterne, hvilket er almindeligt i denne type undersøgelse (cf. Greenwald et al., 2003).

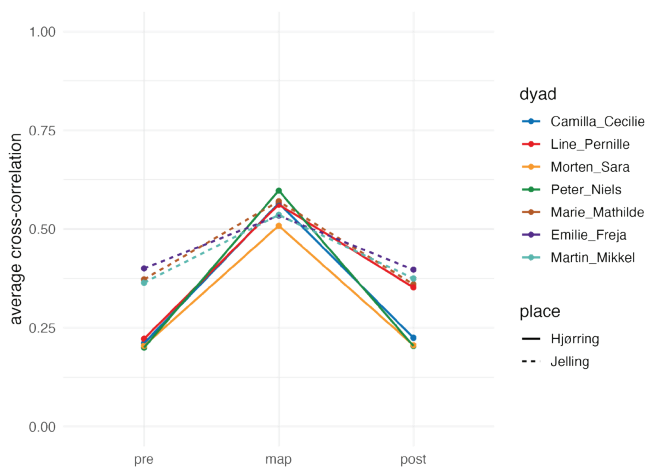
4. Resultat af de to undersøgelser

I det følgende vil jeg kort skitsere hvordan jeg har analyseret de indsamlede data og efterfølgende præsentere resultaterne af undersøgelsen.

4.1 Resultater af amplitudekonvolut-analysen

Som beskrevet i afsnit 3.1, resulterer hver sammenligning af et ordpar i en *match value* mellem 0 og 1. Denne værdi findes ved at køre samtlige ordpar for hver taler gennem et R script (jf. afsnit 3.1.).

Efter at en *match value* er beregnet med krydskorrelation for hvert ordpar i hvert talerpar, beregner scriptet også parrets samlede *match value* henholdsvis før, under, og efter kortøvelsen. Denne værdi repræsenterer den akustiske lighed mellem to talere i et talerpar; jo højere tal, jo mere lighed. Resultaterne per talerpar kan ses i Figur 6



Figur 6: Akustisk lighed efter situation (x-aksen). *Match values* er repræsenteret med et tal mellem 0 og 1 på y-aksen (*average cross-correlation*). Hvert talerpar er repræsenteret af en farve, og de to steder repræsenteres af forskellig linje-skravering.

Det første resultat der kan udledes af figuren, er at talerne generelt er mest ens i løbet af kortøvelsen (*map*) sammenlignet med de to oplæsningssituationer (*pre* og *post*), men at den akustiske lighed ikke overstiger en *match value* på 0,7. Disse resultater stemmer godt overens med tidligere studier som finder værdier mellem 0,35 og 0,75 (Abel & Babel, 2017; Lewandowski & Jilka, 2019), og eftersom vi har at gøre med forskellige talere, ville en *match value* på 1 være mistænkelig. Uanset stor lighed vil der altid være individuel variation.

Hvad figuren også viser, er at deltagerne akustisk set er mindre ens i oplæsningerne, og at denne uenshed konsekvent ligger på omkring 0,2 for alle talerpar i Hjørring og 0,4 for talerne i Jelling (med en enkelt undtagelse som jeg vil vende tilbage til). Dette tyder på at der generelt er mere individuel variation i måden talerne har oplæst ordene

fra listen end i deres udtale under kortøvelsen. I forlængelse af dette kan den konsekvente forskel mellem Jelling- og Hjørring-parrene tyde på at der er mindre individuel variation blandt Jelling-talerne end blandt Hjørring-talerne. Forskellen i resultaterne for oplæst tale og interaktion kan også tolkes som en genreforskel som afspejler at der sandsynligvis er mere reduktion i udtalen i løbet af interaktionen hvilket påvirker amplitudeudsvingene og dermed lighedsmålene.

Havde der været akkommodation som rækker ud over interaktionen, skulle den akustiske lighed være steget i oplæsningen efter kortøvelsen. Det gør den overordnet set ikke, med undtagelse af parret Line & Pernille, markeret med rød i Figur 6. Disse to taleres lighedsmål stiger fra ca. 0,2 før kortøvelsen til 0,35 efter kortøvelsen. Selvom ingen af tallene overstiger deres *match value* under kortøvelsen (ca. 0,55), så tyder dette på at de to talere har nærmet sig hinanden akustisk som følge af interaktionen.

Om ligheden mellem talerne stiger i løbet af interaktionen, fx fra kortsæt 1 til kortsæt 2 – en interaktion på gennemsnitligt ca. en halv time – har det desværre ikke været muligt at efterprøve eftersom antallet af datapunkter for hvert sæt var for få.

4.2 Resultater af AXB-analysen

Resultaterne fra AXB-studiet er trukket ud af systemet på soscisurvey. de og bearbejdet således at der kan laves statistik på dataene. For at overskueliggøre analysen er den opdelt i tre datagrupper. Den første gruppe sammenligner resultaterne for ligheden mellem en talers oplæste udtale før og efter kortøvelsen og samtalepartnerens kortudtale. Det er denne gruppe som er genstand for den kommende analyse. I denne datagrube måles det – som det var tilfældet i den forudgående analyse – om deltagerne er blevet påvirket af deres samtalepartner i løbet af kortøvelsen, og om dette i så fald har påvirket deres udtale efter endt øvelse. Her måles det altså hvor mange procent af gangene en udtale efter kortøvelsen (post) vurderes som tættest på partnerens udtale i løbet af øvelsen (task). Resultaterne kan ses i Tabel 1 og er baseret på 5078 unikke vurderinger foretaget af de 68 deltagere.

I tabellen findes også resultater for sammenligningen af udtalen hhv. før og efter øvelsen med talerens egen udtale under kortøvelsen, det som hedder *pre-task vs task* og *post-task vs task* i Tabel 1. Fra

et akkommodationsperspektiv er førstnævnte sammenligning mellem *pre* og *post* mest interessant fordi den viser om en eventuel påvirkning fra en partner fortsætter ud over interaktionens rammer, sådan som det tilsyneladende kun var tilfældet for ét talerpar i forudgående analyse. At en talers udtale i løbet af kortøvelsen er tættere på partnerens udtale i løbet af kortøvelsen end en af de to oplæsningsudtaler, er ikke overraskende; det peger nok nærmere på en genreforskel som også blev tydelig i resultaterne for amplitudekonvolutterne. Derfor er fokus her på de resultater der fremgår af første linje i Tabel 1.

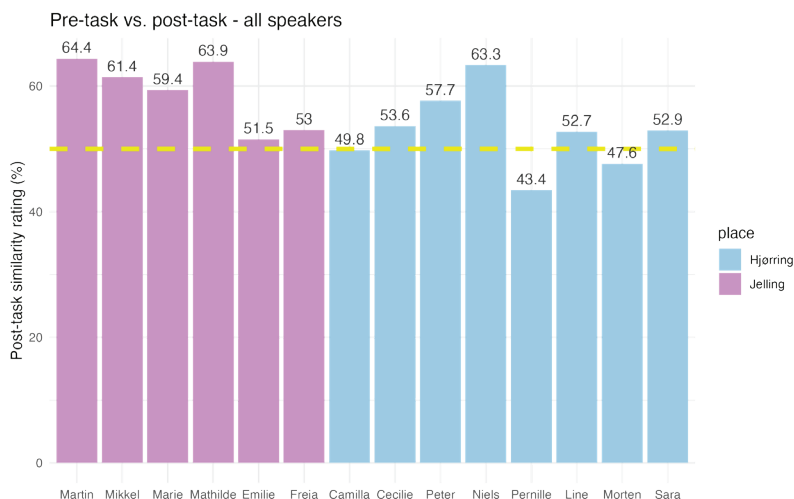
Af tabellen fremgår det at resultatet af lighedsvurderingerne er på 55 % for *pre-task* vs *post-task*. Dette betyder at i 55 % af tilfældene har lytterne vurderet at forekomsten fra oplæsningen *efter* kortøvelsen (*post*) lød mest som samtalepartnerens udtale i løbet af kortøvelsen, hvilket peger på en akkommodationseffekt. 55 % er ikke et overvældende resultat, men det stemmer godt overens med lignende undersøgelser i akkommodationslitteraturen (cf. Babel et al., 2014; Pardo, 2006; Pardo et al., 2013). Dataene er desuden *fittet* til lineære modeller for de tre forskellige datakombinationer som alle bekræfter at effekterne er signifikant større end tilfældet (50 %).

condition	Hjørring	Jelling	all
pre-task vs post-task	53 %	59 %	55 %
pre-task vs task	57 %	63 %	59 %
post-task vs task	57 %	63 %	59 %

Tabel 1: Overordnede resultater af AXB-analysen.

I første linje af Tabel 1 kan vi også se at deltagerne i AXB-undersøgelsen i en større andel af tilfældene har vurderet Jelling-deltagerens udtale efter endt øvelse som tættest på partners kortudtale. Dette antyder at Jelling deltagerne akkommoderer mere end Hjørring-deltagerne, eller at de er mere ens i det hele taget. Ser vi på de individuelle resultater for hver af talerne i Figur 7, bliver det klart at det tilsyneladende er to talerpar, Martin & Mikkel og Marie & Mathilde, der driver effekten hos Jelling-talerne, mens billedet er noget mere broget i Hjørring. Til gengæld viser resultaterne for Line & Pernille her ikke tegn på konvergens – måske nærmere det modsatte

– hvilket indikerer at den akkommodation der blev fundet i analysen af amplitudekonvolutterne, dækker over andre sproglige træk end dem deltagerne i perceptionsundersøgelsen har lyttet efter i deres vurdering af lydlighed mellem talerne.



Figur 7: Individuelle AXB-vurderingen for alle talere i pre-task vs. post-task.

5. Konklusion og fremtidige perspektiver

Denne artikel har introduceret to metoder til måling af akkommodation mellem to talere i interaktion; en akustisk og en perceptuel. Resultaterne viser både individuelle, kontekstuelle og regionale forskelle. Sammenholdt med hinanden viser resultaterne at lyttere finder konvergens (øget lighed) hvor det akustiske mål ikke gør; det eneste tegn på konvergens i det akustiske mål var hos parret Line & Pernille. Denne analyse viste til gengæld at oplæsning og interaktion må betragtes som forskellige genrer, også i denne henseende.

Den perceptuelle analyse viste en generel effekt af at blive udsat for en samtalepartners udtale, hvilket tyder på konvergens. Dette viser dog ikke om der er tale om en reel tilpasning til partneren eller blot et akustisk sammenfald hvor begge deltagere har fået talt sig varme og derfor udtaler *post-task*-ordlisten på en måde der er tættere på stilen brugt i interaktionen – og dermed opfattes af lytteren som mere ens.

Generelt viser resultaterne tydelige regionale forskelle som tyder

på at deltagerne i Jelling både er mere ens end Hjørring-talerne og tilpasser sig mest. Dette stemmer godt overens med indtrykket af at Hjørring-talerne generelt udviser større sproglig variation end Jelling-talerne og også kommer fra et større opland hvor flere talere er opvokset uden for Vendsyssel (blandt andre Line & Pernille). Jelling-talerne kommer alle fra det sydøstjyske område, med undtagelse af én taler som kommer fra Djursland. Det bekræfter desuden fund af Kim et al. (2011) som viser at kortere sproglig afstand mellem talere afføder mere sproglig tilpasning end lang sproglig afstand. Samtidig viser begge undersøgelser individuelle forskelle mellem talerne, og det bliver tydeligt at der er stor variation i hvem der tilpasser sig hvem og hvordan.

Resultaterne peger på at der stadig er lang vej igen for at forstå den dynamiske proces som en sproglig interaktion mellem to mennesker udgør. Undersøgelser på forskellige niveauer finder forskellige resultater, og meget tyder på at lyttere bruger andre (eller mange samtidige) faktorer til at vurdere lydlig lighed.

Litteratur

- Abel, Jennifer & Babel, Molly (2017). Cognitive Load Reduces Perceived Linguistic Convergence Between Dyads. *Language and Speech*, 60(3), 479-502. <https://doi.org/10.1177/0023830916665652>
- Aguilar, Lauren, Downey, Geraldine, Krauss, Robert, Pardo, Jennifer, Lane, Sean & Bolger, Niall (2016). A Dyadic Perspective on Speech Accommodation and Social Connection: Both Partners' Rejection Sensitivity Matters. *Journal of Personality*, 84(2), 165-177. <https://doi.org/10.1111/jopy.12149>
- Anderson, Anne H., Bader, Miles, Bard, Ellen Gurman, Boyle, Elizabeth, Doherty, Gwyneth, Garrod, Simon, Isard, Stephen, Kowtko, Jacqueline, McAllister, Jan, Miller, Jim, Sotillo, Catherine, Thompson, Henry S. & Weinert, Regina (1991). The HCRC Map Task Corpus. *Language and Speech*, 34(4), 351-366.
- Babel, Molly (2010). Dialect divergence and convergence in New Zealand English. *Language in Society*, 39(4), 437-456.
- Babel, Molly (2012). Evidence for phonetic and social selectivity in spontaneous phonetic imitation. *Journal of Phonetics*, 40(1), 177-189. <https://doi.org/10.1016/j.wocn.2011.09.001>

- Babel, Molly & Bulatov, Dasha (2012). The Role of Fundamental Frequency in Phonetic Accommodation. *Language and Speech*, 55(2), 231-248. <https://doi.org/10.1177/0023830911417695>
- Babel, Molly, McGuire, Grant, Walters, Sophia & Nicholls, Alice (2014). Novelty and social preference in phonetic accommodation. *Laboratory Phonology*, 5(1), 123-150. <https://doi.org/10.1515/lp-2014-0006>
- Cohen Priva, Uriel & Sanker, Chelsea (2020). Natural Leaders: Some Interlocutors Elicit Greater Convergence Across Conversations and Across Characteristics. *Cognitive Science*, 44(10), 1-34. <https://doi.org/10.1111/cogs.12897>
- Goldinger, Stephen D. (1998). Echoes of echoes? An episodic theory of lexical access. *Psychological Review*, 105(2), 251-279. <https://doi.org/10.1037/0033-295X.105.2.251>
- Greenaway, Ruth Elizabeth (2017). Chapter 13—ABX Discrimination Task. I L. Rogers (Red.), *Discrimination Testing in Sensory Science* (s. 267-288). Woodhead Publishing. <https://doi.org/10.1016/B978-0-08-101009-9.00013-7>
- Greenwald, Anthony G., Nosek, Brian A. & Banaji, Mahzarin R. (2003). Understanding and using the Implicit Association Test: I. An improved scoring algorithm. *Journal of Personality and Social Psychology*, 85(2), 197-216. <https://doi.org/10.1037/0022-3514.85.2.197>
- Grønnum, Nina (2022). *DanPASS - Danish Phonetically Annotated Spontaneous Speech [Web page]*. <https://danpass.hum.ku.dk/>
- He, Lei & Dellwo, Volker (2016). A Praat-Based Algorithm to Extract the Amplitude Envelope and Temporal Fine Structure Using the Hilbert Transform. *Interspeech 2016*, 530-534. <https://doi.org/10.21437/Interspeech.2016-1447>
- Hilbert, David (1912). *Grundzüge einer allgemeinen Theorie der linearen Integralgleichungen*. B. G. Teubner.
- Kim, Midam, Horton, William S. & Bradlow, Ann R. (2011). Phonetic convergence in spontaneous conversations as a function of interlocutor language distance. *Laboratory Phonology*, 2(1). <https://doi.org/10.1515/labphon.2011.004>
- Lewandowski, Natalie (2012). *Talent in nonnative phonetic convergence* [ph.d.-afhandling]. Universität Stuttgart. <http://elib.uni-stuttgart.de/handle/11682/2875>
- Lewandowski, Natalie & Jilka, Matthias (2019). Phonetic Convergence, Language Talent, Personality and Attention. *Frontiers in Communication*, 4, 1-19. <https://doi.org/10.3389/fcomm.2019.00018>
- MacLeod, Bethany (2013). *What is the default behaviour in accommodation: Convergence or maintenance?* 1-5. <https://doi.org/10.1121/1.4800626>
- Moore, Brian C. J. (2008). The Role of Temporal Fine Structure Processing in Pitch Perception, Masking, and Speech Perception for Normal-Hearing and Hearing-Impaired People. *Journal of the Association for Research in Otolaryngology*, 9(4), 399-406. <https://doi.org/10.1007/s10162-008-0143-x>
- Namy, Laura L., Nygaard, Lynne C. & Sauerteig, Denise (2002). Gender Differences in Vocal Accommodation: The Role of Perception. *Journal of Language and Social Psychology*, 21(4), 422-432. <https://doi.org/10.1177/026192702237958>

- Pardo, Jennifer S. (2006). On phonetic convergence during conversational interaction. *The Journal of the Acoustical Society of America*. <https://doi.org/10.1121/1.2178720>
- Pardo, Jennifer S., Jordan, Kelly, Mallari, Rolliene, Scanlon, Caitlin & Lewandowski, Eva (2013). Phonetic convergence in shadowed speech: The relation between acoustic and perceptual measures. *Journal of Memory and Language*, 69(3), 183-195. <https://doi.org/10.1016/j.jml.2013.06.002>
- Pardo, Jennifer S., Urmache, Adelya, Gash, Hannah, Wiener, Jaclyn, Mason, Nicholas, Wilman, Sherilyn, Francis, Keagan & Decker, Alexa (2019). The Montclair Map Task: Balance, Efficacy, and Efficiency in Conversational Interaction. *Language and Speech*, 62(2), 378-398. <https://doi.org/10.1177/0023830918775435>
- Pardo, Jennifer S., Urmache, Adelya, Wilman, Sherilyn, & Wiener, Jaclyn (2017). Phonetic convergence across multiple measures and model talkers. *Attention, Perception, & Psychophysics*, 79(2), 637-659. <https://doi.org/10.3758/s13414-016-1226-0>
- Rathje, Marianne (2008). *Generationssprog i mundtlig interaktion. En sociolingvistisk undersøgelse af generationsspecifikke sproglige og interaktionelle træk i tre generationers talesprog* [ph.d.-afhandling]. Det Humanistiske Fakultet, Københavns Universitet.
- Rathje, Marianne, Hougaard, Tina Thode & Jensen, Eva Skafte (2021). E-mailhilsner: Relation, variation og akkommodation. *Globe: A Journal of Language, Culture and Communication*, 13(2021), 47-85.
- Sonderegger, Morgan, Bane, Max & Graff, Peter (2017). The Medium-Term Dynamics of Accents on Reality Television. *Language*, 93(3), 598-640.
- Tilsen, Sam (2009). Subphonemic and cross-phonemic priming in vowel shadowing: Evidence for the involvement of exemplars in production. *Journal of Phonetics*, 37(3), 276-296. <https://doi.org/10.1016/j.wocn.2009.03.004>
- Wade, Travis, Dogil, Grzegorz, Schütze, Hinrich, Walsh, Michael & Möbius, Bernd (2010). Syllable frequency effects in a context-sensitive segment production model. *Journal of Phonetics*, 38(2), 227-239. <https://doi.org/10.1016/j.wocn.2009.10.004>
- Yu, Alan C. L., Abrego-Collier, Carissa, & Sonderegger, Morgan (2013). Phonetic Imitation from an Individual-Difference Perspective: Subjective Attitude, Personality and “Autistic” Traits. *PLoS ONE*, 8(9), 1-13. <https://doi.org/10.1371/journal.pone.0074746>