

EN STOR SPROGTEKNOLOGISK ORDBOG FOR DANSK

- med særligt fokus på håndtering af flertydighed i en niveaudelt ordbog

Anna Braasch og Bolette Sandford Pedersen (Center for Sprogteknologi)
e-mail: anna@cst.ku.dk, bolette@cst.ku.dk

0 STO-projektet

Center for Sprogteknologi har taget initiativ til et omfattende projekt hvis formål er at skabe en stor dansk sprogteknologisk ordbog (STO). En sådan ordbog er der brug for i datalingvistisk forskning (f.eks. i systemer til sproglig analyse) og sprogteknologiske anvendelser (f.eks. stavekontrol, maskinoversættelse, taleteknologi). Det er meget tidkrævende og dermed omkostningstungt at udføre det planlagte arbejde. Den færdige ordbog skal derfor være et basisprodukt mht. leksikalsk dækning og lingvistisk beskrivelse; den skal kunne opfylde mange forskellige slags behov ved at kunne tilpasses forskelligartede formål. På denne måde vil vi sikre at ordbogen får et bredt anvendelsespotentiale både i forskningen og i sprogindustrien. Samtidig er det vigtigt at fastholde, at netop på grund af dette meget brede sigte skal ordbogens struktur og oplysningsindhold tilgodese meget forskellige lingvistiske og leksikografiske krav; dette medfører at vi må træffe nogle valg som er anderledes end ved traditionelle ordbøger.

1 Ordbogens faglige grundlag

Projektets faglige udgangspunkt er den danske PAROLE-'ordbog' både hvad metode, udarbejdet materiale og erfaringer angår. Denne ordbog, som er en leksikalsk datasamling i elektronisk form lagret i en database, er udarbejdet i det EU-støttede PAROLE (LE2-4017) projekt, hvis formål var at skabe ensartede sproglige ressourcer (ordbøger og korpora) for alle 12 officielle EU-sprog. Den danske ordbog er udarbejdet af CST i samarbejde med Det Danske Sprog- og Litteraturselskab i perioden 1995-98 og den indeholder ca. 20.000 opslagsord forsynet med morfologiske og syntaktiske beskrivelser. I det efterfølgende SIMPLE (LE4-8346) projekt bliver ca. halvdelen af dette ordforråd forsynet med semantiske beskrivelser.

PAROLE-ordbogens størrelse er imidlertid utilstrækkelig både til realistiske sprogteknologiske applikationer og for et ordbogsmodul i forskningsrettede datamatiske applikationer; men materialet er velegnet til at blive udbygget betydeligt. Dette sker i STO-projektet; dets mål er at udbygge ordforrådet til ca. 55.000 opslagsord, og dermed bliver antallet af semantiske læsninger ca. 100.000¹.

I PAROLE-projektet er der anvendt en generel og såkaldt teorineutral beskrivelsesmetode: på den ene side skulle metoden være velegnet til så forskellige sprog som græsk og dansk, på den anden side skulle den ikke baseres på en bestemt slags lingvistisk retning (f.eks. HPSG, LFG, GPSG). Metoden danner basis for en bredt anvendelig, såkaldt multifunktionel ordbog. Alle tolv PAROLE-ordbøger er opbygget efter samme beskrivelsesmodel og er harmoniseret mht. oplysningstyper og deres repræsentation. Denne model er lingvistisk set meget detaljeret og med hensyn til strukturen ganske kompliceret, fordi den skal tilgodese en bred vifte af forskellige sprog og applikationer. Vi kan forenkle modellen en del i STO da den jo kun skal omfatte ét sprog. Samtidig foretager vi nogle justeringer med henblik på en bedre beskrivelse af specifikke sproglige fænomener som for eksempel partikelverber. Når vi foretager ændringer i modellen, bestræber vi os dog på at opretholde kompatibiliteten med PAROLE, så STO eventuelt senere kan sammenkobles med andre etsprogsordbøger der bliver udarbejdet efter lignende principper hos andre sproggrupper.

2 Den niveaudelte beskrivelsesstruktur

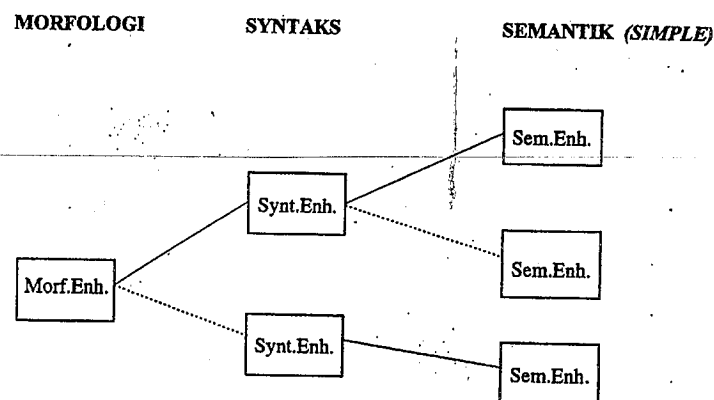
Makrostrukturen i den anvendte model afviger fra den der anvendes i traditionelle ordbøger: sidstnævnte består af sammenhængende artikler, dvs. opslagsord og den dertil hørende, for ordbogen udtømmende, leksikografiske beskrivelse af opslagsordets morfologiske, syntaktiske og semantiske egenskaber.

Informationsindholdet er i stedet opdelt i oplysningstyper som struktureres på tre adskilte *beskrivelsesniveauer* (jf. Figur 1) svarende til morfologiske, syntaktiske og semantiske egenskaber; disse beskrivelser udgør ordbogens morfologiske, syntaktiske og semantiske enheder. Enhederne kædes sammen på forskellig vis og derved produceres en *niveaudelt* ordbog. Denne beskrivelsesmodel sikrer en særlig fleksibilitet for modulær datamatisk

¹ For mere information om STO-projektet, se også Braasch, Mægaard & Pedersen (1998) samt Braasch, Christensen, Olsen & Pedersen (1998).

håndtering af data (f.eks. opdatering og konvertering) og for sprogteknologiske anvendelser, men den producerede ordbog er ikke umiddelbart læselig for en menneskelig bruger. Datamatiske anvendelser, som er ordbogens primære anvendelsesområde, stiller særlige krav om en formaliseret, systematisk, eksplicit og præcis lingvistisk beskrivelse; derfor vil det producerede materiale efter konvertering (f.eks. til et traditionelt ordbogsformat) også kunne anvendes til andre leksikografiske og lingvistiske formål.

Figur 1: De tre beskrivelsesniveauer



2.1 Den niveaudelte ordbog og dens indhold (oplysningstyper)

En morfologisk enhed giver mindst oplysning om opslagsordets stavning, bøjning, ordklasse, og køn; en syntaktisk enhed indeholder oplysninger om ordets konstruktionspotentiale (funktionel og kategoriel valens mm.) og syntaktiske funktion, brug af hjælpeverbum osv. Endelig indeholder en semantisk enhed oplysninger om semantisk klasse, tematiske roller, selektionsrestriktioner, domæne mm.

Ved en sådan formalistisk opdeling i adskilte beskrivelsesenheder opstår nogle fundamentale lingvistiske spørgsmål, f.eks. hvor skal skellet mellem beskrivelsesniveauer lægges? Morfologi og syntaks har et overlappende område, morfosyntaksen, der omfatter eksempelvis ordklasse og køn. Disse oplysninger er relevante - og er derfor tilstede - på begge beskrivelsesniveauer.

Et andet spørgsmål: hvor detaljerede bør de prototypiske mønstre, f.eks. valensmønstre, være? På den ene side har vi kravet om lingvistisk præcision, hvilket vil sige høj detaljeringsgrad i beskrivelsen, og det resulterer i opsplitning i flere syntaktiske enheder: jo højere detaljeringsgrad, jo flere enheder. Dette medfører på den anden side flere mulige analyser af sætninger (især hvor optionelle led er udeladt), og det resulterer i en u hensigtsmæssig overgenerering i datamatiske anvendelser. En sådan flertydighed kan først elimineres ved at inddrage semantiske oplysninger, dette afspejler syntaksens og semantikens tætte forbindelse, men det svækker også ordbogens modulære anvendelsesmuligheder.

Det er generelt et spørgsmål om strategi hvordan man udnytter den niveaudelte model. Det gælder både hvordan man tilordner oplysninger til et bestemt niveau, hvilken detaljeringsgrad man vælger, og hvordan man opdeler et opslagsords beskrivelse i niveaubestemte enheder på grundlag af formaliserbare lingvistiske skel.

2.2 Kodningsstrategi

Vi har valgt en grundlæggende strategi: der er tale om én enhed så længe der ikke registreres forskellige værdier for et bestemt egenskab på det pågældende beskrivelsesniveau. Denne strategi kalder vi 'split så sent som muligt'. Eksempelvis håndteres to homonymer, der har ens morfologiske egenskaber, som én morfologisk enhed, men hvis der er forskel i deres syntaktiske distribution eller valensmønstre, så splittes de op i to eller flere syntaktiske enheder. Hvis der først på semantisk niveau er en formaliserbar forskel, splittes der ikke på syntaktisk niveau, men først i semantikken. Strategien har således den fordel at flertydigheder først introduceres på det niveau hvor de kan håndteres ved formel analyse; derved undgås uønsket flertydighed og overgenerering. I afsnit 3 gives et par detaljerede eksempler på de problemstillinger der er forbundet med denne strategi.

3 Homonymi og polysemi i den niveaudelte ordbog

Den første problemstilling vi vil tage fat på for at illustrere hvad niveaudeling af leksikalsk information betyder for strukturen i ordbogen, er spørgsmålet om repræsentation af homonymi og polysemi.

Et standardkriterium for at skelne mellem homonymi og polysemi er *etymologi* (Lyons 1977:550-569): hvis to betydninger af et ord har samme forhistorie, altså stammer fra det

samme ord, taler vi om polysemer; ellers om homografer. Et andet, knap så stringent kriterium er sammenhængen mellem to betydninger: hvis to betydninger er relateret til hinanden, kalder vi dem polysemer; hvis ikke, er der tale om homografer.

I de fleste traditionelle ordbøger anvendes det sproghistoriske kriterium, og vi finder som oftest en klar begrebsmæssig (og grafisk) distinktion mellem homografer og polysemer². Homografer har hver deres opslagsord, ofte markeret med romertal som illustreret nedenfor (jf. f.eks. Nudansk Ordbog 14. udg.):

(1) *Homografer i den traditionelle ordbog*

pande I (i betydningen 'hoveddel')
pande II (i betydningen 'køkkenredskab')

Til forskel herfra udtrykkes polysemer som forskellige betydninger af det samme opslagsord:

(2) *Polysemer i den traditionelle ordbog*

afstemning

1. som i 'nøje afstemning af farverne'
2. som i 'holde afstemning om et politisk spørgsmål'

Ofte ses også en yderligere underinddeling af relaterede betydninger, f.eks. placeres overførte betydninger eller underbetydninger af et ord ofte som en underenhed til en specifik betydning.

I den niveaudelte, sprogteknologiske ordbog medtages de sproghistoriske forhold ikke idet de skønnes ikke at være relevante for den datamatiske anvendelse. Her arbejdes som nævnt med strategien 'split så sent som muligt' og derfor vil homografen 'pande' kun repræsenteres med én enhed på hhv. det morfologiske og syntaktiske niveau, hvor der ikke kan registreres nogen forskel; med andre ord: de to ord bøjes ens og har samme syntaktiske distribution. På det semantiske niveau derimod må skellet nødvendigvis indføres idet der refereres til to meget forskellige entiteter med meget forskellig struktur ('pande' i betydningen 'hoveddel' er f.eks. naturskabt, hvorimod 'pande' i betydningen 'køkkenredskab' er et menneskeskabt redskab udviklet til at tilberede mad på). Lidt anderledes forholder det sig med homografen 'love' i betydningerne 'forpligte sig til' og 'prise'. Disse to betydninger har forskellig syntaktisk distribution som det ses i hhv. 'han lovede mig at komme' og 'Gud være lovet' og derfor må der splittes i to enheder allerede på syntaksniveau:

² En undtagelse er Retskrivningsordbogen som er opbygget efter andre principper.

(3) *homografer i den niveaudelte ordbog*

MORFOLOGI	SYNTAKS	SEMANTIK
pande	pande	pande
		pande
love	love	love
	love	love

Hvis vi ser på håndteringen af polysemer i den niveaudelte ordbog, ser vi at de tilsyneladende beskrives helt parallelt til homograferne. Igen fastholdes strategien om at opsplitte i forskellige enheder så sent som muligt i analysen; 'mus' i betydningerne 'en computermus' og en 'levende mus' splittes først på semantikniveau hvorimod 'afstemning' splittes allerede på syntaks, idet der er forskellig syntaktisk distribution ('afstemning af farverne' vs. 'afstemning om et politisk emne').

(4) *polysemer i den niveaudelte ordbog*

MORFOLOGI	SYNTAKS	SEMANTIK
mus	mus	mus
		mus
afstemning	afstemning	afstemning
	afstemning	afstemning

Spørgsmålet er så om det i den semantiske beskrivelse implicit eller eksplicit vil blive udtrykt om to betydninger er semantisk relaterede eller ej. For *STO*'s vedkommende er det endnu ikke klarlagt præcis hvordan den semantiske repræsentation skal se ud, og vi må derfor sige at vi endnu ikke har fuld klarhed over hvorvidt dette vil være tilfældet. Vi eksperimenterer i øjeblikket med at anvende Pustejovskys (1995) 'Generative Leksikon', hvor der for hver semantisk enhed indkodes en qualiastruktur der angiver forskellige dimensioner af ordets betydning, såsom dets tilhørsforhold i et semantisk hierarki (f.eks. human vs. ikke-human), dets interne struktur (er der f.eks. tale om en del-helhedsrelation), dets funktion (f.eks. instrumentfunktion) og dets tilblivelse (f.eks. naturskabt eller menneskeskabt³). Med en sådan semantisk beskrivelse er det overvejende sandsynligt at polysemer vil få en semantisk

³ Eksperimentet foregår inden for SIMPLE-projektet, som er et EU-projekt med deltagelse af en lang række europæiske lande. Hvert land har til opgave at etablere 10.000 semantisk kodede enheder på basis af Pustejovskys qualiastruktur. De 10.000 enheder skal alle være indeholdt i *PAROLE*-ordbogen. SIMPLE kan således ses som en videreudvikling af *PAROLE*-projektet.

beskrivelse der har mange fælles træk, hvorimod homografen 'pande' vil få tilskrevet to meget forskellige semantiske beskrivelser. Dog må det anses for usandsynligt at vi får en så fintmasket semantik at vi i den semantiske beskrivelse kan se at de to betydninger af 'mus' er relaterede. At en computermus ligner en almindelig mus og har fået navnet deraf, vil næppe være et semantisk træk som vil vægtes i den formelle beskrivelse.

Imidlertid er der en lang række tilfælde af *systematisk* polysemi hvor relationen vil blive givet *explicit* ved hjælp af et 'link' mellem to semantiske enheder. Systematisk polysemi dækker over de tilfælde hvor grupper af ord kan undergå de samme betydningsændringer. En foreløbig liste over sådanne tilfælde af systematisk polysemi i dansk gives nedenfor:

(5) klasser af systematisk polysemi

beholder/kvantitet	'kop'
dyr/mad	'lam'
genstand/åbning	'dør'
handling/resultat	'bygning'
sted/folk	'Danmark'

Hvis vi ser bort fra den systematiske polysemi, må vi konkludere at skellet mellem homonymi og polysemi delvis udviskes i den niveaudelte ordbog. Vi skønner med andre ord at hvis der alligevel ikke er tale om betydningsrelationer der kan udnyttes formelt, så er det heller ikke en oplysning vi har brug for at repræsentere formelt (dette skønner vi f.eks. er tilfældet med relationen mellem de to betydninger af 'mus'). Dette betyder dog ikke at oplysningen ikke kan være relevant for den menneskelige bruger af ordbogen. Med den menneskelige bruger mener vi de folk der på et givet tidspunkt skal sætte sig ind i ordbogens struktur, enten fordi de skal anvende den til en bestemt applikation, eller fordi de evt. skal udvide ordbogens dækningsgrad. Derfor har vi besluttet at indføre oplysningen om homonymi vs. polysemi som en *brugeroplysning*, således at homograferne i et brugersfelt markeres med romertal som det ses nedenfor med homografen 'pande':

(6) brugeroplysning ved homografer

MORFOLOGI	SYNTAKS	SEMANTIK
pande I, II	pande I, II	pande I pande II

4 Partikelverber og partikelkonstruktioner i den niveaudelte ordbog

Den anden lingvistiske problemstilling som vi arbejder meget med i øjeblikket, og som også er meget illustrativ for den niveaudelte ordbog, er håndteringen af konstruktioner med partikler. Et væsentligt skel som man er nødt til at foretage - og som man også er nødt til at foretage i traditionelle leksika, omend der ikke er tale om nogen meget klar grænse - er en distinktion imellem det vi vælger at kalde partikelverber, hvor vi ønsker at give en leksikaliseret fortolkning, og det vi kalder partikelkonstruktioner, hvor vi snarere ønsker at give en grammatisk fortolkning ved at opfatte partiklen som en del af verbets valensmønster.

Hovedkriteriet for sådan et skel er traditionelt *gennemskuelighed*. Hvis enten verbets betydning eller partiklens betydning er uigennemskuelig - altså afviger fra grundbetydningen i verbet alene eller i partiklen alene - så må vi vælge den leksikaliserede løsning, som i tilfældet 'vaske op'⁴. Hvis betydningen derimod er nogenlunde forudsigelig ud fra verbets og partiklens grundbetydninger - og her er der altså tale om et skøn - så foretrækker vi valensfortolkningen, som f.eks. i tilfældet 'grave noget op/ned'. Et andet kriterium, som kan være relevant at anvende, er hyppighed. Nogle verber optræder så hyppigt med en bestemt partikel at det synes rimeligt alene på baggrund heraf at opfatte det som en leksikaliseret størrelse. Som nævnt er dette skel også relevant for den traditionelle ordbog, idet man må beslutte om en konstruktion med partikel skal medtages som et sublemma til hovedartiklen eller beskrives via en valensbeskrivelse (f.eks. 'gå + RETNING').

I den niveaudelte ordbog arbejder vi i øjeblikket med to forskellige repræsentationsmodeller der etablerer det samme skel på forskellig vis. Målet med de to alternative modeller er at skabe et erfaringsgrundlag der gør det muligt på et lidt senere tidspunkt at vælge én forhåbentlig optimal løsning. I den ene model etableres skellet imellem hhv. partikelkonstruktion og partikelverbum på syntaksniveau:

(7) Model 1: skellet skabes på syntaksniveau

MORFOLOGI	SYNTAKS	SEMANTIK
grave	grave (+part)	grave (+retn)
vaske	vaske op	vaske op

⁴ Leksikalisering er bl.a. nødvendig i forbindelse med maskinoversættelse, idet vi ikke ønsker 'vaske op' oversat kompositionelt til f.eks. engelsk (*wash up).

Vi har stadig ønsket om at splitte så sent som muligt, og det vil sige at alt hvad der formelt kan tolkes som flertydighed helst etableres så sent som muligt. Således lader vi morfologiniveau være uberørt af en mulig leksikalisering og fastholder kun én morfologisk enhed for hvert verbum. I modellen ovenfor splitter vi til gengæld på syntaksniveau: ved partikelkonstruktionen fortolker vi partiklen som en del af valensmønstret, og partiklen angives som en option i valensfeltet. Ved partikelverbet derimod udskiller vi et nyt leksem 'vaske op' og beskriver dette i en syntaktisk enhed for sig.

I nogle tilfælde skaber denne model imidlertid en flertydighed som egentlig ikke burde være nødvendig at etablere allerede på syntaksniveau; som i tilfældet 'gå op' hvor begge fortolkninger er mulige; både en gennemskuelig og en ikke-gennemskuelig ('gå i retningen opad' vs. 'regnestykket går op'):

(8) en semantisk flertydighed etableres på syntaksniveau

MORFOLOGI	SYNTAKS	SEMANTIK
gå	gå (+part)	gå (+retn)
	gå op	gå op

Sådanne eksempler kan bruges som argumentation for først at etablere skellet på semantikniveau, idet der jo dybest set er tale om en semantisk betinget forskel. En sådan model illustreres nedenfor:

(9) Model 2: Skellet skabes på semantikniveau

MORFOLOGI	SYNTAKS	SEMANTIK
vaske	vaske (+part)	vaske vaske op
grave	grave (+part)	grave (+retn)
gå	gå (+part)	gå (+retn) gå op

Umiddelbart virker denne model tiltalende og bedst i tråd med vores strategi om at splitte så sent som muligt; på den anden side forekommer det ikke helt intuitivt korrekt ud fra et syntaktisk synspunkt at opfatte partiklen 'op' i 'vaske op' som optionel til 'vaske'.

5 Konklusion

Vi har nævnt flere årsager til hvorfor det er hensigtsmæssigt at etablere en niveaudelt ordbog bestående af hhv. morfologiske, syntaktiske og semantiske enheder i *STO*. Først og fremmest har vi skullet sikre at ordbogen kan anvendes til forskellige formål. I og med at et af de største problemer inden for datamatisk sprogbehandling udgøres af *flertydighedsproblematikken*, er det indlysende at netop denne problematik må have en stor indflydelse på ordbogens struktur. Med den niveaudelte ordbog sikrer vi at en applikation der kun skal anvende de morfologiske informationer i ordbogen (f.eks. stavekontrol), helt undgår at støde ind i den flertydighedsproblematik som udgøres af hhv. polysemer, homografer og partikelverber. Ligeledes undgår vi at en applikation der skal anvende ordbogen til grammatikkontrol, skal håndtere semantisk flertydighed og forholde sig til to betydninger af 'pande'. Ved maskinoversættelse og natur-sprogsgrænseflader er det sandsynligt at hele ordbogen inklusive dens semantik skal tages i anvendelse, men også her vil det ofte være hensigtsmæssigt med en modular opbygning. Kort sagt: hvis vi ønsker en applikationsneutral ordbog, er det nødvendigt med en strengt modular opbygning af de leksikalske oplysninger.

Dette krav står imidlertid, som vi har vist i artiklen her, af og til i modsætning til lingvistiske fænomeneres beskaffenhed. Det er ikke altid let at afgøre hvad der er hhv. morfologisk, syntaktisk og semantisk information, idet niveauerne på mange planer er sammenflettede eller tæt forbundne. Specielt partikelverberne viste sig her at ligge i et grænseområde imellem niveauerne.

Litteratur

- Braasch, A., B. Maegaard, B.S. Pedersen. 1998. *En stor dansk sprogteknologisk ordbog - et nationalt projekt. Datalingvistisk Årsmøde 1997*. Southern Denmark Business School.
- Braasch, A., A. B. Christensen, S. Olsen & B.S. Pedersen. 1998. A Large-Scale Lexicon for Danish in the Information Society. *Proceedings from the First International Conference on Language Resources and Evaluation*, Granada 1998.
- Lyons, J. 1977. *Semantics*. Cambridge University Press.
- Pustejovsky, J. 1995. *The Generative Lexicon*. The MIT Press, Cambridge, MA, London.

Slutbemærkning

CST har taget initiativ til dette projekt, har udført en stor del af forarbejdet og forventer også at udføre en stor del af arbejdet, herunder styringen af projektet. For at kunne drage nytte af

den forskning og ekspertise der findes forskellige steder i det faglige miljø, lægges der stor vægt på et bredt samarbejde. Vi håber på tilbagemeldinger fra interesserede kolleger. Det kunne bidrage til at skabe værdifulde kontakter til mulige samarbejdspartnere.

Peter Widell og Mette Kunøe (udg.):
7. Møde om Udforskningen af Dansk Sprog
Århus 1999

LEKSIKALSK LANGTIDSAKKOMMODATION BLANDT DANSKERE I NORGE¹

Af Randi Benedikte Brodersen (Universitetet i Bergen)

1 Indledning

Denne artikel indeholder følgende afsnit: 1. Indledning, 2. Teoretisk grundlag, 3. Mål, materiale og metoder, 4. Leksikalsk akkommodation i 11 danske informanternes talesprog, 5. Informanternes sproglige praksis, idealer og holdninger, 6. Sproginterne og sprogksterne faktorer i informanternes akkommodation og 7. Afsluttende bemærkninger.

Emnet er 11 danskeres leksikalske akkommodation – eller tilpasning – til norsk. Disse danskere var mine informanter² i foråret og sommeren 1998 og blev da interviewet af en norsk interviewer i et pilotprojekt som er en del af mit *dr.art.*-projekt ved Universitetet i Bergen om *Sproglig variation, sproglig tilpasning og sprogholdninger blandt danskere i Norge.*

Leksikalsk akkommodation er en sprogbrugers mere eller mindre bevidste eller ubevidste sproglige modificering i forhold til en anden sprogbrugers sprog og en kommunikationsstrategi som en sprogbruger evt. kombinerer med andre sproglige og ikke-sproglige kommunikationsstrategier.

Den form for akkommodation der er aktuel i mit projekt, er sproglig langtidsakkommodation (efter Trudgill 1986:3) blandt danskere som har boet i Norge i en årrække (dvs. i mindst 3 år). I fokus er dermed danskeres sproglige variation og tilpasninger til norsk som resultat af deres direkte kontakt med norske sprogbrugere under langvarige eller permanente ophold i Norge.

Den leksikalske akkommodation til norsk blandt danskere i Norge er interessant af flere grunde; bl.a. fordi den ikke har været undersøgt tidligere og fordi det er den tidligste, mest markante og mest udbredte form for sproglig tilpasning blandt mine informanter og andre

¹ Denne artikel er en revideret version af mit foredrag på 7. Møde om Udforskningen af Dansk Sprog. En udvidet version er Brodersen 1998.

² Tak til hver enkelt informant for alle ord og alle svar, og tak til min norske interviewer for at have stillet alle spørgsmål.