

NyS 65

NYDANSKE SPROGSTUDIER

NyS 65
NYDANSKE SPROGSTUDIER

© NyS og artiklernes forfattere, 2024

Redaktion:

Jonas Nygaard Blom, Jan Heegård, Tina Thode Hougaard, Alexandra Kratschmer,
Jessie Leigh Nielsen, Michael Nguyen, Jacob Thøgersen

Oplysninger om indsendelse af bidrag findes bagest i nummeret.

Udgives i samarbejde med

Dansk Sprognævn

Adelgade 119 B

5400 Bogense

Grafisk opsætning: René Malchow

ISSN 2246-4522 (online)

Indhold

Forord	5
BOLETTE SANDFORD PEDERSEN, NATHALIE C. HAU SØRENSEN, SUSSI OLSEN & SANNI NIMB Evaluerings af sprogforståelsen i danske sprogmodeller – med udgangspunkt i semantiske ordbøger	8
EA LINDHARDT OVERGAARD & ULF DALVAD BERTHELTSEN Kan AI reproducere fagdisciplinær stemme? Et komparativt korpusbaseret studie af GPT4's evne til at reproducere fagdisciplinær stemme i AI-genereret sprogvidenskabelig prosa	41
JONAS NYGAARD BLOM, ALEXANDRA HOLSTING & JESPER TINGGAARD SVENDSEN På sporet af chatbottens sproglige fingeraftryk En sproglig ophavsanalyse af tekster skrevet af danskstuderende og ChatGPT	79
TINA THODE HOUGAARD Emojier som parasproglige resurser i skrevet onlineinteraktion	111
IB ULBÆK Anmeldelse af Peter Widell: <i>Sproget i verden. Tekstforståelsens filosofiske grundlag</i> . Aarhus Universitetsforlag, 2023. (355 sider)	145
HOLGER JUUL Bogstav-lyd-forbindelser i dansk: The Ultimate Collection Anmeldelse af Christian Becker-Christensen: <i>Dansk Retskrivning: Bogstav – Lyd – Bogstav. En grafemisk beskrivelse af dansk ortografi</i> . Syddansk Universitetsforlag, 2021. (416 sider)	160
Engelske abstracts	171
NyS 67 – Temanummer Sproglige perspektiver på klima og vejr	175

Forord

2023 var året hvor kunstig intelligens (AI) og *large language models* (LLM) for alvor satte deres aftryk på almindelige sprogbrugeres liv og i særdeleshed på forskning og undervisning. Chatbotten ChatGPT genererede avisoverskrifter (i begge betydninger) og truede med (afhængig af standpunkt) at revolutionere eller undergrave forskning, undervisning og eksamination. Pludselig blev vi opmærksomme på at det at skrive en sammenhængende og forståelig tekst ikke krævede et menneske endsige en lang uddannelse.

Nu, i 2024, er vi så småt ved at vænne os til den nye verdensorden, og vi bevæger os fra spekulation til levede erfaringer. LLM'er er ikke længere en hypotetisk *disruption* af undervisning eller arbejdsmarked. LLM'er har fundet vej ind i undervisningslokaler fra folkeskole til universitet. Vi diskuterer nu ikke længere hvordan vi skal dæmme op for snyd med hjælp fra AI, men hvordan studerende skal undervises i at bruge LLM'er i lærings- og eksamenssammenhæng. Vi taler mindre om hvordan LLM'er vil overtage humanisters arbejdsfunktioner, og mere om hvordan vi træner humanister til en fremtid hvor de skal arbejde tæt sammen med AI.

Vi begynder måske også at få øjnene op for de mere interessante og foruroligende spørgsmål som AI og LLM'er kaster af sig. Hvad siger LLM'er fx om menneskets sprogevne og menneskets status som den eneste art der bruger sprog? Med GPT4 er vi vidne til at sprogrelaterede opgaver nu bliver løst på en måde der ofte ikke længere kan skelnes fra hvad der er produceret af mennesker. Vi begynder at nærme os drømmen om *Artificial General Intelligence* – maskinen der kan løse en lang række kognitive opgaver på niveau med et menneske. Hvis maskinen løser opgaver tilsyneladende som et menneske, er det så forkert at sige at maskinen 'tænker'? Eller omvendt: Hvis maskinen ikke tænker, hvorfor siger vi så at mennesker 'tænker'? Vi vil ikke gå for langt ind i denne filosofiske debat, som vækker minder om tidligere tiders debatter om behaviorisme vs. kognitivismen. Men vi peger på den fordi den er rele-

vant for terminologien i bidragene til dette nummer: Nogle bidrag bruger en terminologi som italesætter LLM'er som 'en maskine'. De taler fx om "tekster der er genereret" eller om "maskinens output". Andre bidrag bruger en terminologi som italesætter LLM'er som en intelligent aktør. De taler fx om at LLM'er "ræsonnerer" eller "forstår". I redaktionen finder vi denne variation fascinerende og illustrativ for det tidspunkt vi befinder os på i vores forståelse af LLM'er.

Under overskriften *Sproget, mennesket og teknologien* kan vi denne gang bringe fire artikler og to anmeldelser. De tre første artikler går meget direkte til spørgsmålet om LLM'er og deres evne til at producere sprog som ikke afviger fra det et menneske kan.

Bolette Sandford Pedersen, Nathalie C. Hau Sørensen, Sussi Olsen og Sanni Nimb giver i deres artikel *Evaluering af sprogforståelsen i danske sprogmodeller – med udgangspunkt i semantiske ordbøger* forslag til hvordan man kan afprøve og evaluere, *benchmarke*, danske LLM'ers evne til at håndtere tekstsemantisk-logiske opgaver som bl.a. synonymhåndtering, inferenser og følgeslutninger, eller, med forfatterens egne ord, LLM'ernes evne til at 'ræsonnere' og 'forstå'. Evalueringerne er baseret på en række eksisterende danske ordbøger og teknologiske sprogresurser som flere af forfatterne har udviklet over de seneste år. Forfatterne konkluderer at deres benchmarksystem er tilstrækkeligt udfordrende for de testede LLM'er, at de 'begår fejl', men også at ChatGPT 4.0 præsterer rigtigt godt.

I artiklen *Kan AI reproducere fagdisciplinær stemme? Et komparativt korpusbaseret studie af GPT4's evne til at reproducere fagdisciplinær stemme i AI-genereret sprogvidenskabelig prosa* sammenligner Ea Lindhardt Overgaard og Ulf Dalvad Berthelsen et korpus af sprogvidenskabelige tekster skrevet af GPT 4.0 med et korpus af menneskeskrevne sprogvidenskabelige tekster. De to korpusser sammenlignes for deres 'stemmes' *stillingtagen*, *engagement* og brug af *fagspecifikt ordforråd*. Analysen viser at GPT i høj grad er i stand til at efterligne tekster skrevet af mennesker, eller sagt på en anden måde, at den har lært at udtrykke sig med en passende akademisk stemme.

Jonas Nygaard Blom, Alexandra Holsting og Jesper Tinggaard Svendsen undersøger i deres artikel *På sporet af chatbottens sproglige fin-*

geraftryk. En sproglig ophavsanalyse af tekster skrevet af danskstuderende og ChatGPT om det er muligt at skelne sprogligt mellem tekster skrevet af nystartede danskstuderende på universitetet og ChatGPT. Forfatterne sammenligner blandt andet de to tekstkorpussers forskellige mængde af afvigelser fra retskrivningsnormen, udbredelsen af interaktionelle resurser samt afsnits- og periodelængde. Forfatterne konkluderer at hvis de studerende og ChatGPT får stillet samme opgave, har chatbotten dels en mere ensartet måde at skrive på, dels færre ortografiske fejl end de studerende. Det er dog ifølge forfatternes undersøgelse muligt i en vis udstrækning at sløre at en tekst er skrevet af ChatGPT ved at prompte den til at efterligne de studerendes skrivestil.

Tina Thode Hougaards bidrag *Emojier som parasproglige resurser i skrevet onlineinteraktion* skiller sig ud fra de øvrige ved ikke at handle om GPT og LLM, men i stedet om menneskers brug af teknologiens muligheder til at opfylde deres kommunikative behov. Computermediert kommunikation var i udgangspunktet tekstbåren, men med tiden udviklede sprogbrugere forskellige former for ikoniske tegn som ”:-)” som senere blev konventionaliseret med et stigende antal *emojier*. Hougaard undersøger det kommunikative potentiale for den slags tegn og konkluderer at de samlet set tjener den funktion at få den talendes krop *virtualiseret* i den computermedierede samtale.

De to sidste bidrag til dette nummer af NyS er Ib Ulbæks, lektor emeritus, anmeldelse af Peter Widells bog *Sproget i verden. Tekstforståelsens filosofiske grundlag* (udgivet i 2023) og Holger Juuls, leder af Center for Læseforskningen ved Københavns Universitet, anmeldelse af Christian Becker-Christensens bog *Dansk Retsskrivning: Bogstav – Lyd – Bogstav. En grafemisk beskrivelse af dansk ortografi* (udgivet i 2021). Anmeldelser der optræder i temanumre af NyS, er ikke i udgangspunktet knyttet til temaet.

Vi ønsker læserne god læselyst!

Evaluering af sprogforståelsen i danske sprogmodeller – med udgangspunkt i semantiske ordbøger

BOLETTE SANDFORD PEDERSEN, NATHALIE C. HAU
SØRENSEN, SUSSI OLSEN & SANNI NIMB

ABSTRACT

Artiklen beskriver hvordan vi har udviklet en række datasæt – et såkaldt benchmark – til at evaluere forskellige aspekter af sprogforståelse i danske sprogmodeller. Vores antagelse er at den viden der allerede er beskrevet i en række eksisterende danske ordbøger, kan opfattes som 'ground truth' for semantikken i det danske ordforråd. Vores metode går derfor ud på at 'vende' de semantiske ordbøger om og bruge dem til at generere et benchmark der afprøver modellernes evne til at forstå dansk. Mere specifikt undersøger vi hvor godt modellerne i) forstår synonymi, nærsynonymi, og hvornår noget er semantisk associeret, ii) skaber inferens i relation til begrebsmæssig viden og nedarvning af egenskaber fra overbegreb til underbegreb, iii) laver korrekte følgeslutninger i forbindelse med specifikke handlinger og hændelser, iv) skelner mellem centrale betydninger af et ord i kontekst og v) håndterer positiv og negativ konnotation eller 'sentiment' i løbende tekst. Vi afprøver vores datasæt på ChatGPT 3.5 turbo og på ChatGPT 4.0 og kan se at datasættene har en passende sværhedsgrad i forhold til hvad modellerne er i stand til at håndtere, om end ChatGPT 4.0 opnår særdeles gode resultater for flere af datasættene.

EMNEORD: danske sprogmodeller; evaluering; sprogforståelse; semantiske ordbøger; benchmark; ChatGPT

1 STORE SPROGMODELLER – HVOR GODT FORSTÅR DE DANSK?

Store sprogmodeller indgår som den helt centrale komponent i de nye intelligente sprog tjenester der er ved at revolutionere vores videnssamfund. Chatbotter som ChatGPT er taget i brug i store dele af samfundet – også på dansk – på trods af at vi endnu ikke helt forstår hvordan de virker, og hvor høj deres abstraktionsevne egentlig er.

Selv om de danske tekster der genereres af chatbotterne, ved første øjekast ser imponerende og temmelig overbevisende ud, afslører et nærmere eftersyn hurtigt en del huller i forhold til deres evne til at håndtere det danske ordforråd og dets semantik på en nuanceret og sprogligt kohærent måde. En del af disse mangler skyldes de sproglige og kulturelle bias der er opstået på grund af de engelsksprogede tekster som modellerne i høj grad er trænet på. Man kan derfor antage at disse problemer vil mindskes og måske med tiden forsvinde når og hvis sprogmodellerne i højere grad trænes på dansksproget materiale, og i mindre grad baseres på oversættelse eller 'sprogtransfer' fra engelsk og andre store sprog.

Det samme gælder for så vidt den problematik der opstår fordi modellerne primært er trænet på løbende tekst. I tekst formulerer vi typisk forgrundsviden, dvs. det der træder frem og er ekstraordinært og kontroversielt, hvorimod den baggrundsviden vi mennesker har og bruger når vi læser og forstår tekst, ikke træder så tydeligt frem da den forudsættes kendt. En del faktuel baggrundsviden kan findes i ordbøger og encyklopædier, men den er typisk sværere tilgængelig hvad angår udnyttelse i store sprogmodeller bl.a. på grund af betalingsmure og ophavsret. Et ofte anvendt eksempel til at illustrere sådanne statistiske bias hvor baggrundsviden er fraværende, er *det sorte får*. Sprogmodeller har længe fejlagtigt svaret *sort* hvis man spurgte hvilken farve et får har, selv om der i ordbøger og encyklopædier står at et får typisk er hvidligt. Det skyldes at det sorte får er det der slår igennem statistisk i teksterne. Udtrykket, som i sig selv er en metafor for noget der stikker ud (og ikke lever op til fællesskabets forventninger), bliver altså en ny metafor for statistisk bias eller skævvridning i sprogmodellerne, hvor væsentlig baggrundsviden ikke får nok vægt. På samme måde kunne man fejlagtigt tro at komikere typisk er kvinder fordi vi (stadig) synes det er relevant at nævne det eksplicit i teksten når de lige netop er kvindelige¹.

Andre problemer bunder i højere grad i sprogmodellernes tekniske design hvor vi endnu mangler helt at forstå hvordan og i hvor høj grad

1 Disse skævvridninger udjævnes hen ad vejen når datamaterialet vokser. Problematikken består imidlertid stadig i mindre, dansktrænede modeller, som vi potentielt også adresserer med vores datasæt så de netop kan bruges til at identificere modellernes begrænsninger.

de abstraherer og 'forstår'² de implicitte referencer vi mennesker anvender når vi bruger og fortolker sprog. Fx når det gælder overført og metaforisk sprog, når vi anvender ironi og sarkasme, eller når vi bruger ordforråd med negativ og positiv konnotation. Når den danske metafor *flødebolle* forklares af ChatGPT som *en der måske er blød, blid, eller ikke alt for modig eller robust* kan man som danskkyndig hurtigt konkludere at det er en temmelig skæv fortolkning; muligvis en fejlloverførsel fra det engelske *creampie* idet metaforen på dansk snarere refererer til en (oftest mands)person med meget høje tanker om sig selv uden at have så meget at have det i, med reference tilbage til det oppustede æggehvideskum som udgør flødebollens indre. Et andet eksempel på fejlloverførsel fra engelsk er æbleskive der af modellen tolkes til at være mad fra en plante, (dvs. bogstaveligt en skive af æble) hvor fænomenet ikke kendes som den kuglerunde kage vi især spiser i Danmark ved juletid.

Hvad enten de sproglige mangler skyldes for lidt dansk input i træningsmateriale, manglende baggrundsviden eller modellernes manglende abstraktionsevne mere generelt, er der brug for store og varierede evalueringsdatasæt, dvs. datasæt der systematisk kan belyse hvor godt sprogmodellerne virker, og hvornår og ved hvilke typer af udvikling de faktisk forbedres. Sådanne datasæt benævnes ofte *benchmarks*.

Det er vores antagelse at det er vigtigt at disse datasæt udvikles med et monolingvalt udgangspunkt fra det sprogsamfund de skal interagere i, og ikke, som det ofte ses, genereres via tilsvarende engelske datasæt. Ved oversættelse af datasæt sker der nemlig ofte det at bias igen introduceres via kildesproget, og at nuancer og mindre hyppigt forekommende eller særegent ordforråd i målsproget overses og dermed ikke afprøves. I denne sammenhæng er det vigtigt at understrege at relevante benchmarks også i høj grad handler om at ramme præcist i forhold til såkaldt 'aboutness' (jf. fx Hershcovich et al. 2022); dvs. de skal teste netop de aspekter der er særligt relevante for det givne sprog og det givne sprogsamfund, og ikke i hovedsagen evaluere aspekter som er mere relevante i fx et engelsk sprogsamfund.

2 Når vi i denne tekst anvender termerne 'forstå' og 'ræsonnere', refererer vi til sprogmodellens evne til at forholde sig nogenlunde adækvat til specifikke sprogforståelsesopgaver. Vi går således ikke nærmere ind i den sprogfilosofiske fortolkning af hvad sprogforståelse i øvrigt indbefatter af fx hensigt, kontekst osv.

Vores benchmark tager udgangspunkt i det faktum at de allerede etablerede semantiske ordbøger og ressourcer der findes på dansk, udgør valideret semantisk viden; vi betragter dem således som en slags 'ground truth' eller 'guldstandard'. I vores undersøgelse bliver de så at sige 'vendt om' og brugt til at generere et benchmark der afprøver udvalgte sprogmodellers evne til at forstå dansk. Og som gør det på en omfangsrig og struktureret måde der også inddrager de mindre prototypiske dele af sproget og fx de mindre hyppige ord.

Vi regner med at vores benchmark som i øjeblikket indbefatter i alt knap ca. 4200 testeksemplere fordelt over 26 datasæt³, kan anvendes til at afdække i hvor høj grad sprogmodellerne i) forstår synonymi, nærsynonymi og hvornår noget er semantisk associeret, ii) skaber inferens i relation til begrebsmæssig viden og nedarvning af egenskaber fra overbegreb til underbegreb og iii) laver korrekte følgeslutninger i forbindelse med handlinger og hændelser, iv) skelner mellem centrale betydninger af et ord og v) kan håndtere positiv og negativ konnotation på en adækvat måde.

Via vores mangeårige arbejde med udvikling af danske semantiske ordbøger ved vi at netop disse informationstyper er indeholdt i de leksikalske beskrivelser på en nogenlunde systematisk måde som giver mulighed for semiautomatisk at generere omfangsrige evalueringsdata. Som det vil fremgå, fokuserer vi på danske data og evaluering af modeller der fungerer på dansk, men vi antager at vores metodeudvikling og erkendelser i den forbindelse også vil være relevante for andre sprog. Særligt for de mindre og mellemstore sprog er der et udtalt behov for at få genereret sådanne datasæt semiautomatisk (og dermed med lave omkostninger) og i stor skala, så udviklingen af sprogmodeller for disse kan følges og vurderes nøje.

Vi indleder i afsnit 2 med en kort introduktion til sprogmodellerne virkemåde og til hvordan man typisk evaluerer dem i en sprogteknologisk kontekst; dels på dansk, dels på andre sprog. Herefter introducerer vi i afsnit 3 de semantiske ordbøger som vi betragter som 'ground truth' for vores benchmark. I afsnit 4 beskriver vi hvordan vi udvikler de enkelte datasæt på baggrund af information fra ordbøger-

³ Data er tilgængelige på <https://github.com/kuhumcst/danish-semantic-reasoning-benchmark> og vil løbende blive opdateret.

ne, og for hver sprogforståelsesopgave diskuterer vi desuden hvor godt de løses af udvalgte 'state of the art'-modeller. I afsnit 5 sammenligner vi sprogmodellernes løsning af de syv sprogforståelsesopgaver, og i afsnit 6 konkluderer vi og beskriver de videre perspektiver for at udvikle benchmarket.

2 HVORDAN VIRKER SPROGMODELLERNE, OG HVORDAN EVALUERER VI DEM?

En central bestanddel i store sprogmodeller er de såkaldte statistiske ordprofiler, også kaldet *word embeddings*, som udgøres af numeriske vektorer, en slags lister med numeriske værdier der repræsenterer hvilke ord et givent ord *typisk optræder sammen med* i et tekstkorpus. Hver ordprofil repræsenterer med sin vektor et punkt i et flerdimensionelt vektorrum, og ord med afvigende distribution vil ligge langt fra hinandens ord med lignende distribution og dermed lignende betydning⁴ vil ligge tæt på hinanden i vektorrummet. Synonyme ord som fx *sprinte* og *spurte* har således ordprofiler der ligger tæt på hinanden fordi de typisk optræder sammen med lignende ord i et korpus⁵.

Efterhånden som man i de senere år har øget antallet af parametre til lagring af oplysninger om ordenes distribution, har man også mangedoblet den sproglige information som sprogmodellen kan regne sig frem til og gør brug af. Det betyder i korte træk at modellerne repræsenterer mere tekstlig kontekst på et højere generaliseringsniveau, og at der i højere grad kan tages højde for langdistanceafhængigheder i det tekstlige input. Vi ser således at det hyppige danske fænomen med partikelverber hvor partiklen ofte strander langt fra sit verbum, nu kan fortolkes langt bedre af modellerne, som i eksemplet: *du kan slå alle de ord du ikke kender, op i ordbogen online*. Her er pointen at *op* skal fortolkes som en del af partikelverbet *slå op*, og at betydningen af verbet således ikke er gennemsigtig i forhold til betydningen af ordets to enkeltdele, hhv. *slå* og *op*.

4 Vi henholder os her til den distributionelle hypotese om at ord der har lignende distributionel kontekst, også har lignende betydning, jf. fx Firth (1957), Levin (1991) og relateret til NLP fx Lenci & Sahlgren (2023).

5 For en mere uddybende forklaring for lægmand kan vi henvise til fx Lee & Trott (2023) eller andre umiddelbart tilgængelige beskrivelser i blogposts og på wikipedia.

Modellens abstraktionsevne understøttes af det *lagdelte* design (som også kaldes neurale lag efter inspiration fra den menneskelige hjerne), hvor hvert lag tager en sekvens af ordprofiler som input, beregner yderligere information ud fra distributionen og videregiver disse udvidede vektorer til næste lag. En simpel googlesøgning afslører at ChatGPT 3.5 fx indeholder 175.000.000.000 sådanne parametre på tværs af 93 neurale lag, mens ChatGPT 4.0 indeholder ca. 1.8.000.000.000.000 parametre på tværs af 120 lag. GPT-Modellerne benævnes også *transformermodeller*, hvilket indikerer at der for hvert lag transformeres en inputsekvens til en outputsekvens med et højere generaliseringsniveau, og at dette sker ved hjælp af såkaldte 'self attention heads', som forstærker signalet mellem to eller flere relaterede ord (reelt tokens) i en sætning og dermed muliggør at mere generel information genereres og deles på tværs af inputtet. Hvert lag giver plads til endnu et abstraktionsniveau, og dermed øges den generelle viden om bl.a. ordenes morfologiske, syntaktiske og semantiske egenskaber, om forholdet mellem ord og ordsekvenser, og ikke mindst om hvilket ord der typisk følger efter som det næste. Det er væsentligt at bemærke at de modeller der i øjeblikket fungerer bedst for dansk, er flersproglige, og at de rent faktisk lærer en del om dansk sprog via andre sprog, særligt engelsk. Dette er en meget stor fordel da mange sproglige fænomener går igen fra sprog til sprog; ulempen er så at der opstår sproglige bias fordi det er svært at styre modellerne så de kun overfører *relevant* viden fra det ene sprog til det andet. Generelt gælder det, ikke overraskende, at jo mere tekst der trænes eksplicit med for det pågældende sprog, desto bedre fungerer modellen for netop det sprog.

Udover at analysere tekst bruges sprogmodeller i dag typisk produktivt til netop at forudsige hvilke ordsekvenser der hyppigt kommer efter et givent ord. Så taler vi om en *generativ* eller *selvproducerende* anvendelse af sprogmodellen, som er den vi ser i netop ChatGPT, hvor modellen selv skaber nyt indhold på en relativt selvstændig og kreativ måde, dog i udgangspunktet naturligvis baseret på den tekst som modellen er trænet med. Generative sprogmodeller og også generativ kunstig intelligens mere generelt kan altså forstås som systemer der selv skaber indhold, i modsætning til systemer der kun *analyserer* data.

Evaluering af sprogmodeller og relevante eksisterende benchmarks

Når man udvikler og vedligeholder sprogmodeller, er det almindeligt at anvende standardiserede og omfangsrige evalueringsdatasæt så man løbende kan afprøve i hvor høj grad og på hvilke områder modellerne udvikler sig. Datasættene giver mulighed for automatisk og systematisk at måle hvor godt de hver især løser en lang række opgaver, bl.a. med henblik på at sammenligne dem med hinanden. Opgaverne kan være designet så de tester modellernes evne til at udføre forskellige praktiske sprogtenester ('extrinsic evaluation'), som fx at oversætte mellem to sprog eller at lave et resumé af en tekst. Men de kan også designes så de mere generelt tester modellernes sprogforståelse fx i forhold til at kunne 'ræsonnere' om forskellige forhold, til at forstå nuanceforskelle mellem ord, og til at forstå konsekvensen af en beskrevet handling (samlet set også kaldet 'intrinsic evaluation'). Vores datasæt fokuserer på sidstnævnte af to grunde: dels mangler vi i høj grad sådanne datasæt for dansk, hvilket vanskeliggør en nuanceret evaluering af hvor godt danske sprogmodeller udvikler sig; dels er det her at vi antager at vores semantiske ordbøger udgør en helt særlig ressource.

Der findes altså allerede forskellige evalueringsdatasæt for dansk der løser praktiske sprogopgaver som fx tekstklassifikation, tekstresumering og maskinoversættelse. Flere af disse er tilgængelige via Digitaliseringsstyrelsens sprogteknologiportal 'sprogteknologi.dk', og de samles desuden i den skandinaviske evalueringsplatform *ScandEval* (Nielsen 2023) i de tilfælde hvor der findes tilsvarende datasæt for svensk og norsk.

Der eksisterer også datasæt rundt omkring i virksomheder (fx medievirksomheder) og ved de forskellige forskningsinstitutioner som kan bruges til at evaluere delelementer i sprogmodellernes performans. Disse indbefatter datasæt for leksikalsk-semantisk nærhed (fx Nielsen & Hansen 2017, Schneidermann et al. 2020) samt betydningsopmærkede korpora der kan anvendes til at vurdere hvor godt sprogmodellerne skelner mellem ordbetydninger i kontekst (Pedersen et al. 2016, 2023).

I det internationale miljø findes der pakker af datasæt og evalueringsplatforme der samlet set indbefatter begge typer evaluering. Det gælder fx GLUE- og SUPERGLUE-testsuiten (Wang et al. 2018, 2020) samt den nyligt udviklede svenske SUPERLIM-platform (Berdicevskis et al. 2023) lavet med SUPERGLUE som forbillede. Udfordringen er

at modellerne hele tiden forbedres; derfor må testapparatet følge med og gøres sværere. Således indeholder den seneste version af SUPERGLUE otte relativt komplekse sprogforståelsesopgaver, herunder beregning af koreference, entydiggørelse af flertydige ord, udredning af følgeslutninger mv. For norsk findes der for nuværende den såkaldte *NorBench* (Samuel et al. 2023), der blev præsenteret ved NODALIDA-konferencen i Thorshavn i 2023. Den indeholder dog først og fremmest et antal sprogjenesteopgaver af typen maskinoversættelse og spørgsmål–svar og ikke specifikke sprogforståelsesdatasæt som dem vi præsenterer her.

3 HVILKEN RELEVANT VIDEN RUMMER DE SEMANTISKE ORDBØGER?

De semantiske ordbøger som vi afprøver som grundlag for vores benchmark, er *Den Danske Begrebsordbog* (Nimb et al. 2014, Nimb 2016), det danske *wordnet*, *DanNet* (Pedersen et al. 2009), *Det Danske FrameNet-leksikon* (Nimb et al. 2017)⁶, og den semantiske komponent i *Det Centrale OrdRegister for dansk*, også kaldes *COR.SEM*⁷ (Nimb et al. 2022a). Alle de her nævnte ordbøger indeholder formaliserede semantiske oplysninger der mere eller mindre eksplicit bygger på beskrivelser af lemmaer i *Den Danske Ordbog* (Hjort & Kristensen 2003-2005; <https://ordnet.dk/ddo>).

Fra **Den Danske Begrebsordbog** kan vi udlede hvilke ord mennesker anser for at være synonyme og nært semantisk og tematisk beslægtede. Ordbogen indeholder ca. 100.000 lemmaer og 130.000 betydninger fra Den Danske Ordbog fordelt over 22 kapitler og 888 afsnit. Hvert afsnit er navngivet tematisk og indeholder en række ordklasseopdelte og dernæst semantisk inddelte grupper der hver især består af betydningsmæssigt relaterede ord og udtryk, hvor den indbyrdes afstand afspejler graden af semantisk nærhed. Nogle af grupperne indledes af såkaldte 'nøgleord', der fungerer som en slags overskrift. Et nøgleord er i langt de fleste tilfælde et overbegreb til de følgende ord, men det kan i nogle tilfælde også blot være det mest udbredte ord indenfor en mindre semantisk gruppe. Der opereres med to niveauer af nøgleord: overord-

6 En samlet beskrivelse på dansk af disse ordbøger og principperne bag dem kan findes i Pedersen et al. (2021).

7 <https://ordregister.dk>

nede og underordnede, hvor de overordnede har virkefelt over alle efterfølgende grupper, også dem der indledes med et underordnet nøgleord.

FIGUR 1: ET EKSEMPEL FRA DEN DANSKE BEGREBSORDBOG. OVERORDNEDE NØGLEORD HAR LILLA BAGGRUND (*UDREGNING*, *OPTÆLLING*) OG UNDERORDNEDE NØGLEORD HAR GRÅ BAGGRUND (*NUMMERORDEN*). HVER PRIK ADSKILLER BETYDNINGSGRUPPER.

Afsnit 11.011: Måle, regne

subst

udregning

- udregning, kalkulation, kalkulering, beregning, indregning, opgørelse, regnskab, statistik
- vurdering, taksation, overslag, et slag på tasken
- approksimation, fejleregning, mellemregning, efterregning, gennemregning, omregning, fratræk, gennemregning

optælling

- optælling, tælling, sammentælling, opregning
- initialisering
- nummerering, paginering

nummerorden

- nummerorden, nummerrækkefølge, nummer, nummerering, nr., no.

I figur 1 har det overordnede nøgleord *udregning* virkefelt over tre betydningsgrupper frem til det næste overordnede nøgleord *optælling*, som har virkefelt over fire betydningsgrupper, heraf én gruppe som også hører under det underordnede nøgleord *nummerorden*. Tættest semantiske lighed finder man indenfor samme betydningsgruppe (uanset om den indledes med eget nøgleord eller ej), dernæst blandt de øvrige betydningsgrupper med samme nøgleord i hele det navngivne afsnit eller i det samlede kapitel som afsnittet tilhører. Strukturen er udnyttet i præsentationen af afgrænsede grupper af semantisk beslægtede ord fra *Begrebsordbogen* i online-udgaven af Den Danske Ordbog (<https://ordnet.dk/ddo>) i form af funktionen ‘Ord i nærheden’ (Nimb et al. 2018).

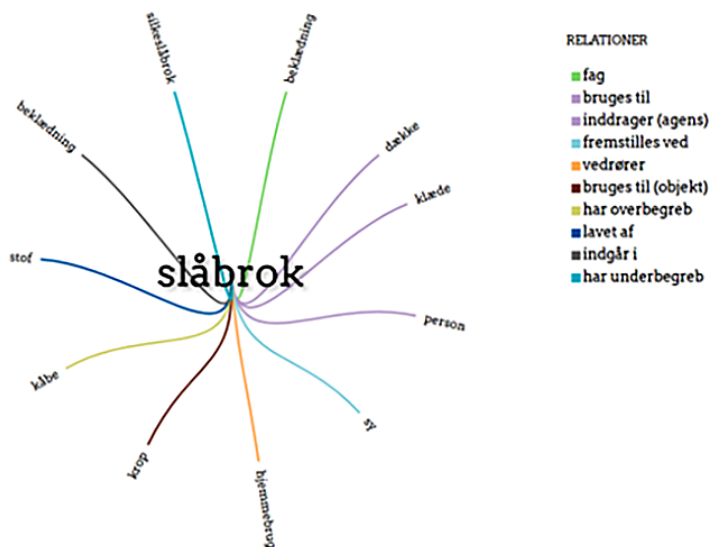
Ud fra DanNets, det danske wordnets, ontologi⁸ og semantiske relationer kan vi udlede en række prototypiske egenskaber (i alt er 70.000 lemmaer beskrevet), dels dem der er specifikke for hvert begreb, dels dem der nedarves fra mere overordnede begreber (fx at fødevarer typisk

8 DanNets ontologi bygger på *EuroWordNet*-ontologien (Vossen et al. 1999), men har derudover en række mere specificerede typer, jf. Pedersen et al. (2009).

har en nærende funktion). Med disse benchmarkdata vil vi gerne afprøve i hvor høj grad sprogmodellerne har tilegnet sig en tilsvarende viden om semantiske nedarvningsrelationer.

DanNet giver også information om begrebers kerneegenskaber angivet på baggrund af Pustejovskys (1998) teori om begrebers såkaldte *qualiastruktur*. Det er fx anført hvordan en entitet er blevet til (et brød fx ved bagning), og hvad nogets typiske funktion eller rolle er (fx at *lyse* for en lampe, eller at *undervise* for en lærer). Figur 2 illustrerer som eksempel en række kerneegenskaber for beklædningsgenstanden *slåbrok* som i udgangspunktet er udledt semiautomatisk fra begrebets definition i Den Danske Ordbog og kombineret med de nedarvede egenskaber fra dets overbegreber (fx at beklædningsgenstande typisk er blevet syet, og at de typisk bruges af mennesker for at tildække kroppen). Denne viden tester vi om sprogmodellerne besidder.

FIGUR 2: BEGREBET SLÅBROK OG DETS KERNEEGENSKABER SOM BESKREVET I DAN-NET BL.A. PÅ BAGGRUND AF DEN DANSKE ORDBOGS DEFINITIONER. ('FAG' SKAL HER FORSTÅS SOM EMNEOMRÅDE.)



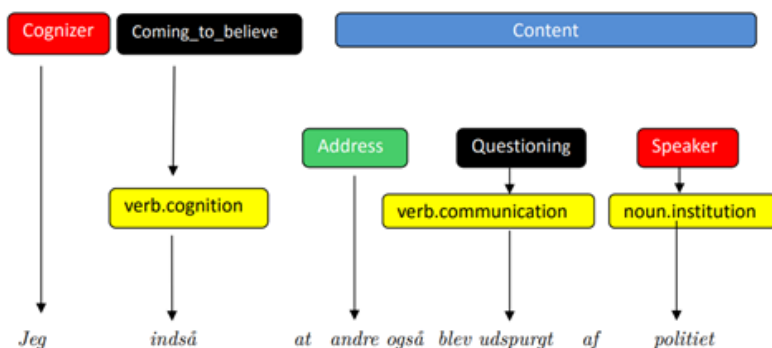
Fra **Det Danske FrameNet-leksikon** kan vi udlede hvilke slags handlinger, transaktioner og situationer danske verber og verbalsubstantiver refererer til, og kombineret med *Berkeley FrameNets* udfoldning af de

semantiske frames (<http://framenet.icsi.berkeley.edu/>) har vi viden om de semantiske roller der typisk forekommer i en ytring med en specifik frame, dvs. hvem der gør hvad med hvem. I en kommunikationsframe er der fx typisk nogen der kommunikerer, nogen der lytter, og noget information der bliver kommunikeret; for bevægelsesframes nogen eller noget der bevæger sig på en bestemt måde og evt. fra et sted til et andet, dvs. en kilde, en sti og et mål.

Tilsvarende kan andre handlinger eller hændelser have særlige følgevirkninger, dvs. en ny situation udløst af at handlingen har fundet sted. Gode eksempler er frames som CUTTING med verberne *klippe* eller *skære* og framen COMMERCE_BUY med verber for at købe. Vi afprøver her i hvor høj grad sprogmodellerne kan tage højde for disse følgevirkninger, og om de inddrager de semantiske roller som de forskellige verber aktiverer i en ytring.

Figur 3 illustrerer hvordan verberne *indse* og *udspørge* udløser hver deres semantiske frame for hhv. kognition og kommunikation (indeholdt i FrameNet-leksikonet), og hvordan de semantiske roller for hver frame realiseres i den givne ytring. Når denne kognition har fundet sted, har ‘jeg’et erkendt noget, som det ikke vidste før, nemlig at også andre var blevet udspurgt af politiet.

FIGUR 3: OPMÆRKNING MED FRAMES (I SORT) OG SEMANTISKE ROLLER FOR HHV. EN KOGNITIONS- OG EN KOMMUNIKATIONSFRAME, UDLØST AF HHV. INDSE OG UDSPØRGE (FRA NIMB ET AL. 2017: 12)



COR.SEM er en helt ny leksikalsk ressource, som i stor udstrækning samler, strømliner og forenkler semantisk information fra de før om-

talte ordbøger. Ordbogen rummer for nuværende 34.000 lemmaer og deres almene betydninger, hvor fokus har været på at beskrive de lemmaer der er enten begrebsmæssigt eller emnemæssigt centrale i dansk. Betydningsinventaret er en forenklet udgave af betydningsinventaret i Den Danske Ordbog (DDO) idet mange betydninger fra ordbogen enten er sammenlagt eller helt frasorteret ud fra en række systematiske principper og en nøje analyse af hvert lemma (jf. Pedersen et al. 2022 for en redegørelse af denne forenkling).

Ressourcen er således særlig værdifuld i forhold til vores benchmark idet den er velegnet til at evaluere modellernes evne til at *skelne mellem de centrale ordbetydninger i kontekst*, og den fungerer umiddelbart som en bedre guldstandard for dette end fx DDO's mere finkornede betydningsinventar. COR.SEM udgør en del af Det Centrale Ordregister⁹ som er et standardiseret ordregister som sikrer at alle lemmaer i dansk tildeles et unikt nummer på tværs af ordbøger og fagterminologier således at det er nemmere at samkøre forskellige leksikalske ressourcer til sprogteknologisk anvendelse, jf. figur 4.

FIGUR 4: DET CENTRALE ORDREGISTER: LIGESOM BORGERE I DANMARK ALLE HAR ET CPR-NR, FÅR ALLE LEMMAER I DANSK TILDELT ET UNIKT NUMMER PÅ TVÆRS AF ORDBØGER (FRA HENRICHSEN 2021).



COR.SEM indeholder også en mere detaljeret version af *Det Danske Sentimentleksikon* (Nimb et al. 2022b). Hvor *Det Danske Sentimentleksikon* mere forenklet angiver positiv eller negativ konnotation på lemma-niveau fra -5 (meget negativ) til +5 (meget positiv) og udelader lemmaer der er tvetydige (fx adj. *fed*), er konnotation i COR.SEM angivet for hver *betydning* (*fed* har fx 2 negative og 3 positive betydninger), dog

⁹ Det Centrale OrdRegister administreres af Dansk Sprognævn og er åbent tilgængeligt: <https://ordregister.dk>.

kun inden for en konnotationsskala fra -3 til +3, hvor -3 altså angiver meget negativ konnotation og +3 meget positiv konnotation. Da vi også har brugseksempler i form af citater fra DDO, kan vi med denne oplysning evaluere et ords sentiment i dets kontekst som en del af vores benchmark.

4 FRA INFORMATION I DE SEMANTISKE ORDBØGER TIL SPROGFORSTÅESESOPGAVER

Vores metode går altså som nævnt ud på at anskue de semantiske ordbøger som vores guldstandard eller 'ground truth' som vi kan evaluere sprogmodellernes evne til sprogforståelse¹⁰ op imod. Vi genererer datasættene semiautomatisk ud fra ordbøgerne og afprøver dem på i første omgang GPT3.5 Turbo og GPT4 via såkaldt prompting¹¹. Det gør vi for at få en første fornemmelse af om vores datasæt er tilstrækkelig udfordrende for modellerne. Vi stiller en række stilistisk stramme spørgsmål (jf. eksemplet i figur 5), hvor svarene kan evalueres automatisk idet sprogsmodellen skal vælge svarmuligheder fra et lukket vokabular. Det kan være et eller flere ord, en sandhedsværdi eller en numerisk værdi. Til hver opgave skriver vi en prompt som er delt op i en overordnet instruktion og en skabelon hvori dataeksempler kan indsættes automatisk. Den overordnede instruktion indeholder en kort forklaring af input til modellen, formålet med opgaven, og hvordan outputtet skal se ud. Nedenfor ses et eksempel fra en af opgaverne (synonymi, figur 5).

FIGUR 5: OPGAVEN TIL CHATGPT SOM DEN ER FORMULERET FOR SYNONYMI

Jeg kommer til at give dig et ord og en liste af mulige synonymer til ordet. Du skal vælge det ord fra listen som er det bedste synonymmatch. Dit svar skal kun være det ene synonym fra listen. Output skal være et enkelt ord og intet andet.

10 Igen er der tale om sprogforståelse på et niveau der angår leksikalsk semantisk og sætningssemantik. Man kan naturligvis forestille sig en lang række andre sprogforståelsesopgaver som inddrager funktionelle, tekstsemantiske eller andre forhold.

11 Flere af disse forsøg er også blevet præsenteret ved LREC-konferencen i Torino i 2024, jf. Pedersen et al. (2024).

Instruktionen er en enkelt sætning med en række pladsholdere som dataeksemplerne kan udskiftes med. Til synonymiopgaven fra figur 5 giver vi sætningen: 'Ordet er {a}, og de mulige synonymy er {b}, hvor {a} er et enkelt ord og {b} en liste af fire mulige synonymy', hvorfra sprogmodellen skal vælge det rigtige. Formålet med skabelonen er automatisk at kunne behandle et stort antal dataeksempler gemt i en fil. Vi har derfor udviklet en grænseflade som i baggrunden har API-adgang til *OpenAI's* generative sprogmodeller (ChatGPT 3.5 Turbo og ChatGPT 4.0). Igennem grænsefladen kan man vælge en fil med en række dataeksempler som bliver behandlet, og resultatet bliver gemt. Grænsefladen tillader også at man nemt kan skifte mellem model og opgavetype.

FIGUR 6: SCREENDUMP FRA LEXIBOT, VORES API-GRÆNSEFLADE TIL CHATGPT.¹²

12 Sprogmodellers 'temperatur' er en parameter hvormed man kan angive graden af tilfældighed og kreativitet i modellens output. En lav temperatur giver konsistente og mere ensartede svar; når temperaturen øges, stiger variationen og kreativiteten i outputtet. Ved højere temperatur øges også risikoen for irrelevante output. Her har vi kun lavet forsøg med standardindstillingen 1.00. Vi vil i senere forsøg undersøge om en lavere temperatur ændrer på resultaterne. Da vi instruerer sprogmodellerne i kun at måtte svare med et enkelt ord og dermed har sat en grænse for hvor kreativt et output de kan levere, er vores hypotese at en lavere temperatur ikke vil medføre større ændringer.

Det skal nævnes at denne opsætning blot er et eksempel på hvordan benchmarket kan bruges på én type sprogmodel; i praksis vil benchmarket kunne bruges på andre sprogmodeller med andre typer opsætninger¹³.

I det følgende gennemgås hver sprogforståelsesopgave for sig, og for hver opgave undersøger vi hvor godt den løses af de to udvalgte sprogmodeller.

Synonymi

Synonymiopgaven går ud på at finde det mest passende synonym til et målord fra en liste med fire mulige ord. Ordenes grad af synonymi varierer i forhold til målordet på en skala fra højeste grad af semantisk beslægtethed, 'snæver' synonymi, som er det korrekte svar, til mindste grad af semantisk beslægtethed (se nedenfor mht. eksempler). Skalaen er automatisk beregnet ud fra strukturen med kapitler, afsnit, betydningsgrupper og nøgleord i Den Danske Begrebsordbog, se ovenfor.

På baggrund af strukturen udtrækkes ordpar i form af to direkte naboord i Begrebsordbogen inden for samme betydningsgruppe. De ordpar der samtidig er synonyme i DDO, repræsenterer snæver synonymi og den højeste grad af semantisk beslægtethed. Når vi har identificeret et synonympar automatisk i DDO-manuskriptet, skal vi finde de resterende tre ord som modellen skal vælge imellem. For at opgaven skal have en vis sværhedsgrad, udvælger vi kandidater som i varierende grad er semantisk beslægtede med synonymparret. Den nærmeste kandidat udvælges blandt ord i en anden betydningsgruppe, men stadig under samme nøgleord i samme afsnit i Begrebsordbogen. Fx kan vi i det navngivne afsnit 'musikinstrumenter' under nøgleordet 'strenginstrument' finde synonymparret *strygeinstrument* og *stryger* og i en anden betydningsgruppe under samme nøgleord finde *elguitar* som en mulig tæt semantisk beslægtet kandidat. Den tredje kandidat findes i et af de andre navngivne afsnit inden for samme kapitel, fx i dette tilfælde kandidaten *musical* i afsnittet 'Teater, optræden', der ligesom afsnittet 'musikinstrumenter' ovenfor også hører under kapitlet 'Kunst og kultur'. Den fjerde og sidste kandidat findes i et af de øvrige 21 kapitler i Begrebsordbogen, som fx *nedfald* fra kapitlet 'Sted og bevægelse', idet

¹³ Ikke alle sprogmodeller instrueres med prompting. Nogle modeller kræver at man først transformerer data til et format modellen kender (fx vektorrepræsentation), og andre modeller klarer sig bedre hvis de først er blevet trænet eller fintunet til den specifikke opgave.

vi forinden har fraserteret de afsnit der er tematisk beslægtede. Den endelige liste kan ses nedenfor i figur 7.

FIGUR 7: EKSEMPEL PÅ DATA FRA SYNONYMI DATASÆTTET.

Målord	Kandidat 1 (korrekt svar)	Kandidat 2	Kandidat 3	Kandidat 4
Strygeinstrument	Stryger	Elguitar	Musical	Nedfald
Kunst og Kultur; Musikinstru- ment; Strenginstru- ment	Kunst og Kultur; Musikinstru- ment; Strenginstru- ment	Kunst og Kultur; Musikinstru- ment; Strenginstru- ment	Kunst og Kultur; Teater, optræden; Teaterform	Sted og bevægelse; Ned; Fald

Alle lister der anvendes i opgaven, består af ordkandidater udvalgt tilfældigt fra de overstående kriterier idet rækkefølgen af de fire kandidater også angives tilfældigt. I alt ender vi med 317 lister (hver med fire kandidater) fra 43 afsnit i Den Danske Begrebsordbog.

Begge modeller klarer opgaven godt med 93,6 % korrekt identificerede synonymer hos ChatGPT 3.5 Turbo og 96,2 % korrekte synonymer for ChatGPT 4.0. En del af ChatGPT 3.5 Turbos fejl skyldes at modellen ikke følger instruktionen og giver selve målordet som output i stedet for at vælge et ord fra listen. Med henblik på at analysere fejlene spørger vi modellerne med en ny prompt der samtidig beder modellen om at forklare sit valg. Her ser vi to mønstre af fejl: 1) den har ikke kendskab til betydningen af et af ordene i synonymparret (fx *mø* i parret '*mø, jomfru*', *minut* i parret '*minut, bueminut*', *citronmelisse* i parret '*citronmelisse, hjertensfryd*' og det historiske militære udtryk *retrate* i parret '*retrate, tappenstreg*'), eller 2) den nægter at svare pga. stødende eller nedsættende sprog (fx i tilfældene '*aboriginer, australneger*' og '*tysskerpige, tyskertøs*'). Det sidste sker også når ChatGPT 4.0 anvendes, dog kun med synonymerne '*sort, neger*'.

I de resterende 11 tilfælde af fejl som ChatGPT 4.0 laver, vælger den ni gange den næst-tætteste ikke-synonyme kandidat (samme afsnit og nøgleord, men anden betydningsgruppe under nøgleordet). I ét tilfælde, for ordet *felt*, vælger den fejlagtigt det ord som er mindst semantisk beslægtet idet det er fra et helt andet kapitel, nemlig ordet *følelsesliv* blandt mulighederne '*område, flig, overlegenhed, følelsesliv*'. Det rigtige

synonym er *område*. Dette kan skyldes en klassisk transferfejl hvor behandlingen af det danske eksempel på en uhensigtsmæssig måde går via engelsk. Modellen sammenblander det engelske verbum *to feel* ('at føle'), som bøjes *felt* i datid, og det danske substantiv *felt* i betydningen 'afgrænset flade eller område'.

Semantisk nærhed i et bredere perspektiv

Vi evaluerer også sprogmodellernes evne til at vurdere grader af semantisk lighed i videre forstand end den rent snævre synonymirelation der afspejles i DDO's synonymangivelser. Det gør vi ved at stille modellerne overfor to såkaldte *word intrusion*-opgaver (Nielsen & Hansen 2017), hvor et afvigende ord skal identificeres i en ordgruppe der i øvrigt består af nært beslægtede ord. Hvert testeksempel består med andre ord af en kernegruppe af nære ord og en afviger idet alle ordene udtrækkes automatisk fra Begrebsordbogen. Efterfølgende udvælger vi manuelt de bedste testeksempler til brug for opgaverne.

Formålet med den første *word intrusion*-opgave er at evaluere semantisk nærhed. I denne opgave er hovedgruppen med de nært beslægtede ord derfor synonyme eller nærsynonymer. De stammer fra den samme betydningsgruppe i Begrebsordbogen, dvs. fra en gruppe af ord der ifølge dens struktur er så tæt beslægtet semantisk som ord kan være. Herefter suppleres den med et afvigende ord fra en anden betydningsgruppe der dog stadig hører under samme overordnede nøgleord, dvs. er forholdsvis tæt semantisk beslægtet med ordene i kernegruppen, men helt sikkert ikke er et synonym.

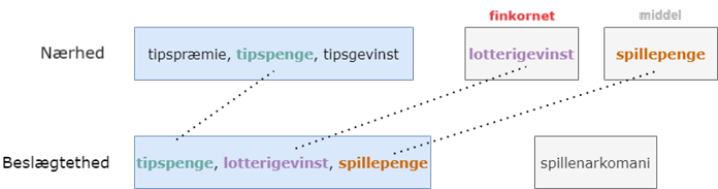
Vi opdeler opgaven i en 'finkornet' og 'grovkornet' underopgave, idet hypotesen er at den finkornede underopgave er sværest at løse. Baggrunden for denne niveauopdeling er at vi gerne vil udfordre sprogmodellerne mest muligt, og vi kender i forvejen ikke til modellerens nuværende niveau. De to underopgaver adskiller sig kun fra hinanden ved at have to forskellige afvigende ord, kernegruppen er den samme som det kan ses i figur 8.

FIGUR 8: EKSEMPEL PÅ EN KERNEGRUPPE OG ET AFVIGENDE ORD FRA HHV. ET GROV- OG FINKORNET DATASÆT OM SEMANTISK NÆRHED

Nært beslægtede ord	Afvigende ord	
tipspræmie, tipspenge, tipsgevinst	Finkornet:	lotterigevinst
	Grovkornet:	spillepenge

I etableringen af det finkornede datasæt tages det afvigende ord fra en anden betydningsgruppe der hører direkte under det samme overordnede nøgleord og derfor må formodes at være forholdsvis lidt afvigende fra de øvrige ord. I det grovkornede datasæt tages det afvigende ord på samme måde fra en af de betydningsgrupper der hører under det samme overordnede nøgleord, men i dette tilfælde hører det desuden under et underordnet nøgleord. Ud fra strukturen i Begrebsordbogen må man derfor antage at det drejer sig om et betydningsmæssigt lidt mere afvigende ord end i det første datasæt. I begge datasæt vil det afvigende ord dog være forholdsvis tæt semantisk beslægtet med de tre øvrige ord da alle fire ord har samme overordnede nøgleord.

FIGUR 9: SAMMENLIGNING AF KERNEGRUPPEN (BLÅ) I DATASÆTTENE OM NÆRHED OG BESLÆGTETHED. BEMÆRK AT AFVIGENDE ORD FRA DATASÆTTET 'NÆRHED' (GRÅ KASSER) GODT KAN FREMTRÆDE I DEN SAMME KERNEGRUPPE I DATASÆTTET 'BESLÆGTETHED'.

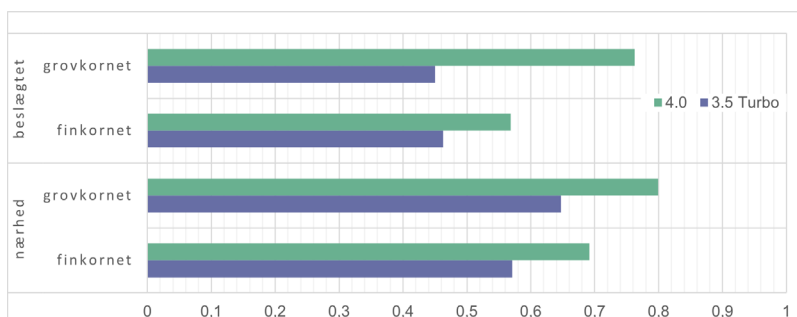


I den anden *word intrusion*-opgave, som vi betegner 'beslægtethed' (eller association), indgår der også to forskellige underopgaver. Endnu engang opstiller vi både en finkornet og en grovkornet underopgave. I det første tilfælde tages det afvigende ord fra en ordgruppe der, modsat ovenfor, hører under et andet overordnet nøgleord, men dog står i samme afsnit. Inden for dette datasæt vil der derfor være tale om en større semantisk afvigelse mellem kernegruppen og det afvigende ord end i den første word intrusion-opgave. Det kan fx være den forskel man finder mellem ord der beskriver handlinger, og ord der beskriver personer (se figur 9).

Til brug for den anden underopgave indenfor 'beslægtethed', den grovkornede underopgave, etablerer vi endnu et datasæt, denne gang bestående af en gruppe ord der blot har det til fælles at de står i samme afsnit i Begrebsordbogen, og hvor det fjerde afvigende ord tages tilfældigt fra et helt andet afsnit, dog et afsnit der stadig hører under samme kapitel. Vi opnår dermed et datasæt med tre ord der er tematisk tæt beslægtet¹⁴, hvorimod det fjerde ord er fjernere beslægtet tematisk set, men dog stadig associativt beslægtet.

Det er vigtigt at understrege at vi manuelt har valideret alle de dataudtræk der indgår i opgaverne for at sikre os at et menneske kan adskille det afvigende ord fra kernegruppen i alle de dataeksempler der anvendes.

FIGUR 10: SVARKORREKTHED I PROCENT FOR CHATGPT 3.5 TURBO (BLÅ) OG CHATGPT 4.0 (GRØN) PÅ TVÆRS AF DATASÆTTENE SORTERET EFTER GRADEN AF SEMANTISK SAMMENFALD. ØVERST ER DER MINDST SAMMENFALD (BESLÆGTETHED, GROVKORNET) OG NEDERST ER DER MEST SAMMENFALD (NÆRHED, FINKORNET).



I figur 10 ses resultaterne af en kørsel af datasættene på hhv. ChatGPT 3.5 Turbo (blå) og ChatGPT 4.0 (grøn). Datasættene er sorteret således at der er mindst semantisk sammenfald øverst (beslægtethed, grovkornet) og mest sammenfald nederst (nærhed, finkornet). Som forventet ser vi generelt at de opgaver vi betegner som 'grovkornede' under begge underopgaver er mindre udfordrende for modellerne. Når vi kigger på selve underopgaverne 'nærhed' versus 'beslægtethed', ser vi at model-

¹⁴ Kernegruppen på dette niveau har nu et emneområde (fx hører *togskinne*, *lokomotiv* og *sporarbejde* alle under emneområdet *jernbane*) eller en kategori (*lilje*, *myrte* og *citrontimian* er alle en type *planteart*) til fælles.

lerne klarer sig bedst ved 'nærhed'. Det kan skyldes at betydningsgrupperne er klart afgrænsede. I datasættene af typen 'beslægtethed' er der langt mere variation mellem ordene da de er blevet udledt automatisk af strukturen i Begrebsordbogen hvorimod de øvrige datasæt er baseret på ordbogens betydningsgrupper og dermed manuelt grupperet ud fra semantisk lighed da den blev skrevet. Flertydighed skaber desuden ikke overraskende problemer i genereringen af kernegrupperne i begge underopgaver af typen 'beslægtethed' hvilket man skal være opmærksom på når man analyserer resultaterne.

Entydiggørelse af ord i kontekst

For at evaluere sprogmodellernes evne til at forstå og adskille betydninger af polyseme ord i en kontekst, genererer vi et såkaldt *Ord-i-Kontekst-datasæt* (af engelsk *Word-in-Context*, Pilehvar & Camacho-Collados 2019). I denne opgave får modellen et målord og to sætninger som begge er korpuseksempler der indeholder målordet. Opgaven går ud på at svare på om målordet er brugt med den samme betydning eller to forskellige betydninger i de to sætninger. Nedenfor ses et eksempel med to forskellige betydninger for målordet *kost*.

- a) *Fødevarer, som indeholder fedt og sukker, er ikke sund **kost**.*
- b) *Han kom fra en familie, hvor skænderier og drillerier var daglig **kost**.*

I sætning a) refererer *kost* til en konkret betydning af *mad/madvare*, mens sætning b) er en overført betydning i DDO og COR.SEM (*noget som man oplever, skal sætte sig ind i eller på anden måde forholde sig til*). For at kunne svare korrekt på om sætning a) og b) indeholder den samme eller to forskellige betydninger af *kost*, skal man altså være i stand til entydiggøre målordets betydning i hver sætning og sammenligne resultatet.

Vi baserer betydningerne i denne opgave på det i mange tilfælde forenkledede betydningsinventar i COR.SEM sammenlignet med DDO, som beskrevet ovenfor. De brugseksempler der angives for størstedelen af COR.SEM's betydninger, anvendes som sætninger i vores Ord-i-Kontekst-datasæt. Vi opretter to versioner. Først laver vi

et rent monosemt datasæt, hvilket måske kan undre da der jo så kun kan være brugseksempler med samme betydning. Men de monoseme data gør det muligt at teste om en model er konsekvent ved entydiggørelse af et målord. Dertil burde de monoseme ord være nemmere at behandle for modellen end de polyseme da de ikke har den samme asymmetri mellem form og betydning. Denne 'lettere' opgave kræver med andre ord at man først entydiggør målordet, og et dårligt resultat vil for dette datasæt være et tegn på at modellen ikke løser opgaven efter hensigten. Dette skete netop for ChatGPT 3.5 Turbo, som konsekvent svarede at sætningerne i det monoseme datasæt indeholdt forskellige betydninger.

For de polyseme ord indsamler vi alle brugseksempler på tværs af betydninger for hvert målord og genererer samtlige mulige kombinationer af sætningspar. Da dette antal varierer både i forhold til antallet af målordets betydninger og i forhold til hvor mange brugseksempler hver betydning har, afgrænser vi antallet af eksempler til maksimalt tre for hver kategori (samme betydning og forskellig betydning).

Ved anvendelsen af ChatGPT 4.0 opnår vi henholdsvis 88 % svar-korrekthed for de monoseme data og 66 % for de polyseme.

Sentimentværdi af ord i kontekst

Vi evaluerer også sprogmodellernes fortolkning af ords såkaldte sentimentværdi, dvs. deres positive eller negative konnotation. I denne opgave giver vi et målord som input suppleret med et brugseksempel med målordet. Opgaven går ud på at svare på om målordet i den givne kontekst er brugt negativt eller positivt på en skala fra -3 (stærkt negativt) til +3 (stærkt positivt). Vi genererer datasættet automatisk baseret på COR.SEM ved at udtrække alle betydninger der både har sentimentværdi og mindst ét brugseksempel. Vi har i samme omgang også manuelt udvalgt en række neutrale betydninger som ikke har sentimentværdi i COR.SEM. Vi validerer manuelt alle eksempler (altså et målord med tilhørende brugseksempel og sentimentværdi) for at være sikre på at sentimentværdien kan identificeres af et menneske i brugseksemplet.

Nedenfor ses eksempler på målord (fed) i kontekst med målordets sentiment.

- a) *Landboforeningerne øjner en historisk **chance** for at få afskaffet jordskatten.* (+2)
- b) *Undervisningen i engelsk som tilvalgsfag **påbegyndes** i 2. semester.* (0)
- c) *Danmark er blevet meget rost for at være én af de første nationer, der afskaffede **slaveriet**.* (-3)

En væsentlig distinktion i denne opgave udgøres af sentimentværdi på hhv. ordniveau og sætningsniveau. Et målord kan forekomme i en kontekst som i sin helhed har en anden sentimentværdi end målordet, fx som i sætning c) hvor målordet *slaveriet* har en sentimentværdi på -3 mens selve sætningen har en positiv konnotation. Netop de eksempler hvor målordets konnotation ikke svarer overens med konnotationen af den samlede kontekst, har givet problemer for modellerne i vores eksperimenter.

Svarene i denne opgave afviger fra svarene i de andre typer opgaver da vi her arbejder med en *skala*. Vi beregner derfor en såkaldt Pearsons korrelationskoefficient, dvs. den lineære sammenhæng mellem to variable. De to variable er i vores tilfælde det korrekte svar (sentimentværdien i COR.SEM) og modellernes output. ChatGPT 3.5 Turbo opnår med denne tilgang en koefficient på 69,6 %, mens ChatGPT 4.0 opnår en koefficient på 81,9 %, hvor det ideelle resultat altså er en koefficient på 100 %.

Inferens om begrebsmæssigt tilhørsforhold

Sprogforståelsesspørgsmål i forhold til begrebsmæssigt tilhørsforhold og semantiske relationer genererer vi automatisk ud fra DanNet-ontologien og de relationer som de ontologiske begreber har indkodet. Vi anvender i første omgang kun skabeloner af typer 'X er en Y', 'X er en slags Y', 'Man bruger X til Y', 'Man laver X ved at Y', 'X kan have en Y' og 'X er en del af Y'. Variablene udfyldes for forskellige ontologiske typer, både konkrete og abstrakte¹⁵.

Figur 11 nedenfor viser eksempler på forskellige slags autogenererede forståelsesspørgsmål på basis af disse skabeloner der altså dels refererer

¹⁵ Begreber der tilhører konkrete ontologiske typer, er fx *hus* og *træ*, mens begreber der tilhører abstrakte ontologiske typer, kunne være *kærlighed* og *filosofi*.

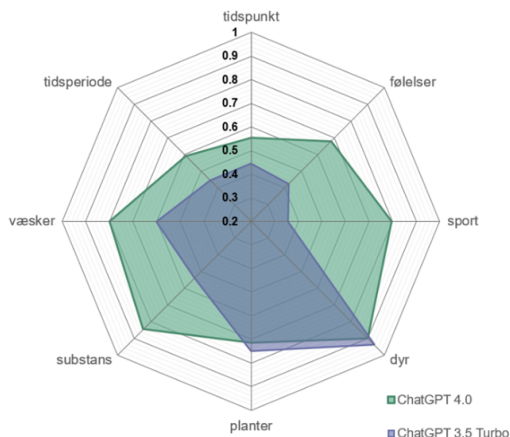
til hvilken overordnet ontologisk kategori et begreb tilhører, dels hvad noget typisk bruges til, hvordan det er blevet til, hvilke dele det har, og/eller hvad det selv er en del af. Udsagnene i del 1 bruges til at forberede chatbotten på spørgsmålstypen, hvor del 2 angives som et 'query' som chatbotten skal besvare som enten sandt eller falsk. Bemærk at der stilles både positive og negative spørgsmål (på en randomiseret måde), noget som sprogmodellerne traditionelt har haft lidt svært ved at skelne mellem.

FIGUR 11: AUTOGENEREREDE UDSAGN OM ONTOLOGISKE TYPER OG SEMANTISKE RELATIONER.

typer	Udsagn del 1	Udsagn del 2 (query)
Feeling; Creature	<i>Sympati er en følelse; Tryghed er en følelse</i>	<i>Q: Spøgelse er en følelse (falsk)</i>
Liquid; Quantity	<i>En bouillon er en væske; En drik er en væske</i>	<i>Q: En slurk er ikke en væske (sandt)</i>
Semiotic; Artifact	<i>Man laver en roman ved at skrive den; Man laver et essay ved at skrive det</i>	<i>Q: Man laver ikke en hat ved at skrive den (sandt)</i>
Food; Liquid	<i>Man laver et tog ved at fremstille det; Man laver en ret ved at tilberede den</i>	<i>Q: Man laver te ved at pochere den (falsk)</i>
Garment; Artifact	<i>Man tager en frakke på for at holde sig varm; Man tager en hue på for at holde sig varm</i>	<i>Q: Man tager en ring på for at holde sig varm (falsk)</i>
Instrument	<i>Man bruger en kniv til at skære med; Man bruger en hammer til at hamre med</i>	<i>Q: Man bruger ikke et rivejern til at rive med (falsk)</i>
BodyPart; Part	<i>En hånd kan ikke have et øje; Et ansigt kan have en mund</i>	<i>Q: Et fly kan have en propel (sandt)</i>

De autogenerede udsagn kan af og til virke pudsige, særligt når flertydige ord er i spil. Fx vil de færreste ved første øjesyn mene at en kylling er en væske, men ordet har faktisk den betydning om (indholdet i) en lille flaske spiritus. Vi overvejer i hvor høj grad vi i kommende eksperimenter manuelt bør luge ud i sådanne pudsigheder og mere sjældne betydninger, eller om de bør være med netop for at dække ordforrådet i alle dets afskygninger.

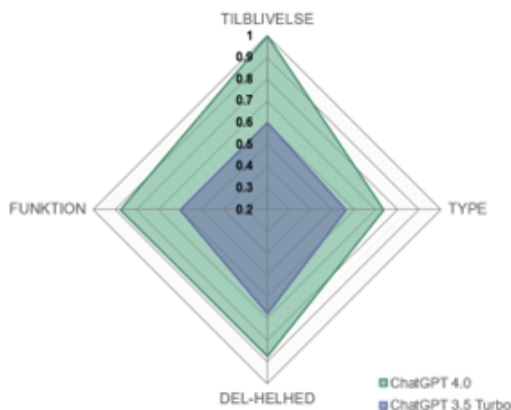
FIGUR 12: SPROGMODELLERNES (CHATGPT3.5 TURBO (BLÅ) OG CHATGPT4S (GRØN)) FORSTÅELSE AF HVILKE BEGREBER DER TILHØRER HVILKE ONTOLOGISKE TYPER.



En kørsel med datasættet på hhv. ChatGPT3.5 Turbo (blå) og ChatGPT4 (grøn) viser i figur 12 ikke så overraskende at forståelsen overordnet set er bedst for mere konkrete begreber som dyr og planter, hvorimod tidlige enheder og abstrakte typer som følelser er vanskeligere for modellerne at håndtere. Selv om ChatGPT3.5 Turbo fungerer en anelse bedre på dyr og planter, så ses der generelt en meget stor forbedring fra ChatGPT3.5 Turbo til ChatGPT4.

Figur 13 angiver hvordan de forskellige semantiske relationer (inkl. ontologisk type) forstås af de samme to sprogmodeller (ChatGPT.5 Turbo (blå) og ChatGPT4 (grøn)). Her ses generelt en god forståelse af hvordan noget er blevet til, hvorimod de andre semantiske relationer er lidt vanskeligere – bl.a. pga. af de mere abstrakte ontologiske typer som trækker resultatet ned. Igen opnår den nyeste af de to modeller de bedste resultater.

FIGUR 13: CHATGPT3.5 TURBOS (BLÅ) OG CHATGPT4S (GRØN) FORSTÅELSE AF SEMANTISKE RELATIONER.



Følgeslutninger vedr. handlinger og hændelser

For at afprøve modellernes evne til at foretage følgeslutninger efter at forskellige hændelser og handlinger har fundet sted, udarbejder vi for hver frametype i FrameNet-leksikonet to standardsætninger. Den ene eksemplificerer en handling/hændelse som et givent verbum eller verbalsubstantiv udløser, og den anden angiver resultatet og fremsættes som et 'query' som er enten sandt eller falsk. Standardsætningerne anvendes til at autogenerere udsagn med de verber og verbalsubstantiver der falder inden for den samme ramme. Igen blandes positive og negative udsagn på en randomiseret måde for at teste sprogmodellernes evne til at håndtere negation.

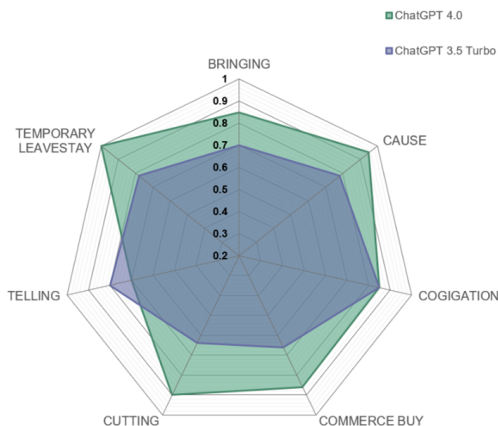
FIGUR 14: UDSAGN FOR FORSKELLIGE FRAMES. ANDET UDSAGN INDIKERER RESULTATET AF HANDLINGEN/HÆNDELSEN OG FREMSÆTTES SOM ET 'QUERY' (SANDT ELLER FALSK).

FRAME	UDSAGN
BRINGING (bringe)	<i>Peter bringer mad ud til Pia</i> Q: <i>Pia har nu mad.</i> (sandt)
CAUSE (forårsage)	<i>Ekspllosionen medførte svære skader på bygningen.</i> Q: <i>Efter eksplosionen var bygningen ubeskadiget</i> (falsk)
COMMERCE_BUY (handel_køb)	<i>Peter købte bogen af Anne</i> Q: <i>Nu ejer Anne bogen</i> (falsk)
CUTTING (skære)	<i>Pia klipper rebet over</i> Q: <i>Pia har nu to kortere reb</i> (sandt)
TELLING (meddele)	<i>Peter fortalte Pia om forlovelsen</i> Q: <i>Nu kender Pia ikke til forlovelsen</i> (falsk)
TEMPORARY_LEAVESTAY (midlertidigt fravær/ophold)	<i>Per tog på ferie i uge 7</i> Q: <i>Per var ikke på arbejde i uge 7</i> (sandt)
COGITATION (overvejelse)	<i>Peter tænkte på Hanne hele tiden</i> Q: <i>Hanne var aldrig i Peters tanker</i> (falsk)

Hvis vi afprøver vores datasæt for følgeslutninger på de samme to sprogmodeller som før, ser vi en generelt god forståelse særligt ved anvendelsen af ChatGPT4 (grøn) som er overraskende god til at tackle 'queries' vedr. konkrete handlinger og hændelser. Kommunikationshandling (som fx i tilfældet Telling) er den type der volder størst problemer, hvilket ikke undrer. Disse er mindre entydige hvad resultat angår; fx er et budskab ikke nødvendigvis forstået selv om det er afsendt (se figur 15)¹⁶.

16 Vi har i opgaven slået to frames sammen: Temporary_Leave og Temporary_Stay til TemporaryLeaveStay. Framen Temporary_Leave beskriver verbale udtryk som fx *tage orlov* og *skulke* mens Temporary_Stay be-skriver ord som *ophold*, *gøre holdt* og *gå i bi*. Framen Cogitation dækker over udtryk for tankevirksomhed såsom *tænke* og *analysere*, men også *skumle* og *dagdrømme*.

FIGUR 15: CHATGPT3.5 TURBOS (BLÅ) OG CHATGPT4'S (GRØN) HÅNDBLING AF FØLGESLUTNINGER VED HANDLINGER OG HÆNDELSER.



5 SAMMENLIGNING AF DE SYV SPROGFORSTÅELES-OPGAVER

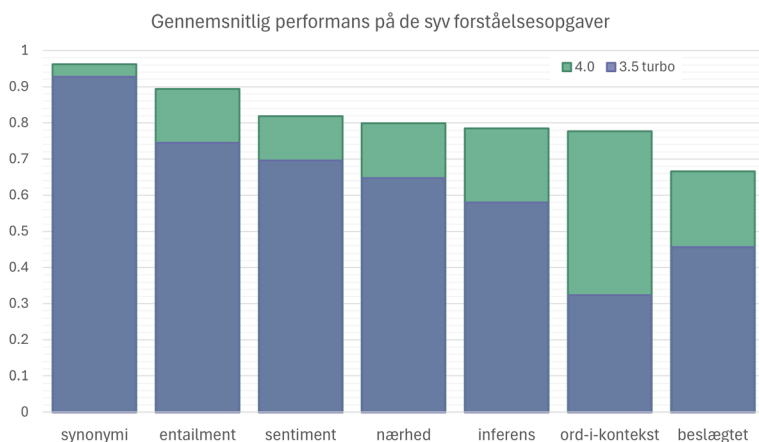
Figur 16 viser hvordan ChatGPT har præsteret gennemsnitligt for hver af de syv opgaver. Synonymi er den opgave som modellerne klarer bedst. Dette er ikke så overraskende da netop semantisk nærhed er noget af det som modellerne er særligt skarpe til at håndtere via de såkaldte word embeddings. Modellerne klarer dog opgaven så godt at vi skal videreudvikle vores testsæt i forhold til sværhedsgraden for reelt at kunne evaluere præstationen. Fx kan man tilføje flere kandidater som semantisk er tættere på synonymiparret i listen af muligheder eller tilføje valør til opgaven således at det korrekte synonym også skal matche valøren i DDO (fx uformelt eller nedsættende).

Det kan umiddelbart overraske at beslægtethed er så relativt svær en opgave for modellerne at løse, men det skyldes at det er svært at skelne imellem om noget ligner hinanden semantisk eller blot er associativt beslægtet. Hvor *kop* og *kaffe* er associativt beslægtet, dvs. nært relateret, er *kaffe* og *te* semantisk tættere på hinanden idet begge refererer til væsker vi indtager, og dermed har samme ontologiske type og samme overbegreb. De sproglige, dvs. de ordmæssige og ordfølgemæssige kontekster for *kop* og *kaffe* er imidlertid næsten de samme som for *kaffe* og *te* (i fx et tekstkorpus). Dette vanskeliggør opgaven for modellerne.

Vi ser også at forståelse af ord i kontekst er en af de svære opgaver for modellerne at løse, også selv om det betydningsinventar der evalueres op imod, er relativt grovkornet (fra COR.SEM). Her er det relevant at nævne at også vi mennesker generelt er uenige i mange tilfælde når vi skal opmærke løbende tekst med ordbetydning (jf. fx Pedersen et al. 2016 om semantisk opmærkning på dansk).

ChatGPT 4 klarer til gengæld inferensopgaven og særligt opgaven om følgeslutninger (kaldet entailment i figur 16) overraskende godt. Også her må vi tænke i mere udfordrende opgaver i fremtidige datasæt.

FIGUR 16: DE SYV FORSTÅELESOPGAVER AFPRØVET PÅ TO VERSIONER AF CHATGPT.



6 KONKLUSION OG VIDERE PERSPEKTIVER

Denne artikel beskriver de første forsøg med at udvikle et benchmark i stor skala som kan bruges til at evaluere danske sprogmodellers evne til at håndtere leksikalsk-semantiske forhold i dansk sprog. Udgangspunktet har været at udnytte de data som vi allerede har tilgængelige via de semantiske ordbøger vi har udviklet gennem flere år, og at betragte dem som vores 'ground truth' eller guldstandard. Vores eksperimenter har vist at metoden er realiserbar, og at de benchmarkdatasæt der udvikles, rent faktisk udfordrer de nyeste sprogmodeller på et passende niveau. Vi planlægger at skalere op så en større del af ordforrådet dækkes, og

så vi opnår en større variation og dermed også en højere sværhedsgrad i de stillede opgaver. På sigt vil vi også gerne kunne evaluere hvor godt sprogmodellerne fortolker mere avancerede og subtile dele af det danske ordforråd, fx metaforer og overført sprog. Vi antager at det især er inden for disse områder at sprogtransfer fra engelsk slår negativt igennem på dansk og genererer upræcist eller forvrøvlet sprog, jf. eksemplerne fra afsnit 1 om *flødebolle* og *æbleskive*. Hvor godt forstår modellerne fx sammenhængen mellem et ords grundbetydning og dets metaforiske søsterbetydning sådan som vi mennesker intuitivt gør når vi fx siger at noget er en *en følelsesmæssig rutsjebane*, eller at noget er *et vindue til verden*? Til dette formål ønsker vi at udnytte data fra Den Danske Ordbog, hvor mindst 6000 lemmaer er beskrevet med dels deres grundbetydning, dels deres afledte metaforiske betydning.

Vores forsøg har imidlertid også givet os anledning til at stille en række metodiske spørgsmål som det vil være interessant at efterforske i det videre arbejde. Ét spørgsmål vedrører i hvor høj grad vores benchmarks i deres nuværende form svarer til hvad almindelige mennesker ville svare på når de bliver stillet over for samme sprogforståelsesopgaver. Konkurrerer sprogmodellerne her i virkeligheden mod 'supermennesker' som præsterer langt bedre end den almindelig sprogbruger? Dette spørgsmål drøftes bl.a. i Tedeschi et al. (2023), og vi overvejer at undersøge det ved at lave nogle stikprøver med en række informanter der ikke er lingvistisk eller leksikografisk skolede.

Et andet spørgsmål som kunne være interessant at undersøge nærmere, vedrører de monolingvalt genererede benchmarks' påståede suverænitet overfor oversatte benchmarks. Her må der tænkes lidt over hvordan vi rent faktisk kan afprøve om vores benchmarks er de oversatte benchmarks overlegne. Vi forestiller os at en grundig analyse af evalueringresultater baseret på fx oversatte sentimentdata sammenlignet med resultater baseret på vores egne monolingvalt genererede sentimentdata kan være med til at belyse dette.

Endelig kan man stille sig selv det spørgsmål om evnen til at løse generelle sprogforståelsesopgaver modsvares af evnen til at løse praktiske sprogtjenesteopgaver a la resumeringsopgaver, spørgsmål-svar-opgaver eller maskinoversættelse. I og med at vi stadig mangler at forstå i dybden hvordan de meget komplekse modeller arbejder, og hvilket abstrak-

tionsniveau de egentlig opnår, er det ikke givet at der er en fuldstændig overensstemmelse mellem praktisk anvendelighed og evne til at klare sprogforståelsesopgaver. En model løser måske ikke en sprogforståelsesopgave særlig godt, men kan sagtens lave et meningsfuldt resumé af en tekst, fx, eller vurdere til perfektion om en anmeldelse er positiv eller negativ. Hele dette spørgsmål vil vi gerne undersøge nærmere på længere sigt, og en af vejene til det er at koble vores datasæt med andre mere taskdrevne datasæt udviklet fx i danske virksomheder.

TAK

Tak til Simon Gray, Center for Sprogteknologi, Institut for Nordiske Studier og Sprogvidenskab, KU, for at udvikle automatiseringsrutiner til DanNet-udtræk og udformning af udsagn til inferensdatasættet.

Bolette Sandford Pedersen, professor
Center for Sprogteknologi, Institut for Nordiske Studier og Sprogvidenskab
Københavns Universitet
bspedersen@hum.ku.dk

Nathalie C. Hau Sørensen, assisterende redaktør
Det Danske Sprog- og Litteraturselskab
nats@dsl.dk

Sussi Olsen, videnskabelig medarbejder
Center for Sprogteknologi, Institut for Nordiske Studier og Sprogvidenskab
Københavns Universitet
saolsen@hum.ku.dk

Sanni Nimb, seniorredaktør
Det Danske Sprog- og Litteraturselskab
sn@dsl.dk

LITTERATUR

- Berdicevskis, A. et al. 2023. Superlim: A Swedish language understanding evaluation benchmark. *Proceedings of the Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 8137–8153. Singapore: Association for Computational Linguistics. DOI:10.18653/v1/2023.emnlp-main.506.
- Firth, J.R. 1957. A synopsis of linguistic theory, 1930–1955. *Studies in linguistic analysis*. Oxford: Blackwell.
- Henrichsen, P.J. 2021. Det Centrale OrdRegister, Del 1. Oplæg ved Sprogteknologisk Konference 2021. Københavns Universitet. <https://cst.ku.dk/kalender/sprogteknologisk-konference-2021/>
- Herscovich, D. et al. 2022. Challenges and strategies in cross-cultural NLP. *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics* (Volume 1: Long papers), 6997–7013. Dublin: Association for Computational Linguistics. DOI: 10.18653/v1/2022.acl-long.482.
- Hjorth E. & K. Kristensen (red.). 2003–2005. *Den Danske Ordbog*. København: Det Danske Sprog- og Litteraturselskab & Gyldendal. <https://ordnet.dk/ddo>.
- Lee T. & S. Trott. 2023. A jargon-free explanation of how AI large language models work. *Ars Technica*. <https://arstechnica.com/science/2023/07/a-jargon-free-explanation-of-how-ai-large-language-models-work/>
- Lenci, A. & M. Sahlgren. 2023. *Distributional semantics*. Cambridge: Cambridge University Press. DOI: 10.1017/9780511783692.
- Levin, B. 1991. *English verb classes and alternations – a preliminary investigation*. Denver: University of Colorado Press.
- Nielsen, D.S. 2023. ScandEval: A benchmark for Scandinavian natural language processing. *Proceedings of the 24th Nordic Conference on Computational Linguistics (NoDaLiDa)*, 185–201, Tórshavn: University of Tartu Library. <https://aclanthology.org/2023.nodalida-1.20>
- Nielsen, F.Å. & L.K. Hansen. 2017. Open semantic analysis: The case of word level semantics in Danish. *Human language technologies. Challenges for computer science and linguistics*, 415–419. Mannheim: Springer.
- Nimb, S. et al. 2014. *Den Danske Begrebsordbog*. Odense: Det Danske Sprog- og Litteraturselskab & Syddansk Universitetsforlag.
- Nimb, S. 2016. Der er ikke langt fra tanke til handling. *Danske Studier* 2016, 25–59.
- Nimb, S. et al. 2017. From thesaurus to framenet. Electronic lexicography in the 21st century. *Proceedings of eLex 2017*, 1–22. Brno: Lexical Computing CZ s.r.o. <https://elex.link/elex2017/wp-content/uploads/2017/09/paper01.pdf>.

- Nimb, S., N.H. Sørensen & T. Troelsgård. 2018. From standalone thesaurus to integrated related words in the Danish dictionary. *Proceedings of the XVIII EURALEX International Congress: lexicography in global contexts*, 916–923. Ljubljana: Ljubljana University Press, Faculty of Arts. <https://euralex.org/publications/from-standalone-thesaurus-to-integrated-related-words-in-the-danish-dictionary/>.
- Nimb, S. et al. 2022a. COR-S – den semantiske del af Det Centrale OrdRegister (COR). *LexicoNordica* 29. <https://tidsskrift.dk/lexn/article/view/134776>.
- Nimb, S. et al. 2022b. A thesaurus-based sentiment lexicon for Danish: The Danish Sentiment Lexicon. *Proceedings of the 13th Language Resources and Evaluation Conference (LREC2022)*, 2826–2832. Marseille: European Language Resources Association. <https://aclanthology.org/2022.lrec-1.302>.
- Pedersen, B.S. et al. 2009. DanNet: the challenge of compiling a wordnet for Danish by reusing a monolingual dictionary. *Language resources and evaluation* 43. 269–299. DOI: 10.1007/s10579-009-9092-1.
- Pedersen, B.S. et al. 2016. The Semdax corpus. Sense annotations with scalable sense inventories. *Proceedings of the Tenth International Conference on Language Resources and Evaluation (LREC 2016)*, 842–847. Paris: European Language Resources Association. <https://aclanthology.org/L16-1136>.
- Pedersen, B.S., S. Nimb & S. Olsen. 2021. Dansk betydningsinventar i et datalingvistisk perspektiv. *Danske Studier* 2021. 72–106. <https://danskstudier.dk/wp-content/uploads/2021/07/danske-studier-2021.pdf>.
- Pedersen, B.S. et al. 2022. Compiling a suitable level of sense granularity in a lexicon for AI purposes: the open source COR-Lexicon. *Proceedings of the 13th Language Resources and Evaluation Conference (LREC2022)*, 51–60. Marseille: European Language Resources Association. <https://aclanthology.org/2022.lrec-1.6>.
- Pedersen, B.S. et al. 2023. The DA-ELEXIS corpus - a sense-annotated corpus for Danish with parallel annotations for nine European languages. *Proceedings of the Second Workshop on Resources and Representations for Under-Resourced Languages and Domains (RESOURCEFUL-2023)*, 11–18. Tórshavn: Association for Computational Linguistics. <https://aclanthology.org/2023.resourceful-1.2>.
- Pedersen, B.S. et al. 2024. Towards a Danish semantic reasoning benchmark - compiled from lexical-semantic resources for assessing selected language understanding capabilities of large language models. *Proceedings of the Fourteenth Language Resources and Evaluation Conference (LREC24)*. Torino: European Language Resources Association.

- Pilehvar, M.T. & J. Camacho-Collados. 2019. WiC: the word-in-context dataset for evaluating context-sensitive meaning representations. *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human language technologies, vol. 1*, 1267–1273. Minneapolis: Association for Computational Linguistics. DOI: 10.18653/v1/N19-1128.
- Pustejovsky, J. 1998. *The generative lexicon*. Cambridge: The MIT Press. DOI: 10.7551/mitpress/3225.001.0001
- Samuel, D. et al. 2023. NorBench – a benchmark for Norwegian language models. *Proceedings of the 24th Nordic Conference on Computational Linguistics (NoDa-LiDa)*, 618–633. Tórshavn: University of Tartu Library. <https://aclanthology.org/2023.nodalida-1.61>.
- Schneidermann, N.S., R. Hvingelby & B.S. Pedersen. 2020. Towards a gold standard for evaluating Danish word embeddings. *Proceedings of the 12th Language Resources and Evaluation Conference (LREC2020)*, 4754–4763. Marseille: European Language Resources Association. <https://aclanthology.org/2020.lrec-1.585/>.
- Tedeschi, S. et al. 2023. What’s the meaning of superhuman performance in today’s NLU? *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics* (vol. 1: long papers), 12471–12491, Toronto: Association for Computational Linguistics.
- Vossen, P. (red.). 1999. *EuroWordNet, A multilingual database with lexical semantic networks*. Amsterdam: Kluwer Academic Publishers. DOI: 10.1017/S1351324903223299
- Wang, A. et al. 2018. GLUE: A multi-task benchmark and analysis platform for natural language understanding. *Proceedings of the 2018 EMNLP Workshop Black-boxNLP: analyzing and interpreting neural networks for NLP*, 353–355. Brussels: Association for Computational Linguistics. DOI: 10.18653/v1/W18-5446.
- Wang, A. et al. 2020. SuperGLUE: A stickier benchmark for general-purpose language understanding systems. *Advances in neural information processing systems 32* (Proceedings, *NeurIPS 2019*). Vancouver. https://proceedings.neurips.cc/paper_files/paper/2019/hash/4496bf24afe7fab6f046bf4923da8de6-Abstract.html.

Kan AI reproducere fagdisciplinær stemme?

Et komparativt korpusbaseret studie af GPT4's evne til at reproducere fagdisciplinær stemme i AI-genereret sprogvidenskabelig prosa

**EA LINDHARDT OVERGAARD &
ULF DALVAD BERTHELTSEN**

ABSTRACT

Formålet med denne artikel er at afdække, i hvilket omfang generative AI-modeller – med GPT4 som eksempel – er i stand til at reproducere fagdisciplinær stemme i dansksproget akademisk prosa. De nye store sprogmodeller kommer med løfter om at forandre vores skrivepraksisser, herunder også akademisk skrivning, men det er stadig uklart, hvad kvaliteten er af de auto-genererede bidrag, ikke mindst når modellerne anvendes på mindre sprog som fx dansk. Vi er særligt interesserede i fænomenet fagdisciplinær stemme, fordi det er et relativt velbeskrevet fænomen, der samtidig kan undersøges kvantitativt gennem analyse af korpus teksters overfladestruktur. Vi fokuserer særligt på tre aspekter af fagdisciplinær stemme, henholdsvis stillingtagen, engagement og fagspecifikt ordforråd og undersøger dette kvantitativt gennem en korpusbaseret komparativ undersøgelse, hvor vi sammenligner et korpus bestående af dansksprogede sprogvidenskabelige artikler med et korpus af AI-genereret akademisk prosa med sprogvidenskabeligt indhold. Analysen viser, at de AI-genererede tekster på nogle områder afviger signifikant fra de autentiske sprogvidenskabelige tekster. For kategorien fagspecifikt ordforråd er forskellene relativt store, og for kategorierne stillingtagen og engagement er forskellene relativt små. I de to sidstnævnte kategorier er forskellene så små, at vi med en vis rimelighed kan sige, at de AI-genererede tekster på disse områder reproducerer fænomenet disciplinær stemme på en måde, der fra et kvantitativt perspektiv er svært at skelne fra det, vi ser i de autentiske tekster.

EMNEORD: generativ kunstig intelligens; chatbots; tekstkvalitet; fagdisciplinær stemme; korpuslingvistik

1 INDLEDNING

De nye store AI-baserede sprogmodeller som fx Open AI's GPT4 og Googles PaLM kommer med løfter om at forandre såvel vores private som vores professionelle skrivepraksisser, herunder akademisk skrivning. Hvis løfterne skal indfris, giver det sig selv, at det output, modellerne genererer, må være af høj sproglig kvalitet. Modellerne er imidlertid stadig nye, og det er endnu uklart, hvad kvaliteten er af de autogenererede tekster, ikke mindst når modellerne anvendes på mindre sprog som fx dansk (Eke 2023, Ghosh & Caliskan 2023). Der er derfor god grund til at interessere sig for, hvordan disse modeller præsterer, når de skal generere dansksproget output. Det er dog ingen trivial opgave at vurdere modellernes tekstkvalitet. Uanset om man arbejder inden for den Halliday-inspirerede SFL-tradition (Halliday 1985, Halliday & Hasan 1976) eller inden for den danske pragmatisk orienterede grammatiktradition (Hansen & Heltoft 2019, Harder 2006, Togeby 2003, Togeby 2014), er det en central pointe, at de sproglige valg, en afsender træffer, i vid udstrækning træffes relativt til kommunikationssituationen, herunder såvel den konkrete kommunikationssituation som den bredere kulturelle kontekst, hvori kommunikationssituationen er indlejret. Den konkrete sproglige realisering af tekster forventes altså at variere, som en konsekvens heraf, med hensyn til fx ordforråd, tekstlængde, syntaktisk kompleksitet og stilleje. De store sprogmodeller er probabilistiske modeller baseret på den såkaldte transformerarkitektur (Vaswani m.fl. 2017). Dette vil lidt forenklet sige, at det output, de genererer, er den mest sandsynlige ordkombination relativt til prompt og træningsdata, samtidig med at modellerne på tekstniveau er i stand til at forholde sig til både den globale og lokale kontekst. Fordi modellerne er trænet på meget store mængder af naturligt sprog, må vi forvente, at den genre- og registervariation, vi finder i det naturlige sprog, afspejles i de data, modellerne er trænet på. Som følge heraf må vi også forvente, at denne variation i et eller andet omfang reproduceres i det autogenererede output, idet sandsynligheden for, at et givet træk optræder i det autogenererede output, groft sagt er baseret på, hvor ofte og i hvilke kontekster dette træk optræder i modellens træningsdata. Kvaliteten af outputtet kan derfor ikke undersøges generisk, men må nødvendigvis undersøges relativt til kontekst og genre. Sagt på en anden måde kan

vi ikke nøjes med at stille det generelle spørgsmål: Er GPT4 god til at skrive dansk? Vi må i stedet spørge mere specifikt, om GPT4 fx er god til at skrive jobansøgninger på dansk, eller om den er god til at skrive festtaler på dansk.

Som et bidrag til at undersøge kvaliteten af autogenereret sprog fokuserer vi i denne artikel på akademisk skrivning. Vi bruger sprogvidenskabelige tekster som eksempel og fokuserer særligt på fænomenet *fagdisciplinær stemme* (Hyland 2005, 2008). Tekstkvalitet er naturligvis et komplekst fænomen, men vi har valgt at fokusere på fagdisciplinær stemme, fordi det er et relativt velbeskrevet fænomen, der samtidig kan undersøges kvantitativt gennem analyse af korpussteksters overfladestruktur (Fløttum m.fl. 2006, Guinda 2012). Formålet med undersøgelsen er at afdække, i hvilket omfang generative AI-modeller, med GPT4 som eksempel, er i stand til at reproducere fagdisciplinær stemme i dansksproget sprogvidenskabelig prosa. I artiklen præsenterer vi en kvantitativ korpusbaseret komparativ undersøgelse, hvor vi sammenligner et korpus bestående af autentiske dansksprogede sprogvidenskabelige artikler med et korpus af AI-genereret akademisk prosa med sprogvidenskabeligt indhold. Korpusset af autentiske sprogvidenskabelige artikler består af 306 artikler, og det AI-genererede korpus består af 300 automatisk genererede artikelfragmenter. Vi fokuserer særligt på to aspekter af fagdisciplinær stemme, henholdsvis stillingtagen og engagement, samt på anvendelsen af fagspecifikt ordforråd, og fænomenerne operationaliseres som en række sproglige træk, herunder bl.a. brugen af pronomener, modalitet og adverbelle udtryk. I forbindelse med sammenligningen benyttes permutationstest for at fastslå, om der er signifikant forskel i forekomsten af de forskellige typer af stemmemarkører mellem de to korpusser.

Artiklen indledes med en præsentation af undersøgelsens teoretiske afsæt samt en redegørelse for operationaliseringen af fagdisciplinær stemme. Herefter følger et metodeafsnit, hvor vi redegør for opbygningen af de to korpusser samt rationalet bag valget af statistiske test. Dette følges af en gennemgang af analysens resultater, og artiklen afsluttes med en diskussion af resultaternes relevans og rækkevidde.

2 BAGGRUND

De seneste år er spørgsmålet om kvaliteten af AI-genereret tekst og anvendelsen af generativ kunstig intelligens blevet stadig mere vigtigt, idet generative AI-værktøjer baseret på store sprogmodeller nu er så veludviklede, at de kan bruges som både kreative sparringspartnere og som værktøjer i forskningsprojekter og endda som medforfattere på videnskabelige tekster. Den videnskabelige debat vedrørende denne udvikling handler især om AI-modellernes evne til at producere sammenhængende og meningsfulde tekster på tværs af forskellige domæner samt dels etiske problemstillinger, dels problematikker relateret til inddragelsen af kunstig intelligens som en aktiv medspiller i forskningsverdenen.

I forskningslitteraturen er debatten fortrinsvis fagdisciplinært funderet, idet brugen af kunstig intelligens oftest undersøges i relation til specifikke videnskabelige domæner. Særligt det sundhedsvidenskabelige område (Clusmann m.fl. 2023, Peng m.fl. 2023), litteraturforskning med fokus på kreativ skrivning (Gunser m.fl. 2022, Hitsuwari m.fl. 2023, Jensen & Sørhaug 2021, Köbis & Mossink 2021) samt generelle akademiske skriveproblematikker (Anderson m.fl. 2023, Dergaa m.fl. 2023, Hinton & Wagemans 2023) fylder i forskningslitteraturen. Disse studier kredser på forskellig vis om især tre overordnede spørgsmål: 1) Om kvaliteten af AI-genereret tekst er sammenlignelig med tekster skrevet af mennesker, 2) om man kan påvise, at en tekst er skrevet ved brug af kunstig intelligens ved at analysere forskellige tekstuelle og lingvistiske aspekter af teksterne, og 3) om AI-modeller på fornuftig vis kan benyttes både til læring for elever og studerende og som værktøj til forskning og på arbejdsmarkedet. Alle er dog enige om, at AI-generering af tekster kræver forbedringer på en række områder, hvis AI-værktøjer skal kunne inddrages i alle aspekter af akademiske og litterære skriveprocesser. Endelig peges der på, at den manglende kvalitet i AI-genereret tekst på nuværende tidspunkt skyldes, at AI-modellerne ikke er trænet på domænespecifikke områder, hvilket bl.a. påvirker modellernes evne til at efterligne menneskelig kreativitet i kunstneriske genrer som fx lyrik (Jensen & Sørhaug 2021: 17). Desuden peges der på, at det er problematisk, at AI-genereret tekst ikke automatisk faktatjekkes, og at modellerne ikke forholder sig kritisk til det indhold,

teksterne er baseret på (Peng m.fl. 2023, Frye 2022). Endelig kritiseres AI-genereret tekst for at mangle afsenderperspektiver (Frye 2022).

Selvom forskningslitteraturen om anvendelse af generativ kunstig intelligens er voksende, er der stadig kun få komparative studier, der sammenligner kvaliteten af AI-generede tekster med tekster skrevet af mennesker. Der argumenteres imidlertid i stigende grad for nødvendigheden af denne type studier (Gunser m.fl. 2022, Hinton & Wagemans 2023, Hitsuwari m.fl. 2023), og vi forsøger med denne artikel at bidrage til feltet med en undersøgelse af GPT4's evne til at generere dansksproget sprogvidenskabelig prosa.

3 FAGDISCIPLINÆR STEMME

Begrebet *stemme* er ikke entydigt defineret, men anvendes om forskellige fænomener i forskellige videnskabelige sammenhænge. Inden for fonetikken interesserer man sig fx for stemmen som fysiologisk og akustisk fænomen (Grønnum 1998), inden for litteraturteorien beskæftiger man sig med stemme som narrativt fænomen (Chatman 1978, Rimmon-Kenan 1983, Wall 1991), og inden for pædagogikken beskæftiger man sig med stemme som et literacy-fænomen relateret til identitet, demokratisk deltagelse og udvikling af skrivekompetencer (MacBeath 2006, Krogh 2012). I denne artikel trækker vi på den funktionelt orienterede forskning i fagdisciplinær skrivning (se fx Biber 2006, Biber & Gray 2015, Hyland 2000, Hyland 2017), herunder særligt Hyland (2005, 2008), der knytter fænomenet stemme til *stillingtagen* (stance) og *engagement* (engagement). Stemme forstås i denne sammenhæng som noget, der får sit tekstlige udtryk i og med afsenderens anvendelse af kommunikative ressourcer med det formål at repræsentere sig selv, sin position og sin relation til andre, herunder modtageren (Hyland 2008: 20). Eksempler på kommunikative ressourcer, der typisk anvendes med dette formål, kunne være modalverber og attitudeadverbier, der bl.a. anvendes til at etablere forpligtende relationer mellem afsender og modtager (Hansen & Heltoft 2019 kap. VI, § 12-13) og til at angive afsenderens holdning til det fremsatte sagforhold (Berthelsen 2007). I forlængelse heraf opfattes stemme som et fænomen knyttet til register og genre. Anvendelsen af kommunikative ressourcer til at udtrykke stemme er således ikke

først og fremmest et stilistisk udtryk for afsenderens kreativitet eller idiosynkrasier, men derimod udtryk for sproglige valg truffet relativt til formål og kommunikationskontekst. Dette giver anledning til den antagelse, at de mønstre, vi ser i anvendelsen af stemme i fagdisciplinær skrivning, i vid udstrækning er udtryk for, at afsenderen tilstræber en tilnærmelse til normer og praksisser accepteret af de fagdisciplinære fællesskaber (Biber & Conrad 2009, Togeby 2014). Som følge heraf kan fænomenet som allerede nævnt ikke undersøges generisk. Af denne grund fokuserer vi i denne artikel på vores egen metier, nemlig akademisk skrivning. Denne indsnævring er imidlertid ikke tilstrækkelig, idet der er store forskelle på skrivekulturen inden for de forskellige akademiske discipliner (Carter 2007, Hyland 2005). For at indsnævre feltet yderligere har vi derfor valgt at fokusere særligt på det sprogvidenskabelige domæne.

Stemme er det vi med Nordahl-Hansen & Kvernbekk (2020) kan kalde en T-term, dvs. en betegnelse for et teoretisk konstrukt, der ikke kan observeres direkte. Det er som følge heraf nødvendigt at fortage en operationalisering af konstruktet i form af en specificering af, hvilke observerbare træk i teksten, der antages at være udtryk for stemme. I sproget kan stemme udtrykkes på en mangfoldighed af måder, hvilket er en udfordring i relation til operationaliseringen, når man ønsker at undersøge fænomenet kvantitativt. For det første må der være tale om træk ved teksten, der kan identificeres ved hjælp af de NLP-værktøjer (Natural Language Processing), der er til rådighed for et givet sprog. Dette forhold vender vi tilbage til i metodeafsnittet. For det andet må der være tale om træk ved teksten, der entydigt kan identificeres. Som eksempel på denne problematik kan nævnes udtrykket 'sikkert'. Det kan dels optræde som både adjektiv og adverbium, dels anvendes på måder, hvor det er den tekstlige kontekst, der afgør, om det har en forstærkende eller modererende effekt. Sammenlign fx brugen af 'sikkert' i 'hun kommer helt sikkert ikke', hvor 'helt' + 'sikkert' har en forstærkende effekt, med 'det er ikke sikkert, at hun kommer', hvor 'ikke' + 'sikkert' har en modererende effekt. Pointen er, at frekvensen af 'sikkert' i de to korpusser ikke i sig selv siger noget om funktionen.

For at kunne balancere mellem disse to kriterier læner vi os op ad

Hyland (2005). Hyland skelner som allerede nævnt mellem to overordnede dimensioner, henholdsvis stillingtagen og engagement. *Stillingtagen* betegner den holdningsmæssige dimension (attitudinal dimension) repræsenteret ved det, man kan kalde tekstens afsenderorienterede træk, mens *engagement* betegner den relationsmæssige dimension (alignment dimension) repræsenteret ved det, man kan kalde tekstens modtagerorienterede træk. Den holdningsmæssige dimension angår afsenderens arbejde med at positionere sig i forhold til det omtalte, hvorimod den relationsmæssige dimension angår afsenderens arbejde med at positionere sig i forhold til modtageren (Hyland 2005: 176).

Hylands analytiske apparat er omfattende, og det er ikke alle kategorier, der umiddelbart kan scores automatisk. Der er dog nogle, der kan, og af disse har vi udvalgt en række kategorier, der alle er prototypiske eksempler på kommunikative ressourcer, der udtrykker stemme. Som udtryk for stillingtagen har vi valgt at arbejde med fire forskellige kategorier: 1) *førstepersonspronomen*, der eksplicit markerer afsenderens tilstedeværelse i teksten, 2) *attitudeadverbialer*, der markerer afsenderens holdning til det omtalte sagforhold og 3) *karakteriserende adjektiver*, der udtrykker afsenderens vurdering af det omtalte. Herudover foretager vi 4) en *sentimentanalyse*. Sentimentanalysen er ikke en del af Hylands analytiske apparat, men er alligevel inkluderet, fordi den giver en indikation af, om en tekst er overvejende positivt eller negativt ladet. Som udtryk for engagement har vi valgt to forskellige kategorier: 5) *andersonspronomen*, der eksplicit indikerer en adressering af modtageren og 6) *modalverber*, der indikerer forskellige grader af engagement og forpligtelse mellem afsender og modtager. Udover disse seks kategorier har vi også valgt at inkludere 7) *klassificerende adjektiver* og 8) *fagspecifikke substantiver*. Disse to mål er ikke direkte knyttet til konstruktet fagdisciplinær stemme, men er alligevel medtaget som supplement til stemmemarkørerne, fordi de kan give en indikation af, i hvilket omfang afsenderen tilstræber en saglig og neutral tone eller en mere subjektiv og holdningsorienteret tone (Baumann & Graves 2010).

4 METODE

I det følgende redegør vi indledningsvis for opbygningen af de to tekstkorpusser, der udgør grundlaget for den sammenlignende analyse, her-

efter for den konkrete operationalisering af konstruktet *fagdisciplinær stemme* og endelig for valget af analysemetoder.

4.1 Opbygning af det autentiske sprogvidenskabelige korpus

Den ene halvdel af datamaterialet består af 306 autentiske sprogvidenskabelige artikler og indeholder i alt 1.116.259 ord. Artiklerne stammer fra konferencerapporterne fra konferencerækken MUDS – Møderne om Udforskningen af Dansk Sprog (www.projekter.au.dk/muds/om-muds). Vi har inkluderet alle dansksprogede artikler inklusive tilhørende abstracts fra årgangene 2000-2020. Udover dansksprogede artikler finder man i konferencerapporten også abstracts uden artikler samt enkelte artikler på svensk eller norsk. Disse er udeladt med henblik på at få et så homogent korpus som muligt. Artiklerne behandler et bredt udsnit af sprogvidenskabelige emner, og de er alle blevet fagfællebedømt inden udgivelse. Vi antager derfor, at dette udvalg, selvom det ikke er repræsentativt i formel forstand, giver et godt afsæt for at undersøge, hvordan fagdisciplinær stemme realiseres i dansksproget sprogvidenskabelig prosa.

4.2 Opbygning af det AI-genererede korpus

Den anden halvdel af datamaterialet er et korpus bestående af 300 AI-genererede tekstfragmenter genereret med OpenAI's store sprogmodel GPT4 (OpenAI 2023). Korpuset indeholder i alt 258.766 ord. I arbejdet med at opbygge det autogenererede korpus har vi forsøgt at balancere mellem på den ene side at generere en tilstrækkelig mængde tekst med henblik på at sikre kvaliteten af de kvantitative analyser, og på den anden side at tage hensyn til, at brugen af store sprogmodeller er relativt ressourcekrævende. For det første har GPT4 en begrænsning på, hvor meget tekst der kan genereres ad gangen, hvilket betyder, at det er relativt tidskrævende at producere mange tekster, og for det andet er brugen af store sprogmodeller meget energikrævende, bl.a. fordi processen er computationelt krævende i forhold til både anvendelsen af processorkraft og lagringsplads (Lauridsen m.fl. 2019, Bender m.fl. 2021). De AI-genererede tekstfragmenter er derfor kortere (3-4 sider per tekstfragment) end de autentiske sprogvidenskabelige artikler, men vi har vurderet, at dette ville være tilstrækkeligt til at sikre kvaliteten af de kvantitative analyser.

4.2.1 De AI-genererede tekstfragmenters struktur og indhold

På nuværende tidspunkt kan GPT4 ikke generere sammenhængende tekst, der i struktur og omfang minder om videnskabelige artikler. Med henblik på at gøre de to korpusser sammenlignelige, såvel strukturelt som indholdsmæssigt, valgte vi derfor at sammensætte hvert tekstfragment af tre selvstændige dele, henholdsvis et abstract, et beskrivende afsnit og et diskuterende afsnit. De enkelte dele blev genereret på baggrund af et sæt sprogvidenskabelige nøgleord, som vi genererede via det sprogvidenskabelige korpus ved hjælp af topic modelling, nærmere bestemt Latent Dirichlet Allocation (LDA) (Blei m.fl. 2003). LDA er en probabilistisk maskinlæringsmodel, der bruges til at afdække, hvilke emner der er indeholdt i en gruppe tekster, og på denne baggrund kategorisere teksterne relativt til, hvor godt emnerne karakteriserer de enkelte tekster. Dette sker ved at se på, hvor ofte ordene optræder sammen, samtidig med at der tages højde for, hvor unikke de enkelte ord er for et givent emne. For eksempel vil et ord som "ord" forekomme på tværs af mange tekster i vores korpus, da vi arbejder med sprogvidenskabelige artikler, og dette ord er derfor ikke emnespecifikt. Omvendt er et ord som "dialekt" knyttet til et specifikt område, og ordet vil være tæt relateret med en række andre ord inden for samme område. For at kunne afgøre, hvor mange emner det sprogvidenskabelige korpus skulle indeles i, evaluerede vi den såkaldte coherence-score for forskellige antal emner (Weston m.fl. 2023). Coherence-scoren fortæller os groft sagt, hvor meningsfulde emnerne er. Analysen viste, at vi fik færrest overlap, hvis vi valgte at inddele teksterne i seks emner. Som resultat af LDA-analysen får vi således seks emner karakteriseret ved en række særligt frekvente nøgleord i lemmaform. Disse nøgleord har vi brugt som afsæt for at prompte GPT4 til at generere tekster. Nøgleordene for de seks emner er følgende:

- Emne 1 = varietet, sproglig, dansk, dialekt, undersøgelse, kort, resultat, rigsdansk, forskellig, respondent
- Emne 2 = sætning, subjekt, dansk, flytning, passiv, ledsætning, kasus, eksempel, rolle, flytte
- Emne 3 = betydning, dansk, moderne, tekst, eksempel, dialogisk, proposition, nydansk, partikel, udtrykke

Emne 4 = dansk, sprog, tale, eksempel, sætning, forskellig,
tekst, form, spørgsmål, sted

Emne 5 = orddannelse, nederlandsk, kvindenavn, participial,
deverbael, amager-hollandsk, forstå-signaler, fri-
sisk, deverbale, amager-register

Emne 6 = navn, fornavn, lemma, stavemåde, dansk, form, antal,
forskellig, sjælden, nyfødt

4.2.2 Prompting

I AI-sammenhæng er en prompt en instruks skrevet i naturligt sprog (i modsætning til i computerkode), der fortæller den generative model, hvilket output man ønsker. I instruksene til GPT4 har vi specificeret 1) at teksten skal bestå af henholdsvis et abstract, et redegørende afsnit og et diskuterende afsnit, 2) hvilken teoretisk ramme teksten skal skrives ud fra og 3) en specifikation af de ti nøgleord. Instruksene sendes herefter til OpenAI's API, hvorefter GPT4-modellen returnerer tre separate tekstsekvenser i form af et abstract, en redegørende tekstsekvens og en diskuterende tekstsekvens. Prompten for et tekstfragment ser ud på følgende måde. For hver sekvens specificeres systemets rolle, herefter det ønskede output:

```
input_1 = [{"role": "system", "content": "sprogforsker der skriver sprogvidenskabelige artikler på dansk"}, {"role": "user", "content": "skriv et videnskabeligt abstract på dansk om syntaktisk flytning, den teoretiske ramme er generativ grammatik, keywords: sætning, subjekt, dansk, flytning, 'passiv, ledsætning, kasus', eksempel, rolle, flytte"}]
```

```
input_2 = [{"role": "system", "content": "sprogforsker der skriver sprogvidenskabelige artikler på dansk"}, {"role": "user", "content": "skriv et videnskabeligt teoriafsnit om syntaktisk flytning, den teoretiske ramme er generativ grammatik, keywords: sætning, subjekt, dansk, flytning, passiv, ledsætning, kasus, eksempel, rolle, flytte, præsenter relevante teoretiske begreber"}]
```

```
input_3 = [{"role": "system", "content": "sprogforsker der skriver sprogvidenskabelige artikler på dansk"}, {"role": "user", "content": "skriv en videnskabelig diskussion om syntaktisk flytning, den teoretiske ramme er generativ grammatik, keywords: sætning, subjekt, dansk, flytning, passiv, ledsætning, kasus, eksempel, rolle, flytte, diskuter fordelene ved relevante teoretiske positioner"}]
```

Med henblik på at opbygge et korpus af passende størrelse blev modellen bedt om at generere 50 tekstfragmenter per instruks. Fordi der er tale om en probabilistisk model, genereres der som følge heraf 50 forskellige tekstfragmenter baseret på samme nøgleord. Når anmodningen om at generere et tekstoutput sendes til API'en, kan man specificere en såkaldt temperatur, hvilket er en indikation af, hvor konservativt modellen gætter. Jo lavere temperatur, jo mere konservativt. I forbindelse med denne undersøgelse valgte vi den højst mulige temperatur med henblik på at sikre en vis sproglig variation i det autogenererede korpus.

4.3 Præprocessering

I det kvantitative arbejde med vores tekstkorpus har det været en forudsætning, at teksterne er maskinlæsbare, rensat for støj og organiseret hensigtsmæssigt. Vi har derfor præprocesseret de rå data (Ondelli 2018: 143-148). Al præprocessering er foretaget ved hjælp af basisbiblioteker i programmeringssproget Python (Van Rossum & Drake 2009), medmindre andet er nævnt.

De sprogvidenskabelige artikler blev præprocesseret på flere niveauer. Efter indsamlingen af artiklerne, som var tilgængelige i tekstscannede pdf-filer, har vi udtrukket tekstscanningslaget og gemt hver artikel separat som rene tekstfiler (.txt). Tekstfilerne er herefter blevet rensat for forstyrrende elementer i tekstscanningslaget, fx sidehoved og -fodder. Vi har efterfølgende fjernet længere citater og transskriptioner. Vi har fjernet disse elementer, da det er svært at skelne mellem stemmemarkører, der er relateret til forfatteren og stemmemarkører benyttet i teksteksampler indeholdt i artiklen, når vi benytter digitale værktøjer i analysen. Korte citater i brødteksten har det ikke været muligt at fjerne. De rensede artikler er efterfølgende samlet i et korpus, hvor metadata som

titel og hvilken udgave af konferencerapporten, artiklen kommer fra, er tilføjet. Hver artikel i korpuset er repræsenteret som én lang streng af tegn. De nødvendige opdelinger i sætninger og tokens foretages først i forbindelse med de enkelte analyser.

De AI-genererede tekster er født maskinlæsbare, idét de er genereret digitalt. Outputtet fra tekstgenereringen er delt i tre dele, et abstract, en redegørende tekst og en diskuterende tekst. Outputtet fra GPT4 er i filformatet JSON (.json), hvor informationer om input og forskellige parametre for instruksen til GPT4 også er indeholdt. Vi har derfor udtrukket selve tekstsekvenserne fra JSON-filerne og derefter gemt dem i txt-format. Ligesom for de sprogvidenskabelige artikler renses de auto-genererede tekster for forstyrrende elementer, i dette tilfælde primært for formateringskoden "\n" (linjeskift). Desuden har vi forenet de tre tekstelementer (abstract, redegørende afsnit og diskuterende afsnit) til én samlet streng af tegn. Herefter har vi samlet tekstfragmenterne i et korpus, hvor prompten samt JSON-filen er tilføjet som metadata. Vi har afslutningsvis forenet de to korpusser til ét samlet korpus, der indeholder samtlige 606 tekster. Da vi i analysen benytter NLP-værktøjer, der tager en samlet streng af tegn som input, udgør dette korpus det endelige datasæt og er udgangspunkt for de følgende analyser. En del af analyserne kræver dog yderligere præprocessering, men dette beskrives nærmere i afsnittet for de enkelte analyser, idet fx opdeling i tokens eller sætninger er inkluderet i den NLP-pipeline vi benytter. De forskellige korpusser og det endelige datasæt er konstrueret ved hjælp af Pandas-biblioteket (McKinney 2010) i Python.

4.4 *Design af analyse*

Vores analyse er frekvensbaseret, hvilket vil sige, at vi for hver tekst tæller en række sproglige træk, der antages at repræsentere konstruktet stemme. Trækkene knytter sig til enten afsenderens *stillingtagen* eller *engagement* eller *anvendelse af fagspecifikt ordforråd*. Inden for stillingtagen ser vi på brugen af førstepersonspronomen, attitudeadverbialer, karakteriserende adjektiver og den emotionelle værdi i teksterne, og trækkene vi knytter til engagement, er andenpersonspronomen og modalverber. Den faglige sprogbrug undersøger vi gennem brugen af klassificerende adjektiver og fagspecifikke substantiver.

For at kunne sammenligne teksterne i de to korpusser har vi udregnet, hvor store andele hver kategori udgør af den samlede mængde ord inden for ordklassen, fx andelen af førstepersonspronomen relativt til det samlede antal pronomen i teksten. Vi har herefter beregnet et gennemsnit for hvert korpus. Vi sammenligner desuden spredningen mellem de to grupper med henblik på at undersøge, i hvilket omfang fordelingen i de to grupper minder om hinanden. Alle analyserne, bortset fra sentimentanalysen, er baseret på en annotering af teksterne foretaget ved hjælp af DaCy (Enevoldsen m.fl. 2021). DaCy er et NLP-værktøj baseret på maskinlæring, der er optimeret til brug på dansksprogede tekster. DaCy er valgt, fordi modellen har opnået et state-of-the-art-performanceniveau på flere parametre, blandt andet på POS-tagging (parts of speech) (Enevoldsen m.fl. 2021: 210). DaCy tager de rå tekstdata som input og returnerer komplekse dataobjekter, der indeholder en lang række metadata om teksten, herunder POS-tags og information om tekstens syntaktiske struktur i form af såkaldte *universal dependency tags* (se <https://universaldependencies.org>), der fx indikerer om et verbum er et hjælpeverbum, og om et nominal er subjekt eller objekt. Fremgangsmåden for de frekvensbaserede analyser gennemgås i følgende afsnit, og metoden for evaluering af analyserne gennemgås efterfølgende. Vi gennemgår metoden for den frekvensbaserede analyse i samme rækkefølge som de analytiske kategorier gennemgås i resultatafsnittet, bortset fra at metoden for kategorierne første- og andenpersonspronomen er slået sammen, ligesom klassificerende og karakteriserende adjektiver er slået sammen under ét afsnit. Disse kategorier er slået sammen her, da vi rent teknisk benytter samme fremgangsmåde i forbindelse med analyserne i de pågældende kategorier.

4.4.1 Første- og andenpersonspronomen

Vi har undersøgt brugen af første- og andenpersonspronomen, fx 'jeg', 'mig', 'vi' og 'du', 'dig', 'I', ved at udtrække en liste for hver tekst, der indeholder samtlige førstepersonspronomen i teksten, og ligeledes udtrække en liste, der indeholder samtlige andenpersonspronomen. Til at måle andelen af de to pronomentyper for hver tekst, har vi udtrukket en liste med samtlige pronomen for hver tekst. Denne liste har vi udtrukket ud fra en opslagsliste indeholdende samtlige pronomen

baseret på DaCy's annotation med POS-tagget "PRON". Ud fra disse lister har vi optalt frekvensen af pronomenerne fra hver gruppe af pronomener og udregnet andelen af førstepersonpronomener og andenpersonpronomener for hver gruppe af tekst. Afslutningsvis har vi beregnet en gennemsnitsandel for hvert af de to korpusser.

4.4.2 *Attitudeadverbialer*

I forbindelse med undersøgelsen af attitudeadverbialerne har vi dels søgt efter attitudeadverbier, dels søgt på bi- og trigrammer, dvs. ordforbindelser bestående af to eller tre ord som fx '*uden tvivl*' og '*ikke med sikkerhed*'. I analysen af adverbier har vi benyttet os af POS-taggene fra annoteringen med DaCy. Da DaCy ikke kan skelne mellem forskellige typer af adverbier, har vi først udtrukket en liste med de 200 mest frekvente adverbier. På denne liste har vi manuelt opmærket alle de adverbier, der er attitudeadverbier, fx '*formentlig*' og '*selvfølgelig*'. For at finde attitudeadverbialer har vi fundet bi- og trigrams i teksterne og ved hjælp af en stopordsliste fjernet ikke-relevante resultater såsom '*sproglig variation*', '*det er*' m.fl. Herefter har vi ud fra lister med de 200 hyppigste bigrams og de 200 hyppigste trigrams manuelt opmærket attitudeadverbialerne. Efterfølgende har vi samlet attitudeadverbier og -adverbialer på en liste og på baggrund af denne liste optalt frekvensen af attitudeadverbialer i hver tekst. Vi har udregnet andelen af attitudeadverbialer i forhold til den samlede mængde adverbier i hver tekst, selvom adverbialerne strækker sig ud over adverbiumkategorien. Dette har ingen betydning for resultaterne, da normaliseringen benyttes til sammenligning på tværs af de to korpusser. Vi har desuden udregnet den gennemsnitlige andel for hvert af de to korpusser.

4.4.3 *Karakteriserende og klassificerende adjektiver*

I analysen af adjektiverne undersøger vi frekvensen af henholdsvis karakteriserende og klassificerende adjektiver i de to korpusser. Vi benytter termene karakteriserende og klassificerende adjektiver i overensstemmelse med Christensen & Christensen (2019: 75-76). De karakteriserende adjektiver anvendes til at angive egenskaber ved en referent, mens klassificerende adjektiver angiver, at den beskrevne referent tilhører en særlig klasse eller et særligt område. For at kunne lave en opdeling af

karakteriserende og klassificerende adjektiver har vi kigget på den samlede gruppe af adjektiver i vores tekstkorpus. Vi har på baggrund af POS-taggene fra annoteringen med DaCy udtrukket en liste med de 200 mest frekvente adjektiver i datasættet. På denne liste har vi manuelt markeret de adjektiver, der er karakteriserende, fx 'vigtig', 'god' og 'central', og de adjektiver der er klassificerende, fx 'sproglig', 'grammatisk' og 'dansk'. På denne baggrund har vi beregnet andelen af begge typer adjektiver per tekst relativt til det samlede antal adjektiver. Den gennemsnitlige andel af karakteriserende og klassificerende adjektiver er efterfølgende beregnet for hvert korpus.

4.4.4 Modalverber

Brugen af modalverber er som i de foregående tilfælde beregnet som en andel relativt til det samlede antal verber i teksten, hvorefter der er beregnet et gennemsnit for hvert af de to korpusser. Vi har optalt antallet af modalverber i hver tekst ved at søge dem frem ud fra en manuelt opskrevet liste af modalverber. Herefter har vi holdt dette antal op mod det samlede antal verber i hver tekst, som er optalt med udgangspunkt i annoteringen af teksterne. Det samlede antal verber er fundet ved at summere antallet af verber med POS-taggene 'AUX' eller 'VERB', som er hjælpeverberne og de resterende verber.

4.4.5 Fagspecifikke substantiver

Listerne med fagspecifikke substantiver er udtrukket på samme måde som listerne med henholdsvis attitudeadverbier og karakteriserende adjektiver. Først har vi udtrukket en liste over de 200 mest frekvente substantiver i det samlede korpus. Herefter har vi manuelt opmærket de fagspecifikke substantiver. Fagspecifik fortolkes i denne sammenhæng som substantiver tilhørende det sprogvidenskabelige domæne (fx 'syntaks', 'sætning', 'fonem'). Herefter har vi først beregnet andelen af fagspecifikke substantiver per tekst relativt til tekstens samlede antal substantiver og herefter en gennemsnitlig andel for hvert af de to korpusser.

4.4.6 Sentiment

Sentimentanalyse er en kvantitativ måde at måle teksters emotionelle værdi baseret på en antagelse om, at ord har en emotionel værdi. Til

at lave analysen har vi benyttet Python-biblioteket Sentida (Lauridsen m.fl. 2019). Sentida er et af de få værktøjer, der er designet til at lave sentimentanalyse på dansksprogede tekster. Sentida har en præcision på over 80 %, hvilket er højere end andre danske sentiment-værktøjer, der i øjeblikket er tilgængelige. Sentida virker ved at 1) tage en sammenhængende tekst og dele op i ord, 2) gennemgå hvert ord og tildele ordet en sentimentscore baseret på et leksikon over sentiment-bærende ord, 3) tjekke for faktorer, der kan ændre ved sentimentværdien, fx negationer, udråbstegn eller adverbierle modifikatorer, og 4) returnere den gennemsnitlige sentimentscore, dvs. alle ords sentimentværdi summeret og divideret med antallet af ord, der bidrager til sentimentværdien (Lauridsen m.fl. 2019: 45). Fx tilskrives 'god' en positiv værdi, mens 'dårlig' tilskrives en negativ værdi samtidig med, at det sikres, at 'ikke dårligt' tilskrives en positiv værdi. Ved at returnere gennemsnittet tages der højde for variation i tekstlængderne. Sentimentværdierne ligger mellem -5 og 5 og er kategoriseret som vist i følgende tabel 1 (Lauridsen m.fl. 2019: 41):

TABEL 1. SKALA FOR SENTIMENTVÆRDIER

Very strong negative emotion	Strong negative emotion	Slightly negative emotion	Weak negative emotion	Very weak negative emotion	Neutral emotion	Very weak positive emotion	Weak positive emotion	Slightly positive emotion	Strong positive emotion	Very strong positive emotion
-5	-4	-3	-2	-1	0	1	2	3	4	5

(Lauridsen m.fl. 2019)

4.5 Statistisk analyse

Med henblik på at beskrive eventuelle forskelle i brugen af stemmemarkører mellem de autentiske og de AI-generede tekster sammenligner vi de to korpusser ved dels at lave en deskriptiv statistisk analyse, dels at foretage en permutationstest.

4.5.1 Deskriptiv statistisk analyse

For hver af de otte analysekategorier har vi som ovenfor beskrevet beregnet en gennemsnitlig frekvensværdi for hvert af de to korpusser, henholdsvis det autentiske og det AI-genererede. Derudover har vi for hver kategori også undersøgt spredningen i observationerne ved at beregne

standardafvigelsen. Gennemsnittene indikerer, hvilke værdier vores observationer er centreret omkring, og siger noget om, hvor hyppigt hvert af de sproglige træk benyttes i de to grupper af tekster. Standardafvigelsen indikerer, hvor meget observationerne er spredt ud i forhold til gennemsnittet.

Ved at sammenstille gennemsnit og spredning i analysen får vi både et billede af, om hyppigheden af den observerede størrelse i gruppen af AI-genererede tekster er på niveau med de sprogvidenskabelige artikler, og om afstanden mellem observationerne i tekster over og under gennemsnittene er forskellig imellem de to grupper af tekster. Vi kan altså sige noget om, hvorvidt variationen i brugen af et sprogligt træk ligner den naturlige variation, man finder i de autentiske sprogvidenskabelige artikler, og ligeledes om spredningen centrerer sig om det samme midtpunkt i de to grupper af tekster.

Som afsæt for sammenligningen af de frekvensbaserede resultater i analysen benyttes histogrammer til umiddelbar visualisering af de numeriske resultater. Histogrammerne viser fordelingen af målingerne inden for hver analysekategori. Man kan altså ud fra histogrammerne se, hvor mange tekster, der har en bestemt andel af den givne analysekategori, og hvor stor afstand der er mellem de højeste og laveste andele i hver analysekategori. Ud fra histogrammerne kan vi se hvordan data er fordelt, og om der er skævhed i fordelingerne. Er et histogram pænt symmetrisk omkring den gennemsnitlige værdi, er gruppen tilnærmelsesvis normalfordelt omkring gennemsnittet, hvorimod skæve histogrammer indikerer, at der ikke er tale om en normalfordeling. Det er interessant at se, om fordelingerne for gruppen af AI-genererede tekster og gruppen af sprogvidenskabelige ligner hinanden, idet sammenligningen giver en indikation af, hvor godt GPT4 kan reproducere den menneskelige brug af det givne sproglige træk, herunder i hvilket omfang variationen på tværs af tekster reproduces.

4.5.2 Permutationstest

I forlængelse af den deskriptive evaluering af resultaterne foretager vi en række permutationstest. En permutationstest er en nonparametrisk statistisk analysemetode, der bruges til at vurdere, om der er signifikant forskel mellem de observerede teststørrelser (fx gennem-

snit) for to grupper. Da vi ikke har grund til at antage, at vores data er normalfordelt, kan vi ikke anvende parametriske test som fx t-test. Permutationstesten er baseret på sampling og er derfor ikke knyttet til bestemte fordelinger af data. I vores analyser benyttes teststørrelserne *forskel i gennemsnit* og *forskel i standardafvigelse*¹ (Chihara & Hesterberg 2019: 52, 421-422).

Med permutationstesten sammenlignes observationer i to grupper ved at permutere observationerne gentagne gange. Det vil sige, at man samler observationerne fra de to grupper i én gruppe og redistribuerer alle observationerne i et nyt datasæt. Herefter beregnes forskellen mellem de to omorganiserede grupper. Denne proces gentages et stort antal gange (typisk 10^3). Afslutningsvis sammenlignes den faktiske observerede forskel mellem de to grupper med alle de permuterede forskelle, hvilket vil sige, at vi beregner en andel for, hvor mange forskelle der er lige så ekstreme eller mere ekstreme end den faktisk observerede forskel. Denne optælling giver en p-værdi, som vi accepterer ved et signifikansniveau på 0,05 (Chihara & Hesterberg 2019: 50-68)².

For hvert af de otte parametre tester vi således, om der er signifikant forskel mellem de to gruppers gennemsnitlige andele og standardafvigelser. Vi tester for lighed mellem grupperne. Det vil sige, at nulhypotesen i begge test er, at der ikke er forskel mellem de to grupper. Hvis resultatet er signifikant, dvs. p-værdien er under vores signifikansniveau på 0,05, betyder det, at nulhypotesen forkastes, eller sagt på en anden måde, at der sandsynligvis er en reel forskel mellem grupperne.

5 RESULTATER

I det følgende gennemgår vi resultaterne af analysen. Indledningsvis kaster vi et eksplorativt blik på histogrammerne med henblik på at få et overblik over datamaterialet. Herefter ser vi nærmere på resultaterne af

1 Denne teststørrelse baserer sig på en F-test, som antager normalfordeling. Da vi kun benytter teststørrelsen og ikke udfører selve F-testen, kan det legitimeres at benytte teststørrelsen i kombination med permutationstesten, selvom vi ikke har nogen antagelse om, at vores data har en bestemt fordeling.

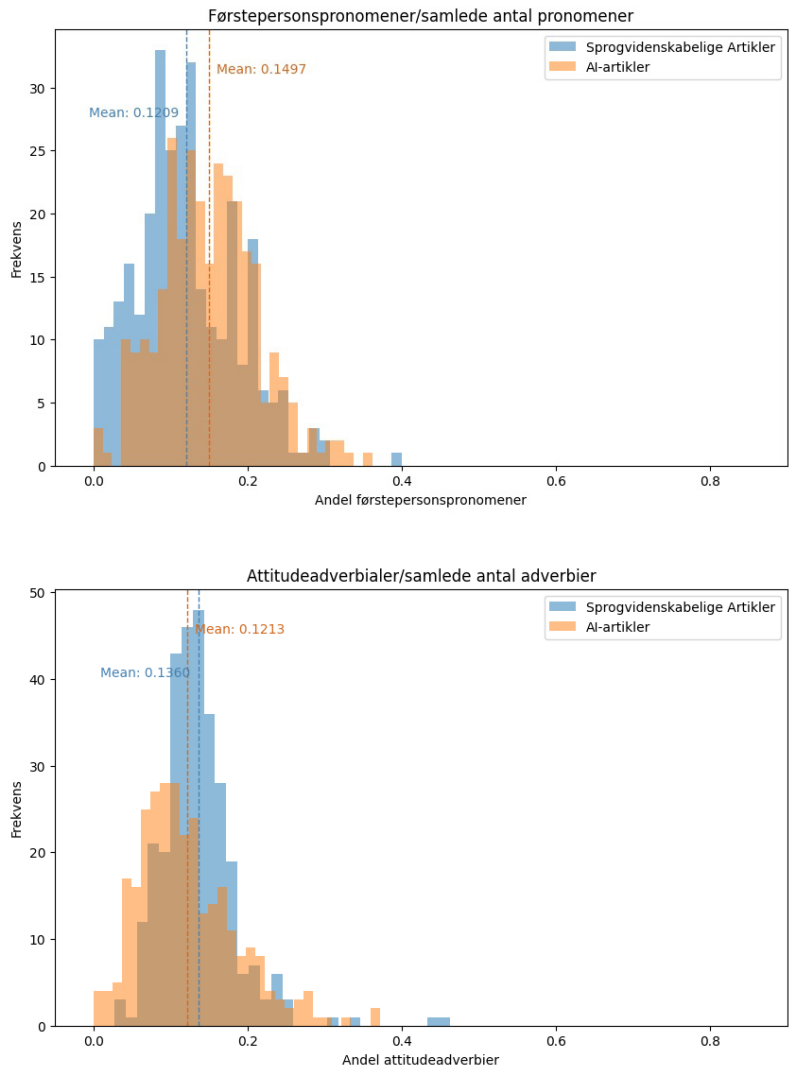
2 For en visuel forklaring af permutationstest, se www.jwilber.me/permutationstest/.

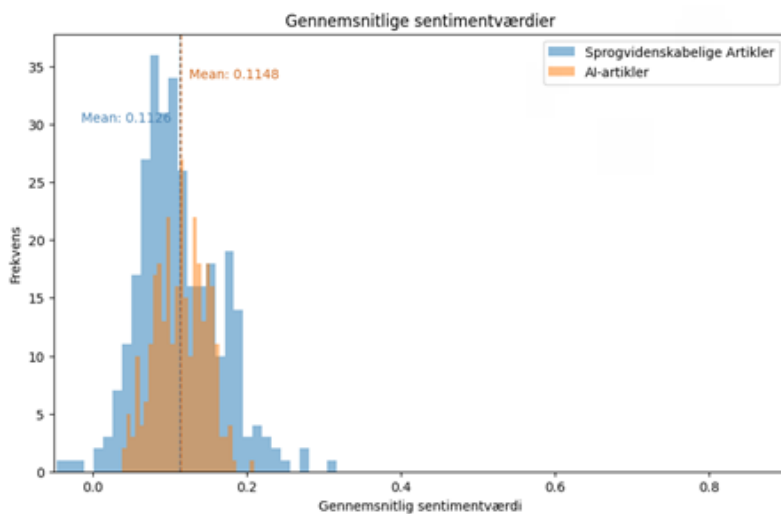
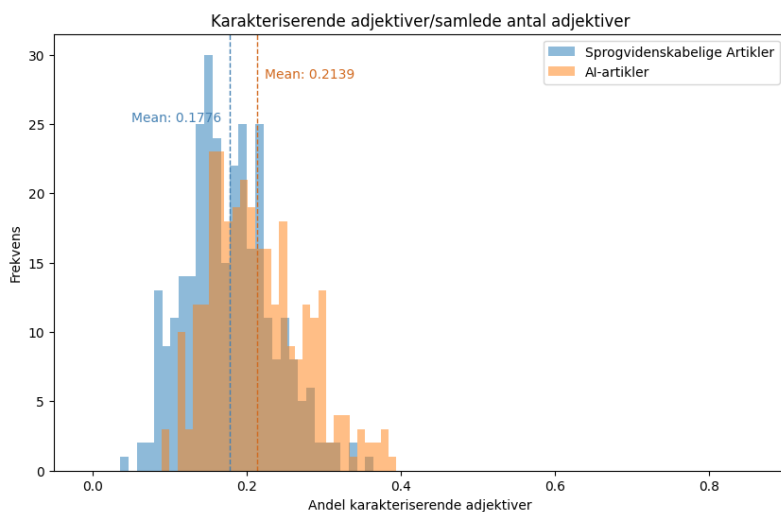
den deskriptive statistiske analyse samt p-værdierne for permutations-testene, og til sidst diskuterer vi resultaterne i relation til fænomenet fagdisciplinær stemme for hver af de tre kategorier *stillingtagen*, *engagement* og *fagspecifikt ordforråd*.

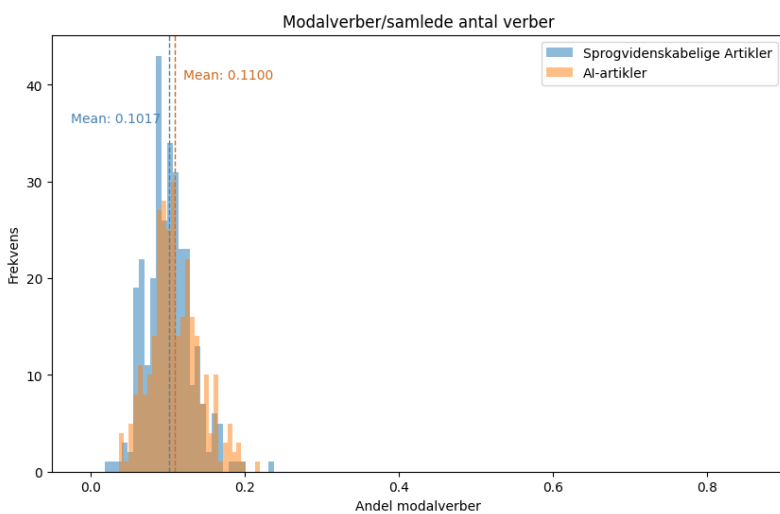
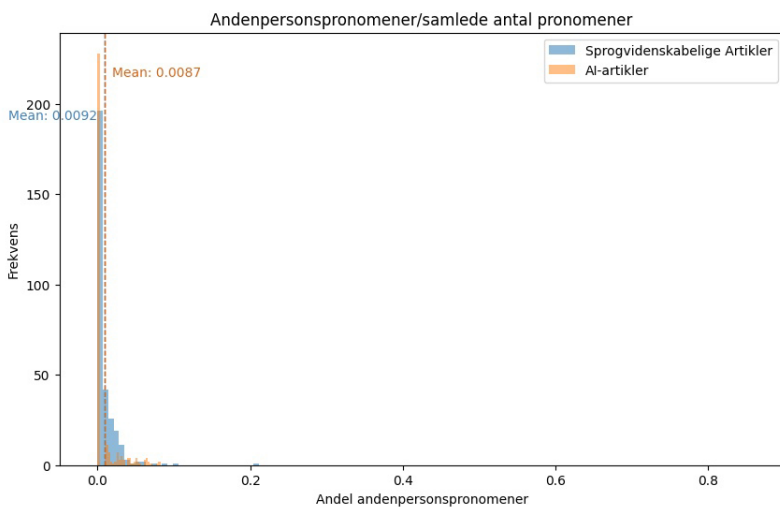
5.1 Histogrammer

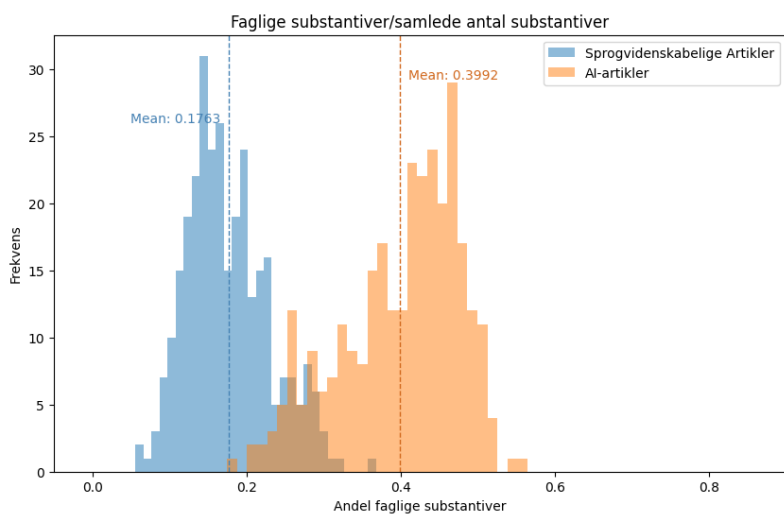
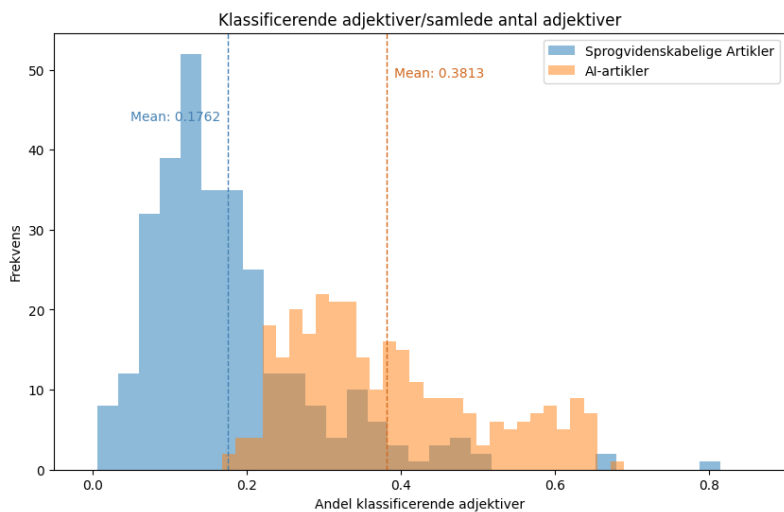
Histogrammer er et godt redskab til at give et visuelt overblik over datamaterialet og en første indikation af, hvad undersøgelsen viser, idet de viser fordelingen af målingerne inden for hver analysekategori. Man kan altså se, hvor mange tekster der har en given andel af de forskellige sproglige træk, og hvor stor afstanden er mellem de højeste og laveste andele i hver analysekategori. Hvert plot indeholder et histogram for gruppen af autentiske sprogvidenskabelige tekster (blåt) og et histogram for gruppen af AI-genererede tekster (orange). De stiplede linjer markerer de to gruppers gennemsnit, og den numeriske værdi er skrevet ud for hver linje, hvilket gør det muligt dels at aflæse, om de to gruppers gennemsnit ligger tæt på hinanden, dels at aflæse i hvilket omfang spredningerne for de to grupper minder om hinanden. X-aksen i histogrammerne viser intervaller med andele inden for den givne analysekategori, og for sentiment-histogrammet viser x-aksen den gennemsnitlige sentimentværdi. Y-aksen viser frekvensen og angiver, hvor mange tekster der falder inden for et givent interval. Hver søjle i histogrammet repræsenterer et interval på x-aksen, og søjlens højde afspejler, hvor mange tekster der falder inden for det givne interval. X-akserne er ensrettet i alle histogrammerne for at lette sammenligningen imellem kategorierne. Y-aksen er derimod justeret i forhold til den enkelte analysekategori for at lette læsningen af de enkelte analysekategorier; en ensretning her ville medføre meget små histogrammer i de fleste kategorier og dermed en mindre detaljeret visualisering.

FIGUR 1: HISTOGRAMMER FOR DE OTTE KATEGORIER









Når vi ser på histogrammerne, er der to ting, der springer i øjnene. For det første ligger gennemsnittene for de autentiske og de AI-genererede tekster i seks ud af de otte kategorier (førstepersonspronomen, attitudeadverbialer, karakteriserende adjektiver, sentimentværdi, andenpersonspronomen og modalverber) relativt tæt på hinanden, og for det andet er der i de samme seks kategorier en meget høj grad af overlap mellem histogrammerne, hvilket indikerer at spredningen er nogenlunde sammenlignelig for de to grupper af tekster i de seks kategorier. De eneste to kategorier, der afviger fra dette mønster, er klassificerende adjektiver og fagspecifikke substantiver. Her ser det til gengæld ud til at være relativt store forskelle både på de to gruppers gennemsnit og på spredningen. De seks første kategorier repræsenterer dimensionerne stillingtagen og engagement, mens de to sidste repræsenterer dimensionen fagspecifikt ordforråd.

Ved første øjekast ser det altså ud til, at GPT4 reproducerer dimensionerne stillingtagen og engagement på en måde, der i høj grad minder om det, vi ser i de autentiske tekster, mens der ser ud til at være en større forskel på brugen af fagspecifikt ordforråd. Her ser det faktisk ud til, at GPT4 i gennemsnit anvender flere fagspecifikke udtryk end det, vi ser i de autentiske tekster. Vi kan imidlertid ikke ud fra histogrammerne afgøre, om der er tale om signifikante forskelle. I det følgende ser vi derfor nærmere på den statistiske analyse af observationerne.

5.2 Statistiske resultater

Tabel 2 viser dels resultaterne af den deskriptive statistiske analyse af de to korpusser, dels resultatet af permutationstesten. Tabellen indeholder de observerede værdier for analysekategorierne førstepersonspronomen, attitudeadverbialer, karakteriserende adjektiver, sentimentværdi, andenpersonspronomen, modalverber, klassificerende adjektiver og fagspecifikke adjektiver, samt en kolonne for henholdsvis den observerede forskel i gennemsnit og den observerede forskel i standardafvigelse for hver kategori. Hver kategori har sin egen række, og de to korpusser er markeret med henholdsvis M (autentiske sprogvidenskabelige artikler) og A (AI-genererede tekster). Værdierne er for alle kategorier angivet i procent, bortset fra sentimentværdien i nederste række, der angiver

den gennemsnitlige sentimentværdi ud fra sentimentskalaen, der som tidligere beskrevet går fra -5 til 5.

TABEL 2: OBSERVEREDE VÆRDIER OG P-VÆRDIER FOR DE TO KORPUSSE

Observerator Kategori	Gruppe	Gen-nemsnit	Forskel gennem-snit	Stan-dardaf-vigelse	Forskel standard-afvigelse	Mini-mum	Mak-si-mum
Førstepersons-pronomener	M	12,09	2,88***	6,66	0,38	0	40,00
	A	14,97		6,28		0	36,14
Attitude-adverbialer	M	13,60	1,47**	5,01	1,42**	2,73	46,26
	A	12,13		6,43		0	37,14
Karakteriserende adjektiver	M	17,77	3,62***	5,59	0,55	3,56	36,50
	A	21,39		6,14		9,01	39,42
Sentimentværdi	M	0,11	0,0	0,052	0,02***	-0,05	0,35
	A	0,11		0,032		0,04	0,21
Andenpersons-pronomener	M	0,92	0,05	1,78	0,04	0	21,11
	A	0,87		1,82		0	8,24
Modalverber	M	10,17	0,83***	3,00	0,23	1,92	23,78
	A	11,00		3,23		3,76	22,02
Klassificerende adjektiver	M	17,62	20,51***	11,25	1,06	0,68	81,37
	A	38,13		12,31		16,83	68,93
Fagspecifikke substantiver	M	17,63	22,29***	5,38	2,35	5,60	36,77
	A	39,92		7,73		17,48	56,41

Signifikansniveau: Markering med * indikerer $p \leq 0.05$, ** indikerer $p \leq 0.01$, og *** indikerer $p \leq 0.001$.

Ser vi på gennemsnittene i tabellen ovenfor, viser det sig, at de AI-genererede tekster har en højere gennemsnitlig andel af førstepersonspronomen, karakteriserende adjektiver, modalverber, klassificerende adjektiver og fagspecifikke substantiver end de autentiske sprogvidenskabelige artikler. For attitudeadverbialerne og andenpersonspronomen forholder det sig omvendt, mens der ikke er nogen forskel på de to gruppers gennemsnitlige sentimentværdi. For andenpersonspronomenne og modalverberne er forskellen på gennemsnittene på under ét procentpoint, og for førstepersonspronomenne, attitudeadverbialerne og de karakteriserende adjektiver ligger gennemsnittene mellem 1,47 og 3,62 procentpoint. For både de klassificerende adjektiver og de fagspecifikke substantiver er afstandene mellem de to grupper på over 20 pro-

centpoint, henholdsvis 20,51 for de klassificerende adjektiver og 22,29 for de fagspecifikke substantiver.

Ser vi herefter på forskellen i spredning, dvs. forskellen mellem den observerede standardafvigelse for de to grupper, gælder det for kategorierne førstepersonpronomen og sentimentværdi, at standardafvigelsen for de AI-genererede tekster er lidt lavere end standardafvigelsen for de autentiske tekster. Man finder altså både autentiske tekster med færre og flere førstepersonspronomen, end man ser i gruppen af AI-genererede tekster. Samme billede gør sig gældende for sentimentværdien, hvor der i gruppen af autentiske tekster findes både mere negative, men også mere positive tekster, end man finder i gruppen af AI-genererede tekster. For de øvrige seks kategorier gælder det modsatte. Her er det de AI-genererede tekster, der spreder sig mere omkring gennemsnittet, hvilket vil sige, at man finder AI-genererede tekster med både flere og færre attitudeadverbialer, karakteriserende adjektiver, andenpersonspronomen, modalverber, klassificerende adjektiver og fagspecifikke adjektiver end i gruppen af autentiske tekster. For alle otte kategorier gælder dog, at der er tale om relativt små forskelle i standardafvigelse, der ligger mellem 0,02 (sentimentværdi) som det laveste og 2,35 som det højeste (fagspecifikke substantiver).

Selvom vi som ovenfor beskrevet har observeret forskelle mellem de to grupper af tekster, er det ikke muligt at vurdere, om forskellene er betydningsfulde alene ud fra sammenligninger af de observerede værdier. Vi har derfor ved brug af permutationstest beregnet p-værdier for de gennemsnitlige forskelle for gennemsnittet og for standardafvigelsen. Resultatet af permutationstestene kan også aflæses i tabel 2, hvor et signifikant resultat er markeret med en eller flere asterisker (* indikerer $p \leq 0.05$, ** indikerer $p \leq 0.01$, og *** indikerer $p \leq 0.001$). P-værdierne for forskel i gennemsnit angiver, om der signifikant forskel imellem de to tekstgruppers gennemsnitlige værdier, og p-værdierne for forskel i standardafvigelse angiver, om der er signifikant forskel i, hvor meget fordelingerne spreder sig omkring gennemsnittet for de to grupper.

Som det fremgår af tabel 2, er der signifikant forskel på gennemsnittene i alle kategorier bortset fra sentimentværdi og andenpersonspronomen, mens der kun er signifikant forskel mellem standardafvigelserne for attitudeadverbialer og sentimentværdi. I de tilfælde, hvor

der er signifikant forskel, er der med stor sandsynlighed en reel forskel på, hvordan de sproglige ressourcer anvendes i de AI-genererede tekster sammenlignet med, hvordan de anvendes i de autentiske sprogvidenskabelige tekster. De signifikante resultater indikerer således, at de målte forskelle mellem tekstgrupperne ikke skyldes tilfældige variationer i teksterne, men at de AI-genererede tekster generelt indeholder en større andel af førstepersonspronomen, karakteriserende adjektiver, modalverber, klassificerende adjektiver og fagspecifikke substantiver og en mindre andel af attitudeadverbialer end de autentiske tekster, samt at spredningen er lidt større i de AI-genererede tekster med hensyn til brugen af attitudeadverbialer og lidt større i de autentiske tekster med hensyn til sentimentværdi.

Signifikansresultaterne siger dog ikke i sig selv noget om, hvorvidt der er tale om bemærkelsesværdige eller vigtige forskelle, men peger blot på, hvor der reelt er tale om målbare forskelle. Hvis man som læser vurderede de to grupper af tekster kvalitativt, ville man måske bemærke forskellen på 20 procentpoint i brugen af klassificerende adjektiver og fagspecifikke substantiver, men ville man bemærke en forskel i brugen af førstepersonspronomen, hvor de autentiske sprogvidenskabelige artikler gennemsnitligt indeholder 12,09 % førstepersonspronomen ud af det samlede antal pronomen, mens de AI-genererede tekster gennemsnitligt indeholder 14,97 %? At forskellene er signifikante, indikerer således ikke nødvendigvis, at de er praktisk relevante, men i første omgang blot, at det er værd at se nærmere på den givne kategori for at undersøge, om en eventuel forskel faktisk er betydningsfuld. Dette gør vi i det følgende, hvor vi diskuterer resultaterne for de tre overordnede dimensioner engagement, stillingtagen og fagspecifikt ordforråd.

5.3 Fagdisciplinær stemme

5.3.1 Stillingtagen

Vi begynder med at se nærmere på stillingtagen, dvs. den holdningsmæssige dimension i teksterne relateret til teksternes afsenderorienterede træk, og spørgsmålet er nu, hvad de statistiske resultater siger om GPT4's evne til at reproducere denne dimension. Som vi så ovenfor, var der signifikante forskelle i gennemsnittet for kategorierne førstepersonspronomen, attitudeadverbialer og karakteriserende adjektiver,

men ikke for sentimentværdien. Der var desuden signifikante forskelle på standardafvigelsen i kategorierne attitudeadverbier og sentimentværdi, men ikke i kategorierne førstepersonspronomener og karakteriserende adjektiver. Mere præcist kan vi sige, at GPT4 på den ene side anvender en lidt højere andel af førstepersonspronomener og en lidt højere andel af karakteriserende adjektiver, sammenlignet med de autentiske tekster, mens den på den anden side anvender en lidt lavere andel attitudeadverbialer sammenlignet med de autentiske tekster. Der er dog samtidig lidt større variation i, hvordan GPT4 anvender attitudeadverbialerne, forstået på den måde, at der er AI-genererede tekster, der indeholder både flere og færre attitudeadverbialer sammenlignet med de autentiske tekster. For sentimentværdien er det derimod omvendt. Her er der autentiske tekster, der har både højere og lavere sentimentværdi sammenlignet med de AI-genererede.

Den statistiske analyse viser altså, at GPT4 både overforbruger og underforbruger de sproglige ressourcer knyttet til stillingtagen. Dette kunne indikere, at en læser vil møde en tekst, hvor afsenderinstansen på den ene side træder lidt tydeligere frem end i de autentiske tekster, dels gennem eksplicit markering af tilstedeværelse via førstepersonspronomener, fx 'jeg' og 'vi', dels gennem tydeligere markering af subjektive vurderinger via brugen karakteriserende adjektiver, fx 'stor', 'vigtig' og 'relevant', og på den anden side træder lidt mindre tydeligt frem i teksten, idet attitudeadverbialer som fx 'formodentlig', 'måske' og 'helt sikkert', der bl.a. bidrager til tydeliggøre afsenderens holdning til de omtalte, anvendes lidt mindre end det, vi så i de autentiske tekster. Vi så dog samtidig, at de observerede forskelle var små, på maksimalt fire procentpoint, hvilket vil sige, at det er tvivlsomt, om en læser i praksis vil registrere disse forskelle. Der vil naturligvis være tekster, der skiller sig ud med både mere og mindre eksplicit afsendertilstedeværelse, men forskellene er formentlig ikke store nok til, at en læser ved gennemlæsning af flere tekster vil bemærke en generel forskel på, hvordan stillingtagen realiseres i henholdsvis de AI-genererede tekster og de autentiske tekster. Vi hæfter os ved dette, da de små forskelle, til trods for at de i visse tilfælde er signifikante, indikerer, at GPT4 trods alt er tæt på at kunne reproducere stillingtagen på en måde, der set fra et kvantitativt perspektiv, minder om det, vi ser i de autentiske sprogvidenskabelige tekster.

5.3.2 *Engagement*

Vi ser herefter på engagement, dvs. den relationsmæssige dimension i teksterne, som er relateret til teksternes modtagerorienterede træk. De sproglige træk, der knytter sig til engagement, er andenpersonspronomen og modalverber, og som vi så ovenfor, var der kun signifikant forskel på den gennemsnitlige andel af modalverber i de to grupper af tekster. Der var ikke signifikant forskel på anvendelsen af andenpersonspronomen, og der var heller ikke signifikante forskelle på spredningen i de to kategorier. Det er altså kun i kategorien modalverber, at GPT4's reproduktion af engagement afviger fra det, vi ser i de autentiske sprogvidenskabelige artikler.

Den statistiske analyse viser her, at GPT4, som vi så med førstepersonspronomen og karakteriserende adjektiver, har et lille overforbrug af modalverber, fx 'kan', 'må' og 'vil'. Dette kunne indikere, at en læser vil møde en tekst, hvor afsenderinstansen, som en konsekvens heraf, fremstår lidt mere engageret i og forpligtet på interaktionen med modtageren sammenlignet med det, vi ser i de autentiske tekster. Som det også var tilfældet med brugen af kommunikative ressourcer knyttet til stillingtagen, er der imidlertid også her tale om en lille forskel, endda en forskel der er endnu mindre, end vi så før, nemlig en observeret forskel imellem de to tekstgrupper på kun 0,83 procentpoint. Resultatet er ganske vist signifikant, men der vil i praksis være tale om en forskel på ganske få ord per tekst, og vi anser det derfor som usandsynligt, at en læser vil bemærke det som en reel kvalitativ forskel på afsenderens engagement i henholdsvis de AI-genererede og de autentiske tekster. Overordnet set indikerer den statistiske analyse altså, idet der kun var signifikant forskel i en ud af de fire kategorier, og idet denne forskel desuden var meget lille, at GPT4 også er i stand til at reproducere engagement på en måde, der, set fra et kvantitativt perspektiv, ligger meget tæt på det, vi ser i de autentiske sprogvidenskabelige artikler.

5.3.3 *Fagspecifikt ordforråd*

Vi ser herefter på den tredje og sidste dimension, nemlig fagspecifikt ordforråd. De to sproglige træk, der knytter sig til denne dimension, er klasificerende adjektiver og fagspecifikke substantiver. Disse to kategorier er som tidligere nævnt ikke direkte relateret til konstruktet stemme,

men er alligevel interessante, fordi anvendelsen af fagspecifikt ordforråd giver en indikation af, i hvilket omfang afsenderen tilstræber en saglig og neutral tone eller en subjektiv og holdningsorienteret tone. Som vi så ovenfor, var der signifikante forskelle i gennemsnittene i begge kategorier, men ingen signifikante forskelle i spredningen. Vi så samtidig, at de signifikante gennemsnitlige forskelle for begge kategorier var relativt store, henholdsvis 20,51 for de klassificerende adjektiver og 22,29 for de fagspecifikke substantiver, og at det i begge tilfælde var de AI-genererede artikler, der havde de højeste andele. Den høje andel af fagspecifikke ord i de AI-genererede tekster kan enten indikere, at GPT4 overforbruger fagterminologi, eller at GPT4 bruger fagterminologi i samme grad som i autentiske tekster, men at der i de AI-genererede artikler i mindre grad findes tekst, der ikke indeholder fagterminologi, fx indledende afsnit eller metatekst, hvilket også vil resultere i en højere andel.

Den statistiske analyse viser således, at der er en reel og relativt stor forskel på, i hvilket omfang de to typer af sproglige ressourcer anvendes i henholdsvis de AI-generede og de autentiske tekster. Det er desuden værd at bemærke, at selv om der ikke er målbare forskelle på spredningen, når vi sammenligner de to grupper af tekster, så er fordelingerne ikke ens. Som de fremgik af plottene i figur 1, så mindede fordelingerne for de øvrige kategorier meget om hinanden. Det var altså ud over få og små målte forskelle også en høj grad af overlap mellem histogrammerne. Dette er ikke tilfældet, når vi se på histogrammerne for de klassificerende adjektiver og de fagspecifikke substantiver. For de klassificerende adjektiver kan vi se, at histogrammet for de autentiske tekster har en tydelig top med mange tekster lidt under gennemsnittet samt en hale mod højre med nogle få tekster med meget høje værdier, mens histogrammet for de AI-generede har en mere flad profil med mange tekster spredt ud over et større interval. For de fagspecifikke substantiver kan vi se, at begge histogrammer har en tydelig top samt en lille hale, men at de to histogrammer er spejlvendt, således at de autentiske tekster har en top, der ligger lidt under gennemsnittet og en hale, der peger mod højre, mens de AI-genererede har en top, der ligger lidt over gennemsnittet og en hale, der peger mod venstre. Dette er vigtigt, fordi det indikerer, at der er en variation i brugen af klassificerende adjektiver og fagspecifikke substantiver, som ganske vist ikke lader sig måle som signifikant forskel

i standardafvigelse, og som vi ikke umiddelbart kan forklare ud fra vores analyse, men som alligevel ser ud til at være interessant.

Set fra et læserperspektiv er forskellene så store, at det formentlig vil blive bemærket, både fordi det vil være tydeligt for en læser, at der i de AI-genererede tekster i højere grad end i de autentiske tekster anvendes fagspecifikt ordforråd, men også fordi den interne variation i grupperne formentlig er så stor, at det også vil opleves som forskellige skrivestile. Dette kunne fx komme til udtryk ved, at tekster genereret af GPT4, som følge af den noget højere andel af fagspecifikt ordforråd, fx 'syntaks', 'ortografisk' eller 'sociolingvistisk', genrelt vil fremstå mere formelle og mindre subjektive end de autentiske tekster. Samlet set viser den statistiske analyse af fagspecifikt ordforråd altså, modsat hvad vi så, da vi undersøgte de to andre dimensioner, at GPT4 reproducerer denne dimension på en måde, der ligger relativt langt fra den brug af fagspecifikt ordforråd, vi ser i de autentiske tekster.

6 DISKUSSION

Formålet med denne artikel har været at bidrage til at undersøge kvaliteten af AI-genereret akademisk prosa, og vi indledte med at spørge, i hvilket omfang generative AI-modeller, med GPT4 som eksempel, er i stand til at reproducere fagdisciplinær stemme i dansksproget sprogvidenskabelig prosa. Vi har undersøgt dette spørgsmål gennem en kvantitativ komparativ analyse af to korpusser bestående af henholdsvis en samling autentiske dansksprogede sprogvidenskabelige artikler og en samling AI-genererede tekster om sprogvidenskabelige emner. I analysen fokuserede vi på tre dimensioner, dels stillingtagen og engagement, der er direkte knyttet til konstruktet stemme, dels brugen af fagspecifikt ordforråd, der er et vigtigt aspekt af fagdisciplinær skrivning, men ikke specifikt er knyttet til stillingtagen og engagement. Resultatet af analysen var ikke entydigt. På den ene side viste der sig at være signifikante forskelle mellem de AI-genererede og de autentiske tekster i alle de tre overordnede dimensioner. På den anden side var der i flere tilfælde tale om endog meget små forskelle, der formentlig knapt ville blive bemærket af en læser. Groft sagt indikerede analyserne, at stillingtagen i form af en eksplicit tilstedeværende afsenderinstans er lidt tydeligere markeret i de AI-genererede tekster, at afsenderinstansen fremstår mi-

nimalt mere engageret i de AI-genererede tekster, og at de AI-genererede tekster i noget højere grad end de autentiske tekster er præget af fagspecifikt ordforråd.

Som tidligere beskrevet foreligger der kun ganske få studier af kvaliteten af AI-genererede tekster, og der foreligger os bekendt heller ikke benchmarkingresultater for tekstkvalitet af dansksproget akademisk prosa. Det kan derfor på den ene side være vanskeligt umiddelbart at tolke disse lidt tvetydige resultater. Generativ kunstig intelligens er på den anden side et område i voldsom vækst, bl.a. hjulpet på vej af støt stigende mængder af data, og der er derfor ingen særlig grund til at forvente, at GPT4 skulle reproducere fagdisciplinær stemme på en måde, der ligger meget langt fra det, vi ser i de autentiske tekster. Vi vælger derfor at tolke resultaterne som et udtryk for, at GPT4 for kategorierne stillingtagen og engagement, på trods af de målte forskelle, er på vej mod at kunne reproducere fagdisciplinær stemme på en måde, der set fra et kvantitativt perspektiv minder om den, vi ser i autentiske sprogvidenskabelige tekster. I forlængelse heraf er det derfor også overraskende, at den måde, hvorpå GPT4 reproducerede fagspecifikt ordforråd, afveg meget fra det, vi så i de autentiske sprogvidenskabelige tekster, endda på en sådan måde, at GPT4 brugte en væsentlig højere andel af fagspecifikke udtryk end det, vi ser i de autentiske sprogvidenskabelige tekster.

Vi kan ikke umiddelbart forklare denne forskel på baggrund af vores kvantitative undersøgelse. En mulig forklaring kunne være, at MUDS-korpusset ikke er repræsentativt for dansksproget sprogvidenskabelig prosa, og at forskellene ville udjævne sig, hvis korpusset havde været større og bredere. Der kan dog være mange andre årsager, og vores metode muliggør ikke uden videre en undersøgelse af disse årsagssammenhænge. Vores undersøgelse er også begrænset i den henseende, at der udelukkende er tale om en kvantitativ undersøgelse. Den siger således ikke noget om teksternes beskaffenhed, set ud fra et kvalitativt perspektiv, fx om indholdet er faktuel korrekt, eller om sproget er sammenhængende og meningsbærende.

Vores analyse giver derimod en indikation af, hvordan GPT4-modellen performer på nuværende tidspunkt set fra et kvantitativt perspektiv. Men dette er selvsagt kun et enkelt skridt på vejen mod et

egentligt overblik over kvaliteten af AI-genereret tekst. Både fordi vi kun har undersøgt et enkelt aspekt, nemlig fagdisciplinær stemme, men også fordi det i det hele taget er en stor udfordring at forstå, hvorfor et givet AI-genereret output ser ud, som det gør. Da generativ kunstig intelligens er baseret på probabilistiske modeller trænet på store mængder tekstdata, vil kvaliteten af AI-genereret tekst variere afhængigt af kvaliteten af træningsdatasættet for den generative model og vil i forlængelse heraf også afhænge af, om den genre og det faglige domæne, der genereres tekst inden for, er velrepræsenteret i træningsdatasættet. Kvaliteten af outputtet vil samtidig afhænge af kvaliteten af den statistiske model, der anvendes, og ikke mindst af måden, hvorpå man prompter modellen til at generere tekst.

Denne treenighed, data/model/prompt, gør det til en overordentlig kompleks opgave at vurdere kvaliteten af AI-genereret output, idet enhver ændring i bare én af de tre vil medføre ændringer i outputtet. Mængden af data samt indførelsen af transformerarkitekturen (Vaswani m.fl. 2017) har revolutioneret måden, vi tænker generativ kunstig intelligens på, og der er ingen grund til at tro, at udviklingen ikke vil fortsætte. Hertil kommer, at prompting ikke er en eksakt videnskab. Der er i bogstaveligste forstand uendeligt mange muligheder for at konstruere en prompt, samtidig med at det ikke på nogen måde er transparent, hvordan promptens design påvirker strukturen af outputtet. Dette er et område, der i sig selv kræver meget mere forskning.

Det er derfor ikke tilstrækkeligt at undersøge kvaliteten af autogenereret tekst ud fra hverken et enkelt fagspecifikt perspektiv eller et mere generisk synspunkt. Vi har tværtimod brug for mange flere undersøgelser, der skal undersøge tekstkvaliteten inden for mange forskellige felter, både inden for de akademiske discipliner og inden for andre områder. Det er desuden vigtigt, at der ikke kun er tale om datalogisk orienterede evalueringer af modellernes performance, men at man også ud fra en bred vifte af lingvistiske og kommunikative kriterier vurderer kvaliteten af de nye generative sprogværktøjer på både et mikro- og et makroniveau. Vi betragter, i forlængelse af dette, vores undersøgelse som et indledende skridt i retning af en systematisk og fortløbende indsats for at monitorere kvaliteten af de til enhver tid tilgængelige generative AI-redskaber.

7 AFSLUTTENDE BEMÆRKNINGER

De nye store sprogmodeller kommer som nævnt med løfter om at forandre vores skrivepraksisser, men hvis dette løfte skal indfries, forudsætter det, at kvaliteten af det autogenererede materiale har et rimeligt kvalitetsniveau. Vi har i denne artikel skitseret én mulig måde, hvorpå sprogvidenskabeligt baserede strategier for at teste kvaliteten af sprogværktøjer baseret på generativ kunstig intelligens kan udvikles. Dette er imidlertid ikke den eneste opgave, vi står overfor. Som vores undersøgelse har vist, har vi i høj grad også brug for mere empirisk viden om, hvordan autentiske tekster faktisk ser ud fra et kvantitativt perspektiv. Hvis vi ikke ved det, har vi reelt ikke noget sammenligningsgrundlag for vurderingen af de AI-genererede tekster. Endelig er der den udfordring, at AI-genererede tekster formentlig i mange tilfælde vil optræde uredigerede, fx i forskellige sammenhænge på de sociale medier og i nyhedsmedierne, men at det på sigt kun vil være en del af den AI-genererede tekst, vi vil støde på i vores kommunikative omgang med hinanden. Selvom ideen om at implementere AI-baserede skriveværktøjer i sin skrivepraksis for nogen kan virke fjern, så er realiteten nemlig den, at det allerede benyttes i vid udstrækning. Fx er AI-baserede værktøjer allerede implementeret i skriveprogrammer som Word og Google Docs. Vi vil altså med stor sandsynlighed se, at AI-modellerne vil blive en form for samarbejdspartner inden for det akademiske område, men også inden for mange andre områder. Det er således ikke nok at undersøge kvaliteten af AI-genererede tekster i sig selv. Fremover vil det i høj grad også være nødvendigt at undersøge kvaliteten af hybridtekster, dvs. tekster skrevet i et samarbejde mellem mennesker og modeller.

Der er naturligvis al mulig grund til at forholde sig både forsigtigt og kritisk til de nye omkalfatrende, generative teknologier. Der er dog samtidig grund til en vis optimisme i den forstand, at kriterierne for, hvad en god tekst er, ikke kan leveres af de generative teknologier selv. De må derimod komme fra de enkelte faglige domæner, herunder sprog- og tekstvidenskaberne. De generative teknologier har måske nok potentiale til at gøre skrivearbejdet nemmere i kvantitativ forstand, men næppe i kvalitativ forstand, hvis tekstens kompleksitet overstiger den, vi finder i indkøbssedler, bageopskrifter og andre skabelonagtige genrer. At bruge generative teknologier konstruktivt til at skabe kvalitetstekster er med

andre ord ikke noget, der kræver mindre viden om sprog og kommunikation sammenlignet med tidligere skrivepraksisser. Tværtimod kræver produktion af vellykkede hybridtekster indgående kendskab til normer og kvalitetskriterier for forskellige genrer og fremstillingsformer og kalder derfor i høj grad på, at sprog- og tekstvidenskaberne engagerer sig kritisk konstruktivt i arbejdet med at udvikle og kvalitetssikre de nye generative teknologier.

Ea Lindhardt Overgaard, ph.d.-studerende
Institut for Kommunikation og Kultur, Nordisk Sprog og Litteratur
elt@cc.au.dk

Ulf Dalvad Berthelsen, lektor
Institut for Kommunikation og Kultur, Nordisk Sprog og Litteratur
udb@cc.au.dk

LITTERATUR

- Anderson, N. m.fl. 2023. AI did not write this manuscript, or did it? Can we trick the AI text detector into generated texts? The potential future of ChatGPT and AI in Sports & Exercise Medicine manuscript generation. *BMJ Open Sport — Exercise Medicine* 9(1). 1–4. DOI: 10.1136/bmjsem-2023-001568.
- Baumann, J.F. & M.F. Graves. 2010. What Is Academic Vocabulary? *Journal of Adolescent & Adult Literacy*. Hoboken: Blackwell Publishing Ltd. 54(1). 4–12. DOI: 10.1598/JAAL.54.1.1.
- Bender, E.M. m.fl. 2021. On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (FAccT '21), 610–623. DOI: 10.1145/3442188.3445922.
- Berthelsen, U.D. 2007. Attitudeadverbialerne i semantisk og pragmatisk perspektiv. *MUDS - Moderne om Udforskningen af Dansk Sprog* 11. 52–63.
- Biber, D. 2006. *University language: a corpus-based study of spoken and written registers*. Philadelphia: Benjamins.
- Biber, D. & S. Conrad. 2009. *Register, genre, and style*. Cambridge: University Press.
- Biber, D. & B. Gray. 2015. *Grammatical Complexity in Academic English: Linguistic Change in Writing*. Cambridge: Cambridge University Press. DOI: 10.1017/CBO9780511920776.

- Blei, D.M., A.Y. Ng & M.I. Jordan. 2003. Latent Dirichlet Allocation. *Journal of Machine Learning Research* 3. 993–1022.
- Carter, M. 2007. Ways of Knowing, Doing, and Writing in the Disciplines. *College composition and communication*. Urbana: National Council of Teachers of English 58(3). 385–418.
- Chatman, S. 1978. *Story and discourse: narrative structure in fiction and film*. Ithaca: Cornell University Press.
- Chihara, L.M. & T.C. Hesterberg. 2019. *Mathematical Statistics with Resampling and R*. John Wiley & Sons.
- Christensen, L.D. & R.Z. Christensen. 2019. *Dansk Grammatik*. Odense: Syddansk Universitetsforlag.
- Clusmann, J. m.fl. 2023. The future landscape of large language models in medicine. *Communications Medicine*. Nature Publishing Group 3(1). 1–8. DOI: 10.1038/s43856-023-00370-1.
- Dergaa, I. m.fl. 2023. From human writing to artificial intelligence generated text: examining the prospects and potential threats of ChatGPT in academic writing. *Biology of Sport* 40(2). 615–622. DOI: 10.5114/biolSport.2023.125623.
- Eke, D.O. 2023. ChatGPT and the rise of generative AI: Threat to academic integrity? *Journal of Responsible Technology* 13. DOI: 10.1016/j.jrt.2023.100060.
- Enevoldsen, K., L. Hansen & K. Nielbo. 2021. DaCy: A unified framework for Danish NLP. *Ceur Workshop Proceedings*, vol. 2989, 206–216.
- Fløttum, K., T. Dahl & T. Kinn. 2006. *Academic Voices: Across languages and disciplines*. Amsterdam: John Benjamins Publishing Company.
- Frye, B.L. 2022. Should Using an AI Text Generator to Produce Academic Writing Be Plagiarism? *Fordham Intellectual Property, Media & Entertainment Law Journal*.
- Ghosh, S. & A. Caliskan. 2023. *ChatGPT Perpetuates Gender Bias in Machine Translation and Ignores Non-Gendered Pronouns: Findings across Bengali and Five other Low-Resource Languages*. arXiv. DOI: 10.48550/arXiv.2305.10510.
- Guinda, C.S. 2012. Proximal Positioning in Students' Graph Commentaries. K. Hyland & C. Sancho Guinda (red.), *Stance and Voice in Written Academic Genres*, 166–183. London: Palgrave MacMillan UK.
- Grønnum, N. 1998. *Fonetik og fonologi: almen og dansk*. København: Akademisk Forlag.
- Gunser, V.E. m.fl. 2022. The Pure Poet: How Good is the Subjective Credibility and Stylistic Quality of Literary Short Texts Written with an Artificial Intelligence Tool as Compared to Texts Written by Human Authors? *Proceedings of the Annual Meeting of the Cognitive Science Society* 44. 1744–1750.

- Halliday, M.A.K. 1985. *An introduction to functional grammar*. London: Edward Arnold.
- Halliday, M.A.K. & R. Hasan. 1976. *Cohesion in English*. London: Longman.
- Hansen, E. & L. Heltoft. 2019. *Grammatik over det danske sprog*. København: Det Danske Sprog- og Litteraturselskab.
- Harder, P. 2006. Dansk Funktionel Lingvistik. *NyS – Nydanske Sprogstudier* 34(34–35). 92–130. DOI: 10.7146/nys.v34i34-35.13459.
- Hinton, M. & J.H.M. Wagemans. 2023. How persuasive is AI-generated argumentation? An analysis of the quality of an argumentative text produced by the GPT-3 AI text generator. *Argument & Computation*. 14(1). 59–74. DOI: 10.3233/AAC-210026.
- Hitsuwari, J. m.fl. 2023. Does human–AI collaboration lead to more creative art? Aesthetic evaluation of human-made and AI-generated haiku poetry. *Computers in Human Behavior* 139. 107502. DOI: 10.1016/j.chb.2022.107502.
- Hyland, K. 2000. *Disciplinary discourses: social interactions in academic writing* (Applied Linguistics and Language Study). Harlow: Longman, an imprint of Pearson Education.
- Hyland, K. 2005. Stance and engagement: a model of interaction in academic discourse. *Discourse Studies*. 7(2). 173–192.
- Hyland, K. 2008. Disciplinary voices: Interactions in research writing. *English Text Construction*. 1(1). 5–22. DOI: 10.1075/etc.1.1.03hyl.
- Hyland, K. 2017. *The essential Hyland: studies in applied linguistics*. London: Bloomsbury Academic. DOI: 10.5040/9781350037939.
- Jensen, C.Z. & E. Sørhaug. 2021. *The Perfect Rap Lyrics - AI Generated Rap Lyrics That Are Better Than Lyrics from Existing Popular and Critically Acclaimed Rap Songs*. NTNU Master thesis. <https://ntnuopen.ntnu.no/ntnu-xmlui/handle/11250/2834977>.
- Krogh, E. 2012. Literacy og stemme – et spændingsfelt i modersmålsfaglig skrivning. S. Ongstad (red.), *Nordisk morsmålsdidaktikk: forskning, felt og fag*. Oslo: Novus.
- Köbis, N. & L.D. Mossink. 2021. Artificial intelligence versus Maya Angelou: Experimental evidence that people cannot differentiate AI-generated from human-written poetry. *Computers in Human Behavior* 114. 1–12. DOI: 10.1016/j.chb.2020.106553.
- Lauridsen, G.A., J.A. Dalsgaard & L.K.B. Svendsen. 2019. SENTIDA: A New Tool for Sentiment Analysis in Danish. *Journal of Language Works - Sprogvidenskabeligt Studentertidsskrift* 4(1). 38–53.

- MacBeath, J. 2006. Finding a voice, finding self. *Educational Review*. Routledge 58(2). 195–207. DOI: 10.1080/00131910600584140.
- McKinney, W. 2010. Data Structures for Statistical Computing in Python. *Proceedings of the 9th Python in Science Conference* 56–61. DOI: 10.25080/Majora-92bf1922-00a.
- Nordahl-Hansen, A. & T. Kvernbekk. 2020. Construct Validity in Scientific Representation: A Philosophical Tour. *Nordisk tidsskrift for pedagogikk & kritikk* 6. 88–99. DOI: 10.23865/ntpk.v6.1704.
- Ondelli, S. 2018. Treat Texts as Data but Remember They Are Made of Words: Compiling and Pre-processing Corpora. A. Tuzzi (red.), *Tracing the Life Cycle of Ideas in the Humanities and Social Sciences*, 133–150. Springer International Publishing. DOI: 10.1007/978-3-319-97064-6_7.
- OpenAI. 2023. ChatGPT. <https://openai.com/chatgpt>. (8 januar, 2024).
- Peng, Y. m.fl. 2023. AI-generated text may have a role in evidence-based medicine. *Nature Medicine*. Nature Publishing Group 29(7). 1593–1594. DOI: 10.1038/s41591-023-02366-9.
- Rimmon-Kenan, S. 1983. *Narrative fiction: contemporary poetics*. Reprint. London: Methuen.
- Togeby, O. 2003. *Fungerer denne sætning?: funktionel dansk sproglære*. København: Gad.
- Togeby, O. 2014. *Bland blot genrerne - ikke teksterne!: om sprog, tekster og samfund*. Frederiksberg: Samfundslitteratur.
- Van Rossum, G. & F.L. Drake. 2009. *Python 3 Reference Manual*. Scotts Valley, CA: CreateSpace.
- Vaswani, Ashish, m.fl. 2017. Attention is all you need. *Advances in neural information processing systems* 30. https://proceedings.neurips.cc/paper_files/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html
- Wall, B. 1991. *The narrator's voice: the dilemma of children's fiction*. London: MacMillan Press Limited. DOI: 10.1007/978-1-349-21109-8.
- Weston, S.J. m.fl. 2023. Selecting the Number and Labels of Topics in Topic Modeling: A Tutorial. *Advances in Methods and Practices in Psychological Science*, 6(2). 1–13. DOI: 10.1177/25152459231160105.

På sporet af chatbottens sproglige fingeraftryk

En sproglig ophavsanalyse af tekster skrevet af danskstuderende og ChatGPT

JONAS NYGAARD BLOM, ALEXANDRA HOLSTING & JESPER TINGGAARD SVENDSEN

ABSTRACT

Chatbotterne har for alvor meldt deres ankomst i uddannelsessystemet og skaber lige nu muligheder men også udfordringer. Et særligt presserende problem er, at der ikke findes nogen fuldt pålidelig metode til at spore, om en tekst er skrevet af en chatbot. Det betyder på den ene side, at studerende kan bruge chatbotter til at skrive opgaver, uden at det kan opdages, og på den anden side, at studerende kan risikere at blive mistænkt for at have brugt chatbotter, selvom det ikke er tilfældet. I dette studie bidrager vi til den aktuelle forskning i chatbotters sproglige særtræk ved at undersøge, om det er muligt at skelne sprogligt mellem tekster skrevet af danskstuderende og ChatGPT. Resultaterne viser, at der generelt set er visse forskelle på de studerendes og chatbottens skrivestil, men samtidig, at det ved hjælp af supplerende prompts er muligt at sløre chatbottens sproglige fingeraftryk. Dette understreger vanskeligheden ved at spore tekster skrevet af kunstig intelligens. Afslutningsvis diskuterer vi betydningen af dette for de danske uddannelsesinstitutioner.

EMNEORD: chatbotter; ChatGPT; kunstig intelligens; sproglig ophavsanalyse

1. INDLEDNING

Med tilkomsten af chatbotter drevet af generativ kunstig intelligens står samfundet midt i en af de største teknologiske omvæltninger i nyere tid. Udviklingen påvirker professioner og uddannelser, hvor tekstproduktion udgør en central opgave, heriblandt danskfaget, hvilket fører til nye faglige muligheder men også udfordringer. Chatbotterne behersker allerede i deres nuværende versioner flere af de praktiske skrivefærdigheder og analytiske kompetencer, som mange studerende normalt skal

tillære sig under uddannelse og senere sikre sig job på baggrund af. Det medfører, at tekstproducerende fag som dansk står i en vanskelig situation. Her er det vigtigt at finde en balance mellem at anvende den nye teknologi, så de studerende får en tidssvarende uddannelse og bliver konkurrencedygtige på fremtidens arbejdsmarked, og samtidig at bevare de sproglige og analytiske kompetencemål, som er fagets kerne og eksistensgrundlag, ikke mindst fordi disse kompetencer gør det muligt for de studerende at forholde sig selektivt og kritisk til de tekster, som chatbotterne producerer.

Brugen af chatbotter risikerer dog hurtigt at spænde ben for fagets traditionelle læringsmål, især i tilfælde, hvor de studerende skal arbejde selvstændigt derhjemme med skriftlige opgaver eller eksamineres skriftligt. Her kan det være fristende for nogle studerende at udnytte teknologien til at gøre arbejdet for dem, også selvom det i mange tilfælde slet ikke er tilladt, fx ved folkeskolens afgangsprøver (Børne- og Undervisningsministeriet 2023a) og til stedprøver på flere af landets universiteter, medmindre andet står eksplicit nævnt i fagbeskrivelsen (Aarhus Universitet 2023, Uniavisen 2023, Syddansk Universitet 2023). En undersøgelse foretaget for Akademikerbladet i 2023 viste, at hver tiende universitetsstuderende har brugt chatrobotter til at snyde (Akademikerbladet 2023). Dette behøver i sagens natur ikke betyde, at hver tiende *dansk*studerende har gjort det samme, men det er på den anden side heller ikke utænkeligt, at det forholder sig sådan. I den aktuelle situation er der samtidig en risiko for, at undervisere og eksaminatorer begynder at mistænke studerende for at bruge chatbotter, også selvom det ikke nødvendigvis er tilfældet. I uddannelsessystemet har udviklingen skabt et akut behov for viden om forskelle på og ligheder mellem studerendes og chatbotters skrivestil, som kan bidrage til at afgøre, om opgaver er skrevet af eller med hjælp fra en chatbot.

Vores studie tager afsæt i denne aktuelle problemstilling og har som formål at bidrage til udforskningen af, hvad der karakteriserer den AI-genererede skrivestil. Vi undersøger forskellige aspekter af chatbotens dansksproglige skrivestil og afdækker, om det ved hjælp af sproglig ophavsanalyse er muligt at skelne mellem chatbotters og studerendes tekster. Nærmere bestemt laver vi en sammenlignende analyse af to tekstkorpusser; det ene korpus består af opgaveløsninger skrevet af ny-

startede danskstuderende på SDU, og det andet korpus af opgaveløsninger genereret af ChatGPT v. 4.0 (OpenAI 2023) ud fra et prompt med samme opgaveformulering. Vi sammenligner altså maskinens tekster med menneskets for at se, om chatbotten har en særlig sprogbrug, der adskiller sig fra de studerendes.

I de følgende afsnit forklarer vi først, hvordan chatbotterne fungerer, og vi gennemgår de redskaber og metoder, som man i andre sammenhænge har forsøgt at udvikle til at spore tekster skrevet af chatbotter, og opsummerer, hvad forskningen har vist om deres anvendelighed og pålidelighed. Dernæst redegør vi for korpusserne og vores teoretiske og metodiske tilgang, hvorefter vi præsenterer vores analyser og resultater. Afslutningsvis diskuterer vi vores fund, redegør for mulige fejlkilder og konkluderer.

2. SPROGMODELLER OG SPORINGSREDSKABER

Generative chatbotter som ChatGPT drives af store sprogmodeller (Large Language Models), der er trænet med omfattende mængder af tekster og som ved hjælp af kunstig intelligens og neurale netværk er i stand til at afkode sproglige mønstre og selv generere tekster som svar på *prompts*, dvs. instruktioner fra brugerne (Lingard 2023). Sprogmodeller er en særlig type statistiske modeller, der forudsiger sandsynligheden for, hvad det næste ord er i en given sætningsstreng og tekstuel sammenhæng (Wu et al. 2023). Selvom chatbotterne er bygget op om en sådan maskinel statistisk logik, evner de at mime menneskelig kommunikation så overbevisende, at det i mange tilfælde kan være svært at se forskel på maskinens og menneskets sprog. Ved at lade chatbotter gennemføre en række kognitive sprogtest, som man normalt bruger til at forstå menneskers processering af sprog, har man ligefrem påvist en række ligheder mellem menneskers og chatbotters sprogforståelse, selv hvad angår helt grundlæggende syntaktisk og semantisk processering (Cai et al. 2023). Blandt andet derfor er det ikke nogen enkelt sag at spore, hvornår en tekst er genereret af en chatbot.

Der findes os bekendt endnu ingen applikationer på markedet, der med sikkerhed kan afgøre, om en tekst er skrevet af en chatbot eller et menneske, også selvom talrige aktører i techindustrien har arbejdet på at udvikle sporingsapplikationer, fx OpenAI, GPTZero, Copyleaks, Originality, Winston AI, GLTR, Crossplag og Content at Scale. Forskningen inden

for området har vist, at applikationerne har svært ved at opdage, om en chatbot står bag teksten, og at de har en tendens til at give falsk positive svar og derfor risikerer at opfatte menneskelige tekster som kunstige (Chaka 2023, Elkhatat 2023, Khalil & Er 2023, Otterbacher 2023, Perkins et al. 2023). Applikationerne har ligefrem en tendens til systematisk at vurdere menneskelige tekster som kunstige, når forfatteren har skrevet på et andet sprog end sit modersmål (Liang et al. 2023). Konsekvensen kan være, at man kommer til at anklage sagesløse studerende for at snyde, hvad der allerede har været eksempler på i medierne (Rolling Stone 2023). Samtidig er udviklingen af sporingsredskaber til mindre sprog som dansk langt fra så veludviklet som inden for hovedsproget engelsk.

Techindustriens tilgang til sporing er ikke altid lige transparent, blandt andet fordi firmaerne – ligesom producenterne af chatbotter – har det med at holde på deres forretningshemmeligheder, og fordi firmaerne hurtigt risikerer, at folk kan prompte sig uden om sporingsapplikationers alarmsystem, hvis det er alment kendt, hvordan de fungerer. Det står dog klart, at applikationerne typisk bygger på dybdelærning af teksttrænede algoritmer, der blandt andet beregner graden af sproglig kompleksitet og forudsigelighed ('perplexity') og sproglig variation ('burstiness') (Chaka 2023).

Blandt forskere finder man også forsøg på at spore chatbotternes sproglige fingeraftryk. Særlig lovende og relevant for vores studie er Desaire et al. (2023), der har udviklet en målemetode, som med en anslået træfsikkerhed på 99 % kan skelne en genrevariant af engelsksprogede videnskabelige tekster, der er skrevet af mennesker, fra tilsvarende tekster skrevet af ChatGPT.

Forfatterne anvender i deres analyser fire mål til at afgøre, om de videnskabelige tekster er skrevet af ChatGPT eller mennesker: 1) afsnitslængde, 2) periodelængde¹, 3) særlig tegnsætning og 4) højfrekvente ord. Resultaterne fra undersøgelsen viser blandt andet, at de menneske-

1 I forlængelse af Desaire et al. (2023) forstår vi periodelængde som det antal ord, der står mellem to punktummer eller tegn med punktums vægt. Vi opfatter samtidig afsnit som typografiske helheder i teksten, der adskilles af linjespring og/eller linjeskift og evt. indryk. Dog betragter vi eksempler, der er fremhævet vha. linjespring, og bulletpoints adskilt af linjespring som en del af det typografiske afsnit. Vi har udeladt overskrifter i optællingerne. Det samme gælder citater fra den teoritext, som opgaven tager afsæt i, da indlejrede sætninger fra andre tekster ville kunne forurene resultaterne.

lige tekster generelt indeholder flere meget lange perioder og flere meget korte perioder og samtidig har større standardafvigelse i afsnits- og periodelængde end chatbottens tekster. Dette er også bekræftet af andre studier (AlAfnan & MohdZuki 2023, Muñoz-Ortiz et al. 2023). De menneskelige tekster indeholder desuden flere adversative og kausale koblingsord som *although*, *however*, *but* og *because*.

Noget tilsvarende observeres af Herbold et al. (2023), der i et større sammenlignende studie af essays skrevet af gymnasieelever og ChatGPT påviser, at chatbotten generelt bruger færre diskursmarkører. I studiet viser der sig også andre forskelle på den menneskelige og maskinelle skrivestil. Blandt andet anvender chatbotten færre epistemiske markører, men til gengæld flere nominaliseringer, og den har samtidig en tendens til at strukturere sine tekster på en ensartet måde med en tilsyneladende fast argumentationsstruktur.

Altså synes der – i hvert fald i de undersøgte engelsksprogede gener – at være målbare træk, der kan være med til at sandsynliggøre, om teksterne er skrevet af en chatbot eller et menneske.

3. TEORETISK OG METODISK TILGANG

For at kunne undersøge chatbottens skrivestil og identificere eventuelle sproglige forskelle på chatbottens og de studerendes tekster har vi valgt at anvende en kombineret kvantitativ og kvalitativ tilgang til sproglig ophavsanalyse som forskningsmetode. Sproglig ophavsanalyse anvendes blandt andet i retslingvistikken til at vurdere, hvem der (sandsynligvis) har skrevet en given tekst. Man finder frem til dette ved at sammenligne to korpusser, typisk af kendt og ukendt ophav, og lede efter sproglige normafvigelser og stilistiske særtræk. Man arbejder her ud fra den antagelse, at de fleste sprogbrugere har en unik måde at udtrykke sig på, en særlig idiolekt kendetegnet ved bestemte sociolingvistiske træk (Christensen 2017), og det er den, man forsøger at spore.

I vores tilgang til ophavsanalysen anvender vi indledningsvis de samme kvantitative måleparametre² som Desaire et al. (2023), da denne

2 Dog udelader vi kategorien *særlig tegnsætning*. Det gør vi blandt andet, fordi apostroftegnet i engelsk anvendes mere hyppigt i genitivfunktion end i dansk, hvor apostroffen – normativt set – kun markerer genitiv i ord, der ender på *-s*, *-z*, eller *x* og ved forkortelser uden forkortespunktum. Resultaterne vil derfor ikke være fuldt sammenlignelige.

undersøgelse som nævnt har vist lovende resultater. Fremgangsmåden ligger her inden for rammerne af stilometrisk ophavsanalyse (Grieve 2023), der har til formål at identificere højfrekvente ord og særlige stiltræk, der kendetegner en given skribents skrivemåde. Vi har dog valgt en mere omfattende stilometrisk tilgang end Desaire et al. (2023) ved at lave en korpuslingvistisk analyse af både nøgleord og kollokationer, dvs. forbindelser af ord, der optræder hyppigt sammen. Vi forsøger på den måde at spore særlige frasemønstre, der potentielt kan være en del af chatbottens sproglige fingeraftryk.

Længdeanalyserne er lavet vha. en automatiseret optælling, mens analyserne af ordfrekvenser og kollokationer er lavet i programmet Ant-Conc (Anthony 2023). I analyseafsnittet forklarer vi mere om, hvordan vi har gjort dette i praksis.

Som en del af vores kvantitative analyser har vi desuden valgt at optælle antallet af sprogfejl, da vi fra forudgående studier (Blom et al. 2017) ved, at danskstuderende ofte afviger fra retskrivningsnormen, mens det – os bekendt – endnu ikke er undersøgt, i hvor høj grad ChatGPT efterlever standardnormerne for dansk skriftsprog. Forskelle i sprogriktighed kan derfor være en potentiel markør. Vi har kodet sprogfejlene manuelt i korpusserne og optalt dem derefter.

For at styrke sporingen yderligere supplerer vi de kvantitative analyser med en mere dybdegående kvalitativ analyse. En sådan kombineret metodisk tilgang er nødvendig med tanke på, at man lige nu oplever en form for “våbenkapløb” (Otterbacher 2023) mellem dem, der arbejder på at udvikle sporingsmekanismer, og dem, der arbejder på at udvikle prompts, der kan imødegå dem. Sporing, der er baseret på relativt enkle kvantitative kriterier som længde og frekvens, vil øjensynlig kunne imødegås med tilsvarende simple prompts – hvilket vi selv tester afslutningsvis – og der kan derfor være brug for, at man supplerer sporingen med andre sproglige markører.

I den kvalitative analyse har vi valgt at undersøge, hvordan chatbotten og de studerende anvender *interaktionelle og interaktive ressourcer* (Hyland & Tse 2004) i teksterne, dvs. hvordan de som afsendere positionerer sig, og hvordan de engagerer modtagerne og guider dem gennem teksten. Disse ressourcer er med til at opbygge en både personlig og interpersonel diskurs og hører til et særligt lag i teksten kaldet

metadiskursen (Hyland 2005), der rækker ud over tekstens propositionelle udsagn. I og med chatbotten – til dels afhængigt af, hvordan den er promptet – ofte vil mangle det (fulde) situationelle overblik, er det interessant at undersøge, om den anvender de interaktionelle og interaktive ressourcer anderledes end de studerende, der typisk har et mere fyldestgørende overblik.

De interaktionelle ressourcer giver afsenderen mulighed for at referere til sig selv (fx *jeg har valgt*), bruge styrkemarkører (fx *måske* som 'hedge' eller *tydeligvis* som 'booster') og attitudemarkører (fx *desværre*) til at signalere sin indstilling til det propositionelle indhold (Hyland & Tse 2004) og engagere læseren (fx *som vi kan se*).

De interaktive ressourcer muliggør fx helsætningskoblinger (fx *men*), diskursmarkører (fx *afslutningsvis*), henvisninger til informationer i teksten og ydre informationskilder (fx *nedenfor* og *ifølge X*) og forklaringsmarkører (fx *med andre ord*).

I vores analyser har vi ledt efter sådanne interaktionelle markører i korpusserne, annoteret dem og næranalyseret, hvilken funktion de har i teksten. Hvad angår de interaktive ressourcer, har vi valgt at fokusere på sætningskoblinger og diskursmarkører, da tidligere studier (Desaire et al. 2023 og Herbold et al. 2023) som nævnt her har fundet kvantitative forskelle på ChatGPT og studerendes skrivestil.

Samlet set undersøger vi følgende træk i teksterne i vores sproglige ophavsanalyse:

- A. afsnits- og periodelængde
- B. højfrekvente ord og kollokationsmønstre
- C. afvigelser fra retskrivningsnormen
- D. interaktionelle ressourcer: selvreferencer og styrke-, attitude- og engagementsmarkører
- E. interaktive ressourcer: koblingsord og diskursmarkører

4. TEKSTKORPUSSENE

Som nævnt bygger vores undersøgelse på to korpusser. Det første fungerer som *referencekorpus* og består af 29 tekster skrevet af nystartede danskstuderende på Syddansk Universitet i 2015. Korpusset har tidligere været anvendt i andre undersøgelser udarbejdet af denne artikels

forfattere og kolleger (Blom et al. 2017), hvilket giver os den fordel, at vi på forhånd har et generelt overblik over en række sproglige og tekstuelle træk i opgaverne. Korpuset fylder i alt 23.904 ord. I opgaven skulle de studerende først læse en teoretisk artikel om tekstlingvistiske begreber som kohæsion, kohærens, proformer og anaforisk reference. Derefter skulle de redegøre for teorien og anvende den til at analysere og til sidst vurdere en tekst præget af forskellige typer af kohæsi- og kohærensproblemer. Opgaven er på den måde designet til at stige i læringstaksonomisk sværhedsgrad fra redegørelse til analyse og til sidst kritisk vurdering.

Det andet korpus fungerer som *målkorpus*, dvs. det korpus, som vi måler på via sammenligning med referencekorpuset, og som vi forsøger at finde særlige ophavstræk ved. Korpuset rummer ti opgaver med i alt 2.845 ord skrevet af ChatGPT v. 4.0. Vi har promptet chatbotten med den samme opgaveformulering, som de studerende har løst; det gjorde vi den 21. april 2023 fra en og samme IP-adresse. I og med ChatGPT besvarer prompts forskelligt på grund af sprogmodellens non-deterministiske sandsynlighedsdistribution (Fergus et al. 2023, Ouyang et al. 2023), skulle det i princippet give os forskellige tekster. Under indsamlingen af teksterne stod det dog klart, at chatbotten kun til en vis grad varierer sine svar sprogligt, og vi vurderede på den baggrund, at ti tekster var nok til at fungere som et sample, der kunne repræsentere en prototypisk besvarelse samt sproglige variationer herover.

5. DEN KVANTITATIVE ANALYSE

I de følgende afsnit gennemgår vi resultaterne fra vores analyser. Først de kvantitative, dernæst de kvalitative. Derefter opsummerer vi vores fund og diskuterer dem.

5.1 Afsnits- og periodelængde

I skema 1 ses målingerne af de korteste og længste afsnits- og periode-længder i korpuset. De er beregnet vha. automatisk optælling. Længdevariationen (fra korteste afsnit og korteste periode til længste afsnit og længste periode) er størst i korpuset med de studerendes tekster og markant mindre i korpuset med ChatGPT's tekster.

SKEMA 1. KORTESTE OG LÆNGSTE AFSNIT SAMT PERIODER I TEKSTKORPUSSENE

	Målkorpus (ChatGPT)		Referencekorpus	
	Korteste	Længste	Korteste	Længste
Afsnit (ord per afsnit)	39	105	4	396
Afsnit (periode per afsnit)	2	5	1	19
Periode (ord per periode)	7	39	4	80

Ser man på de enkelte tekster, så indeholder 97 % af de studerendes tekster afsnit, der i deres ordlængde overskrider enten de korteste eller længste afsnit i ChatGPT-korpusset. Hvad angår periodelængde, er procentdelen lavere, men dog fortsat i den høje ende: 76 % af de studerendes tekster indeholder perioder, der i deres længde overskrider enten de korteste eller længste perioder skrevet af ChatGPT i korpusset. Resultatet viser, at de studerendes tekster altovervejende er præget af større udsving i de korteste og længste afsnit og til dels også i de korteste og længste perioder. Dette stemmer overens med observationerne i Desaire et al. (2023).

Samme mønster genfinder man i teksternes gennemsnitlige afsnits- og periodelængder og standardafvigelserne herfra. Resultaterne i skema 2 viser, at selvom den gennemsnitlige afsnits- og periodelængde (μ) ligger relativt tæt på hinanden, når man sammenligner de studerendes tekster og ChatGPT's tekster, så er der markant variationsforskel, når man kigger på den mindste (min.) og den største (maks.) gennemsnitlige afsnits- og periodelængde. Standardafvigelsen (σ) viser samtidig, at der er langt større variation i afsnit- og periodelængden i de studerendes tekster sammenlignet med teksterne, der er genereret af ChatGPT. Den højeste gennemsnitlige standardafvigelse finder man i de studerendes tekster, den laveste hos ChatGPT. Det samme gælder minimums- og maksimumsværdierne af standardafvigelserne.

SKEMA 2. GENNEMSNITLIG AFSNITS- OG PERIODELÆNGDE SAMT STANDARDAFVIGELSER

	Gennemsnitlig længde						Standardafvigelse					
	Målkorpus (ChatGPT)			Referencekorpus			Målkorpus (ChatGPT)			Referencekorpus		
	Min.	Maks.	μ	Min.	Maks.	M	Min.	Maks.	σ	Min.	Maks.	σ
Ord per afsnit	57	83	67	31	195	68	8,5	29,9	18,3	12,3	147,1	57,3
Perioder per afsnit	2,8	4,3	3,3*	1,5	11,5	3,8*	0,5	1,6	1,1	0,5	10,6	3,1
Ord per periode	16	22,1	20,1*	12,8	34	18,3*	4,2	8,6	6,2	4,5	17	9,5

Min. og maks. = mindste og største gennemsnitlige længde og standardafvigelse i en korpus tekst.

μ = gennemsnitlig periodelængde på tværs af alle korpussets tekster.

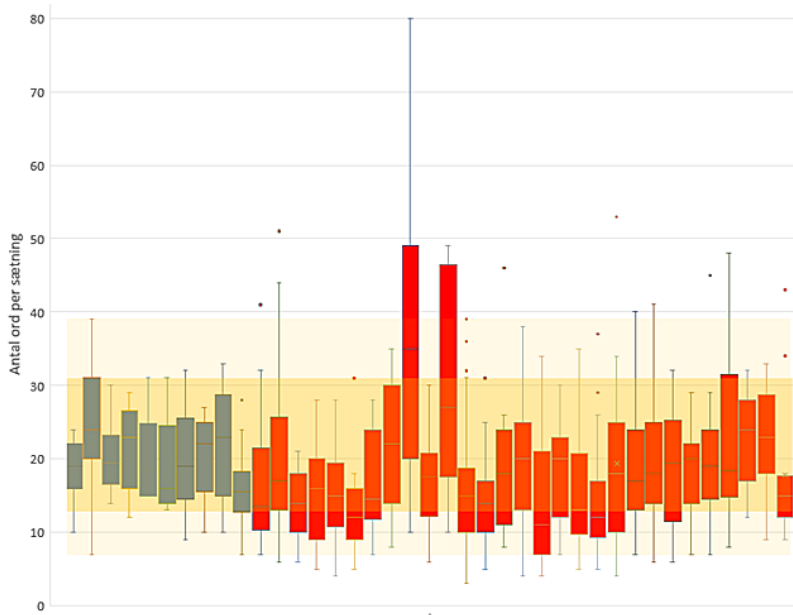
σ = standardafvigelse på tværs af alle korpussets tekster.

** = signifikant forskel på den gennemsnitlige længde ($P < 0.05$), T-testet.*

Ensartetheden i afsnits- og periodelængde synes altså at være en del af ChatGPT’s tekstlige fingeraftryk, hvad andre studier også har påvist (AlAfnan & MohdZuki 2023, Muñoz-Ortiz et al. 2023, Desaire et al. 2023). Til sammenligning er de studerendes tekster generelt markant mere varierede i afsnits- og periodelængde. Dette betyder dog ikke, at man i alle tilfælde kan bruge spredning i helsætningslængde som mål til at skelne de studerendes tekster fra ChatGPT’s tekster. Nogle af de studerendes tekster ligger nemlig inden for eller meget tæt på samme variationsspektrum som ChatGPT-teksterne.

Dette er eksemplificeret i figur 1, der i et boksplot viser den øvre og nedre kvartil for korpussets periodelængder samt de yderste observationer, dvs. de korteste og længste perioder i teksterne. En femtedel af de studerendes tekster ligger her inden for de samme øvre og nedre kvartiler som ChatGPT-teksterne. I disse tilfælde ligner de studerendes begrænsede variation i periodelængde altså chatbottens.

FIGUR 1. BOKSPLOT OVER PERIODELÆNGDERNE I KORPUSSENE



Gråblå farve = ChatGPT-tekster; rød farve = de studerendes tekster; mørkere gul transparent farve = grænsen for ChatGPT-teksternes højeste øvre og laveste nedre kvartil; lysere gul transparent farve = grænsen for de højeste og laveste ydre observationer i ChatGPT-teksterne. Prikkerne viser 'ekstreme' periodelængder (outliers), der er markant afvigende fra den typiske periodelængde i teksterne.

5.2 Nøgleordsfrekvenser og kollokationsmønstre

For at finde frem til potentielle frekvensforskelle i ordvalget mellem de to korpusser har vi lavet en sammenlignende nøgleordsanalyse (keyword-liste) i AntConc. Keyword-funktionen i programmet beregner automatisk og på baggrund af en χ^2 -test, hvilke ord der er usædvanligt frekvente i målkorpusset sammenlignet med referencekorpusset – og omvendt. Resultaterne er listet i nedenstående skema 3.

SKEMA 3. NØGLEORDSANALYSE

	Målkorpus (ChatGPT)	Referencekorpus
<i>til</i>	44	26,7
<i>og</i>	35	17,4
<i>refererer</i>	19,9	4,8
<i>linje</i>	19,5	7,2
<i>pronominer</i>	17,3	6,5
<i>sammenhæng</i>	16,5	2,1
<i>eksempel/eksempler</i>	12	2,1
<i>skabe</i>	9,7	2,2
<i>sproglig</i>	8,6	1,2
<i>nyhedsartiklen</i>	7,9	1,5
<i>dog</i>	7,1	1,9
<i>klar</i>	4,5	0,6
<i>bidrager</i>	4,1	0
<i>entydige</i>	4,1	0,6
<i>set</i>	3,8	0,3
<i>fleste</i>	2,3	0,1
<i>nævnte</i>	2,3	0
<i>man</i>	0	7,9
<i>vi</i>	0	6,3
<i>bliver</i>	0	4,6
<i>jeg</i>	0	2,3
<i>har</i>	0,4	4,9

Ordfrekvens per 1.000 ord; alle forskelle er signifikante ($P < 0.05$)

Frekvensforskellene i den øverste del af skemaet (frem til den tomme skillelinje) skyldes i stort set alle tilfælde en række ensartede frasemønstre i de tekster, der er skrevet af ChatGPT. Eksempelvis indgår præpositionen *til* i en række frasevariationer, hvor præpositionen kombineres med verber som *referere* (*referere til*) og *bidrage* (*bidrage til*) med efterfølgende ensartede styrelser i middelbare objekter. Nedenstående skema 4 illustrerer en række af de ensartede variationer i en analyse af kollokationsmønstre lavet ved hjælp af Antconc. Listen viser de ord, som det centrale søgeord, *til*, hyppigst optræder sammen med i korpusset.

Farvekodningen er vores egen og fremhæver det frasemønster, som er særligt fremtrædende og væsentligt for vores analyse.

SKEMA 4. EKSEMPLER PÅ KOLLOKATIONSKÆDER MED NØGLEORDET *TIL* I MÅLKORPUSSET

Kontekst	Ord	Kontekst
Teksten viser, hvordan disse elementer bidrager	til	at skabe sammenhæng og helhed i en tekst ved ...
De fleste referencer er entydige og bidrager	til	at skabe sammenhæng og helhed i teksten ...
... herunder hvordan pronominer bidrager	til	at skabe sproglig sammenhæng.
... da de fleste referencer er entydige og bidrager	til	en klar forståelse af tekstens indhold ...
... er brugen af pronominer i artiklen vellykket, da de bidrager	til	en klar og entydig sammenhæng mellem tekstelementer.
... i nyhedsartiklen generelt vellykket og bidrager	til	en klar og entydig sammenhæng.

Som det ses, formulerer ChatGPT sig frasemæssigt set relativt ensartet på tværs af teksterne, hvilket man ikke på samme måde ser i de studerendes tekster. Denne ensartethed i ordvalg dækker i flere tilfælde over en tilsvarende ensartethed i chatbottens måde at strukturere sin analyse og bygge sine argumenter op på, hvad vi vender tilbage til senere. Her vil vi til en start blot vise, hvordan chatbotten frasemæssigt set formulerer sig meget ensartet, når den lægger op til at angive forbehold i sin argumentation. Det gør den vha. det argumentative sætningsadverbium *dog*, der indgår i alle chatbottens tekster, jf. eksemplerne i kollokationslisterne i skema 5.

SKEMA 5. KOLLOKATIONSKÆDER MED NØGLEORDET *DOG* I MÅLKORPUSSET

Kontekst	Ord	Kontekst
Der er	dog	også eksempler på uklarheder i referencerne ...
Der er	dog	også eksempler på upræcis eller tvetydig brug af ...
Der er	dog	plads til forbedring i enkelte tilfælde ...
Der er	dog	plads til forbedringer, som i det nævnte tilfælde ...
	dog	er der et eksempel på upræcis reference ...
	dog	er der et eksempel, hvor entydigheden kunne forbedres ...
	dog	er der enkelte steder i teksten, hvor brugen af pronominer ...

Som det ses, indgår *dog* i ensartede frasemønstre sammen med det formelle subjekt *der* og kopulaverbet *være*. Man genfinder ikke dette mønster i samme udstrækning hos de studerende, men det forekommer.

Man finder også faste kollokationsmønstre andre steder i teksterne end ved nøgleordene, fx når chatbotten redegør for teoriens begreber. Her kombinerer den i flere tilfælde lekset *begreb* med den passive form af verbet *præsentere*, jf. eksemplerne i skema 6:

SKEMA 6. KOLLOKATIONSKÆDER MED LEMMAET *BEGREB* I MÅLKORPUSSET

Kontekst	Ord	Kontekst
I [teksten] præsenteres en række væsentlige	<i>begreber</i>	der belyser, hvordan pronominer ...
I [teksten] af Lis Holm præsenteres væsentlige	<i>begreber</i>	der beskriver, hvordan pronominer ...
I [teksten] af Lis Holm præsenteres	<i>begreber</i>	som kohæsion, kohæsiionskæder, proformer ...
I Lis Holms tekst [...] præsenteres de væsentligste	<i>begreber</i>	inden for sproglig sammenhæng og kohæsion ...
I [teksten] af Lis Holm præsenteres	<i>begreberne</i>	kohæsion og proformer ...

Brugen af s-passiv er værd at lægge særlig mærke til, da chatbotten helt undgår at bruge den perifrastiske *blive*-passiv som variant. Til sammenligning anvender de studerende i rigt mål *blive*-passiven (jf. frekvensforskellene på verbet *bliver* i den nederste del af skema 3 ovenfor), for eksempel i kombinationer med verbet *brugt*.

SKEMA 7. KOLLOKATIONSKÆDER MED NØGLEORDET *BLIVER* I REFERENCEKORPUSSET

Kontekst	Ord	Kontekst
Det pronomen som	<i>bliver</i>	brugt mest, er “det”.
Et eksempel, der	<i>bliver</i>	brugt i teksten er følgende citat ...
Generelt hænger teksten meget godt sammen, da der	<i>bliver</i>	brugt mange navne ...
Artiklen indeholder en del proformer, og de	<i>bliver</i>	brugt fornuftigt ...
Ligeledes kan vi gå ned i linje 23, hvor stedordet “ham”	<i>bliver</i>	brugt korrekt som reference.

Til sammenligning bruger chatbotten konsekvent s-passiv i stedet som illustreret i skema 8.

SKEMA 8. EKSEMPLER PÅ BRUGEN AF S-PASSIV I MÅLKORPUSSET

Kontekst	Ord	Kontekst
I linje 22 og 23	bruges	“hans” og “ham” som referencer ...
For eksempel	bruges	“det” (linje 7) til at referere tilbage til kommavaner ...
I linje 46	bruges	“ham” til at referere til en tidligere nævnt person ...
For eksempel	bruges	“de” i linje 3 og 13 for at referere til Dansk Sprognævn ...
Gennem hele teksten	bruges	pronominer som “det”, “de”, “han” og “deres” ...

Dette er et tegn på, at chatbotten i korpuset formulerer sig i et mere formelt stilleje end de studerende, da s-passiven forekommer hyppigere i formelt skriftsprog og *blive*-passiven hyppigere i talenært skriftsprog (Løj & Wille 1985: 17).

En anden afgørende forskel ses i brugen af pronominerne *man*, *jeg* og *vi*, som kun de studerende og ikke ChatGPT bruger. Dette skyldes først og fremmest forskelle i den måde, de studerende og chatbotten positionerer sig sprogligt på i teksterne, hvad vi vender tilbage til senere i analysen.

Samlet set har vi altså fundet en række generelle frekvensforskelle i nøgleord og frasemønstre, der først og fremmest skyldes, at chatbotten formulerer sig relativt ensartet på tværs af sine tekster.

5.3 Afvigelser fra retskrivningsnormen

Korpuset med studentertekster har som nævnt tidligere været anvendt i videnskabelige undersøgelser af universitetsstuderendes skriftsprog (Blom et al. 2017; Rathje 2018), og resultaterne herfra har påvist et højt antal sprogfejl i de studerendes tekster. I gennemsnit begår de studerende 64 sprogfejl per tusind ord i teksterne med en standardafvigelse på 31. Den tekst, der indeholder færrest fejl, rummer i alt 7,8 fejl per tusind ord, mens teksten, der indeholder flest, rummer 145,1 fejl per tusind ord.

I ChatGPT-teksterne forekommer der til sammenligning kun 3 sprogfejl i gennemsnit per tusind ord, og ingen tekster rummer flere sprogfejl end mindstemålet hos de studerende. Det betyder i denne giv-

ne sammenhæng, at en høj frekvens af sprogfejl afslører det menneskelige fingeraftryk i teksten. Det skal i øvrigt bemærkes, at de studerende havde mulighed for sprogteknologisk hjælp i form af automatisk stavekontrol, da de skrev opgaverne.

Blandt de få sproglige afvigelser, som chatbotten laver, er der eksempler på kommafejl, manglende ord, syntaksfejl og idiomatiske fejl, herunder en gennemgående tendens til formuleringen *før at* i tilfælde, hvor det normativt set regnes for (mere) korrekt at skrive *til at* (bemærk, at alle kursiveringer i analysens eksempler er vores egne):

(C7) For eksempel bruges ”de” i linje 3 og 13 *før at* referere til Dansk Sprognævn, og ”det” i linje 11 og 21 *før at* henvise til kommaer og kommaregler. [...] Disse eksempler på pronominer fungerer godt *før at* skabe sammenhæng og gøre teksten mere flydende og læsevenlig.

Overordnet set finder vi altså, at chatbotten skriver væsentligt mere korrekt end de studerende, og at chatbotten kun i et enkelt tilfælde afviger systematisk fra normen.

6. DEN KVALITATIVE ANALYSE

6.1 *Selvreferencer*

Som nævnt havde vi en forventning om, at ChatGPT ville gøre brug af interaktionelle ressourcer på en anden måde end de studerende, der som udgangspunkt burde have et mere fyldestgørende overblik over skrive-situationen og de genreforventninger, der knytter sig til opgavetypen. I korpusserne sporer vi da også en række væsentlige forskelle i den måde, hvorpå chatbotten og de studerende positionerer sig og engagerer læseren. Dog synes forskellene i flere tilfælde at skyldes forskellige tilgange til de sproglige genrekonventioner i akademiske tekster snarere end en særlig opmærksomhed på og bevidsthed om skrivesituationen blandt de studerende.

En markant forskel på de to korpusser ses i den måde, de studerende positionerer sig på, når de anvender metakommunikation til at forklare, hvordan opgaven er bygget op, og hvad de vil gøre eller gør i teksten. Her træder de studerende ofte frem i form af et ekspliceret *jeg*,

der fungerer som en form for skribentrepræsentation, hvis opgave det er at navigere læseren rundt i teksten og skabe overblik.

(S2) *Jeg* vil i denne opgave foretage en analyse af den sproglige sammenhæng i nyhedsartiklen [...]. Dette vil *jeg* gøre med fokus på brugen af pronominer, hvorfor *jeg* vil starte ud med en kort redegørelse for pronominers betydning for en tekst. Dernæst vil *jeg* analysere avisartiklens brug af ordet ”det” og hvad det betyder for tekstens sproglige sammenhæng. Afslutningsvis vil *jeg* kort vurdere, hvor vellykket teksten er i sin brug af pronominer.

(S8) Før *jeg* begynder min analyse af nyhedsartiklen vil *jeg* kort ridse op hvilke begreber *jeg* er blevet introduceret for i den foregående teoritext [...]. Det vil *jeg* gøre da det er ud fra disse rammer *jeg* senere vil analysere den udvalgte nyhedsartikel.

De metatekstlige forklaringer vidner om, at de studerende orienterer sig mod og tager hensyn til læserens behov for overblik i teksten. Samtidig får teksten en personlig og dialogisk stemme (Tardy 2012), når den studerende lader sit tekstlige *jeg* forklare læseren, hvad der gøres i teksten.

Til sammenligning forklarer chatbotten i ingen tilfælde, hvordan opgaven er bygget op, og den forklarer kun i meget begrænset grad, hvad den vil gøre, gør eller har gjort, og i de få tilfælde positionerer den sig udpræget distanceret ved hjælp af passiver og verbalsubstantiver og uden brug af diskursmarkører:

(C1) Nyhedsartiklen [...] *analyseres* med særligt fokus på, hvordan den hænger sammen gennem brugen af pronominer.

(C2) En *analyse* af nyhedsartiklen [...] viser, at artiklen generelt skaber sammenhæng gennem brugen af pronominer og andre kønshæstive elementer.

Det skal bemærkes, at chatbotten godt kan omtale sig selv som *jeg* i andre sammenhænge, så den upersonlige og distancerende positionering skyldes ikke, at den undgår pronomeneret som default. Der er snarere tale

om, at den bagvedliggende sprogmodel i denne sammenhæng har opfattet det som mest sandsynligt, at metasproglige handlinger udtrykkes i passiv eller nominaliseret form i akademiske tekster. Det kan tyde på, at den kunstige intelligens fortrinsvis er fodret med tekster, der har en sprogkonservativ tilgang til tekstgenren, hvor brugen af *jeg* opfattes som en slags brud på akademisk dekorum. Chatbottens brug af s-passiver og nominaliseringer – der også observeres af Herbold et al. (2023) – bidrager altså samlet set til billedet af dens sproglige profil som upersonlig og distanceret.

6.2 Styrke- og attitudemarkører

Den subjektive positionering kan også iagttages i andre sammenhænge, fx ved nedtoning af vurderende påstande, hvorved den studerende angiver, at vurderingen ikke nødvendigvis behøver at være den eneste gyldige, fx:

(S25) I bilag 2 hænger brugen af pronominer relativt fornuftigt sammen *sådan, som jeg ser det*.

På lignende vis positionerer nogle af de studerende sig eksplicit subjektivt, når de bevæger sig fra analyse til vurdering og kommer med 'synsninger' om teksten:

(S8) *Jeg synes* ikke at teksten får brugt alle sine pronominer på en vellykket facon.

(S25) Som helhed *synes jeg*, at teksten lykkedes med at bruge pronominerne på den rigtige måde.

De studerende markerer på denne måde, at det er *deres* vurdering, der her kommer til udtryk, hvilket efterlader plads til, at andre kunne tænkes at mene noget andet. Der er altså dermed skabt et lille epistemisk rum for mulig usikkerhed.

Normativt set regnes den slags ekspliciterede personlige 'synsninger' typisk som et brud på den akademiske genrekonvention. Rienecker et al. anbefaler fx, at man som studerende holder sig til "saglige akademi-

ske fremstillingsformer” (2008: 58), og at man bør være ”meget varsom med at ”mene”” (ibid.) i akademiske tekster. Evaluerende påstande i akademiske opgaver er normativt set kendetegnet ved, at de bør baseres på ”faglige kriterier” (Rienecker et al. 2008: 58), bakkes op af belæg og – i tilfælde af usikkerhed – reguleres med styrkemarkører eller gendrivelses, hvis der er særlige forbehold at tage. Denne typiske argumentationsteknik mestrer ChatGPT, når den på baggrund af sine analyser fremsætter evaluerende konklusioner, fx:

(C8) [*Opsummerende påstand:*] Generelt er teksten vellykket i sin brug af pronominer til at skabe sproglig sammenhæng. [*Belæg med afsæt i analysen:*] Pronominerne refererer entydigt til tidligere nævnte ord og udtryk, hvilket gør det nemt for læseren at følge med i teksten og forstå dens budskab. [*Delvis gendrivelse med afsæt i analysen:*] Dog er der et eksempel på en tvetydig reference i linje 46, hvor ”ham” kan referere til både Jørn Lund og Jacob Thøgersen. [*Alternativ løsning:*] Dette kunne have været undgået ved at gentage det relevante navn i stedet for at bruge et pronomen.

Som optakt til den opsummerende påstand bruger chatbotten adverbiet *generelt* som en nedtonende styrkemarkør, der præsupponerer, at der er – eller kan være – undtagelser til påstanden. Og senere giver chatbotten netop et eksempel på en sådan undtagelse, som den bruger til delvist at gendrive påstanden. Gendrivelsen markeres vha. det argumentative sætningsadverbium *dog*, der som vist i kollokationsanalysen bruges meget hyppigt i denne funktion.

Som oftest afviger de studerende fra denne argumentatoriske idealnorm. Mange undlader for eksempel helt at komme med en opsummerende og evaluerende påstand, hvilket gør, at læseren selv skal udlede den på baggrund af analysen. I andre tilfælde fremsætter de studerende afslutningsvis en vurderende og til tider ladet – nærmest anmelderagtig – påstand uden hverken at nedtone eller gendrive den, fx:

(S2) Det er en dårlig artikel, der ikke lever op til artiklers vanlige struktur.

Igen er dette et eksempel på, at nogle af de studerende afviger fra den normale genrekonvention, når de skal positionere sig i forhold til de påstande, de fremsætter. Det er altså først og fremmest i afvigelserne fra normen, at man bemærker forskellen på chatbotten og de studerende.

6.3 Engagementsmarkører

Den mest direkte måde at engagere læseren på i teksten er ved at bruge pronomenet *du*. Det afholder både de studerende og chatbotten sig dog fra i teksterne. I teksternes redegørende og analyserende dele forsøger de studerende til gengæld flere steder at engagere læseren vha. af pronomenet *vi* og i nogle tilfælde også indirekte direktiver:

(S1) Hvis *vi* kigger på første sætning af artiklen, støder *vi* allerede på *vores* første proform.

(S14) Dette fører *os* til det et andet begreb i teksten, kohæsiionskæde.

(S21) Her kan *vi* se at artiklen gør brug af proform.

Som det ses, anvender de studerende *vi*-formen, når de hjælper læseren med at navigere rundt i redegørelsen og analysen. De engagerer og guider altså læseren på en og samme tid, hvilket er en kombination af interaktionelle og interaktive ressourcer. Brugen af det kollektive *vi* – og fravalget af det individuelle *du* – i studerendes opgaver er også observeret af Hyland (2004: 139-140). Normativt set er det dog omdiskuteret, hvorvidt man i akademiske tekster bør inddrage læseren vha. det kollektive *vi*. Rienecker & Jørgensen (2010: 354) kalder det for eksempel for ”krampagtigt og forkert”, når studerende bruger uspecificeret *vi* (eller *man*) i tilfælde, hvor de i virkeligheden mener *jeg*.

ChatGPT anvender intetsteds de kollektive og inkluderende pronominer, men går i stedet typisk direkte til den propositionelle pointe i sine redegørelser og analyser, ofte i en definerende stil, der parafraserer kildeteksten i en næsten ordret form tenderende til plagiat, fx:

(C3) “Kohæsion er betegnelsen for forbindelsen mellem tekstelementer, når fortolkningen af et element er afhængig af et andet element i teksten.” (jf. kildeteksten: ”Kohæsion er betegnelsen for den forbindelse, der eksisterer, når fortolkningen af et tekstelement er afhængig af et andet element i teksten”).

På dette punkt vælger ChatGPT altså generelt en genrekonservativ og distanceret positionering over for modtageren, mens flere af de studerende kan siges at være enten bevidst genreprogressive eller ubevidste om den traditionelle genrekonvention.

Man ser også, hvordan de studerende – i modsætning til chatbotten – vælger at positionere sig generisk vha. af pronomenet *man*. Det gør de først og fremmest i tilfælde, hvor de har brug for at anlægge et mere alment læserperspektiv, der illustrerer, hvordan en forestillet læser kunne tænkes at forstå et træk ved den analyserede tekst, fx:

(S22) Nu bliver *man som læser* i tvivl om, hvem ham er en anaforsisk reference til.

(S24) I nyhedsartiklen er der brugt mange pronominer, hvilket gør at *man som læser* er tvunget til at forstå handlingen fra starten.

Langt de fleste af referencekorpussets forekomster af *man* kollokerer med *som læser*, så her er der tale om et hyppigt forekommende frasemønster i de studerendes tekster. Til sammenligning forekommer *man* slet ikke i ChatGPT-korpusset.

Selvom ChatGPT ikke gør meget ud af at engagere læseren sprogligt i metadiskursen, ser vi et eksempel, hvor den bryder stilen og giver et indirekte og normativt direktiv:

(C6) *Det er vigtigt at huske*, at en god tekst bør gøre det klart, hvilket andet ord i teksten en proform refererer til, og sikre entydige referencer.

6.4 Koblinger og diskursmarkører

ChatGPT-teksterne er altovervejende præget af en stærkt asyndetisk stil

på helsætningsniveau, dvs. at chatbotten i langt de fleste tilfælde undlader at anvende parataktiske konjunktioner til at binde helsætningerne sammen, og ofte bruger den heller ikke andre koblingsord til at udtrykke semantiske relationer mellem helsætningerne. Et lignende fund finder man som tidligere nævnt hos Desaire et al. (2023), hvad angår adversative og kausale koblingsord.

Den asyndetiske stil resulterer flere steder i et lidt staccato-præget og – fristes man til at skrive – maskinelt informationsflow, hvor sammenhængen mellem helsætningerne alene skabes gennem kohæsion og endoforiske referencer, fx:

(C1) Kohæsion, kohæsionsredskaber og kohæsionskæder er centrale begreber i forståelsen af sammenhæng i tekster. Proformer, som ofte er funktionsord og en del af pronominerne grammatiske ordklasse, er sprogbesparende midler, der refererer til andre ord i teksten og skaber sammenhæng.

Nyhedsartiklen ”Det skal de undersøge” er et eksempel på en tekst, hvor pronominer anvendes for at skabe sammenhæng. Teksten beskriver Dansk Sprognævns planer om at undersøge danskernes kommavaner.

Man ser dog også en lignende stil hos flere af de studerende, særlig i opgavens redegørende passager, så den asyndetiske og koblingsfri stil er ikke et unikt særtræk hos chatbotten. Til gengæld ser man en langt mere udfoldet brug af diskursmarkører i de studerendes tekster. Her bruges markører til at navigere læseren rundt i diskursens flow, fx:

(S13) *Først* vil jeg redegøre for de vigtigste begreber i teoriteksten og *derefter* vil jeg analysere nyhedsartiklen.

Der er i dette tilfælde tale om en kombination af både interaktionel positionering (*jeg*) og interaktiv guidning (*først ... derefter*) i metadiskursen, som man ikke finder hos chatbotten. Selvom det er meget sparsomt med diskursmarkører hos chatbotten, ser vi dog – ligesom Herbold et

al. (2023) – at chatbotten bruger markører til at signalere overgangen fra analysen til den opsummerende og vurderende konklusion:

(C1) *Sammenfattende* er brugen af pronominer i nyhedsartiklen [...] generelt vellykket ...

(C2) *Overordnet set* er nyhedsartiklen vellykket i sin brug af pronominer ...

(C5) *Samlet set* er teksten vellykket i sin brug af pronominer ...

Samme træk forekommer også i flere af de studerendes tekster.

7. OPSUMMERING AF ANALYSENS FUND

Den sproglige ophavsanalyse har vist, at der er en række forskelle på de studerendes og ChatGPT's skrivestil. Chatbotten har generelt set en mere ensartet måde at skrive på end de studerende. Det gælder både i afsnits- og periodelængde, ordvalg, frasemønstre og i den til tider udprægede asyndetiske stil med sparsom brug af koblingsord. En del af dette skyldes muligvis, at de generative sprogmodeller er bygget statistisk op, sådan at chatbotten forudsiger den mest sandsynlige fortsættelse på en given tekststreng. Det skal dog nævnes, at der også er studerende, der skriver relativt ensartet i korpusset, så man kan ikke nødvendigvis afgøre, hvem der har skrevet teksten, alene ud fra dette mål. Analysen viser desuden, at chatbotten kun laver ganske få ortografiske normafvigelser, hvilket står i kontrast til de studerende, der begår langt flere sprogfejl.

Samtidig følger chatbotten en genrekonservativ tilgang til akademisk sprogbrug, hvor den i modsætning til flere af – men ikke alle – de studerende formulerer sig distanceret og helt undgår pronominerne *jeg*, *vi* og også det mere generiske *man*. I stedet anvender den i højere grad s-passiver og nominaliseringer, der trækker det personlige ud af teksten. I modsætning til flere af de studerende undlader den samtidig at eksplicitere personlige synsninger og vælger i stedet i en mere genrekonventionel stil at nedtone ('hedge') og gendrive sine evaluerende påstande. Analyserne viser desuden, at chatbotten i modsætning til flere af de

studerende undlader at metakommunikere om, hvad den gør i teksten, og den anvender heller ikke i samme omfang diskursmarkører til at navigere læseren rundt i teksten.

Det er samtidig en væsentlig observation, at chatbotten har en tendens til at parafrasere opgavens teoretiske sekundærtekst stort set ordret og derfor bevæger sig tæt på eller måske endda over grænsen for plagiat. Skulle en studerende – i det skjulte og dermed udeklareret – vælge at bruge en chatbot til at redegøre for en teori i en opgave, risikerer vedkommende altså en form for dobbeltplagiat. Den studerende plagierer chatbotten, der igen plagierer teoriteksten.

Samlet set dokumenterer vores studie, at det er muligt at finde spor af ChatGPT's sproglige fingeraftryk, når den skriver på dansk, ligesom det også er tilfældet på engelsk (Desaire et al. 2023, Herbold et al. 2023, Muñoz-Ortiz et al. 2023). Dette giver forhåbninger om, at det med tiden måske vil være muligt at udvikle redskaber, der kan spore chatbotten på dens skrivestil. Der er dog en række potentielle fejkilder og visse begrænsninger i vores analyse – og også i de øvrige omtalte studier – der skal tages forbehold for.

8. MULIGE FEJKILDER OG ANALYSENS BEGRÆNSNINGER

For det første kan fundet af de ensartede afsnits- og periodelængder samt frasemønstre i ChatGPT-teksterne skyldes, at teksterne er baseret på samme prompt, og/eller at promptet er sendt fra en og samme IP-adresse. Vi ved, at der er en indbygget tilfældighed i de svar, chatbotten giver (Fergus 2023), men det er samtidig muligt – og for frasemønstrenes tilfælde endda meget sandsynligt – at den sproglige variation mellem svarene ville være større, hvis promptet var forskelligt.

For det andet fokuserer vores studie alene på ChatGPT, der lige nu er den dominerende sprogmodel på markedet. Men der findes også andre konkurrerende chatbotter, der er trænet med andre tekster. Det er muligt, at andre chatbotter har en (helt) anden og mere varieret skrivestil end ChatGPT.

For det tredje udvikler ChatGPT sig hele tiden, i og med sprogmodellen løbende opdateres og fordres med nye tekster. Som påpeget af Herbold et al. (2023) er chatbotterne et "moving target", et mål i konstant bevægelse, og det gør hurtigt forskningsresultaterne på området

forældede og i nogle tilfælde sporingsundersøgelserne vanskelige at replicere. Vi kan allerede nu se, at ChatGPT i dag besvarer vores stillede opgaveformulering anderledes, end da vi indsamlede data.

For det fjerde, så er det muligt at prompte chatbotten til at imitere flere af de sproglige kendetegn, vi har fundet frem til som særkende i studerendes tekster. Det har vi testet ved at supplere opgaveformuleringen med en række prompts. Her følger et par udvalgte eksempler.

Supplerende prompt: “Sørg for, at din besvarelse har en varieret sætningslængde og indeholder nogle korte sætninger på under 5 ord og et par lange sætninger på mellem 60 og 80 ord. Teksten skal samtidig indeholde et helt kort afsnit på under 20 ord, og et langt afsnit på mellem 200 og 250 ord.”

Dette lever chatbotten derefter op til sin besvarelse, så længdevariationen bedre matcher de studerendes. For eksempel forklarer den nu begrebet kohætion i helt korte perioder.

ChatGPT: “Kohætion er essentiel. Det binder teksten sammen.”

Vi har også promptet chatbotten til at forklare, hvad den gør, og bruge *jeg* til at referere til sig selv.

Supplerende prompt: “Forklar indledningsvis, hvordan du besvarer de spørgsmål, der er stillet. Brug ordet “jeg”, når du forklarer, hvad du gør.”

Ud fra denne instruktion anvender chatbotten nu både interaktionelle og interaktive ressourcer og forklarer ligesom flere af de studerende, hvad den vil gøre i teksten, samtidig med den refererer til sig selv, fx:

ChatGPT: “For at besvare de stillede spørgsmål vil jeg først gennemgå de vigtigste begreber fra [teoriteksten] med fokus på, hvordan pronominer skaber sammenhæng i en tekst. Jeg vil derefter analysere brugen af pronominer i nyhedsartiklen [...] for

at vurdere, hvordan de bidrager til tekstens sammenhæng. Til sidst vil jeg vurdere, hvor effektivt pronominer er anvendt i denne artikel.”

Som det ses, findes der altså modtræk, der kan sløre chatbottens sproglige fingeraftryk, både hvad angår kvantificerbare mål og de kvalitative træk i teksten. Samtidig er det muligt for studerende at fodre ChatGPT med deres egne tidligere opgavetekster og bede chatbotten om at imitere skrivestilen i dem, når den besvarer nye opgaver. Og endelig kan de studerende også vælge at kombinere dele af egne tekster med dele af chatbottens eller parafrasere chatbotten frem for at kopiere tekst direkte fra den. Igen kan dette vanskeliggøre en pålidelig sporing.

Endelig – og for det femte – er teksterne i referencekorpusset skrevet af nystartede studerende, hvoraf flertallet må formodes at være relativt urutinerede skribenter. Dette kan forklare, hvorfor der forekommer langt flere sproglige fejl hos de studerende end ChatGPT i vores korpus. Der vil antageligt være danskstuderende længere henne på studiet, som skriver altovervejende eller helt fejlfrit. De studerendes brug af de til tider meget personlige og engagerende positioneringstræk kan muligvis også skyldes manglende rutine og et begrænset kendskab til den mere konservative akademiske genrekonvention. Internationale studier har dog tidligere vist, at såvel studerende som forskere ofte positionerer sig i deres tekster og engagerer læserne (Biber 2006; Hyland 2012). Så selv om den akademiske genre traditionelt set er beskrevet som upersonlig og monologisk snarere end personlig og tilnærmet dialogisk, er dette mere nuanceret end som så.

Samme studier har vist, at der er forskel på studerendes og professionelle forskeres interaktionelle stil, hvilket indikerer, at genreerfaring kan have betydning for den måde, skribenten vælger at positionere sig på. Over tid, og i takt med at de får mere akademisk erfaring, er det muligt, at de studerende vil nærme sig – eller måske fjerne sig yderligere – fra chatbottens stil. Denne udvikling kunne være relevant at belyse med opfølgende studier.

Desuden er der forskelle på genrekonventionerne fra fagområde til fagområde. Hyland (2012: 140) nævner eksempelvis, at brugen af inklusive pronominer er mere udbredt i de ’bløde’ akademiske fagområ-

der, hvor de danskstuderende i vores undersøgelse hører til. Man må derfor forvente, at resultaterne ville kunne tage sig anderledes ud, hvis man undersøgte tekster fra andre fagområder.

9. KONKLUSION OG DISKUSSION

Vores resultater har vist, at det på den ene side er muligt at spore visse sproglige særtræk i ChatGPT's skrivestil, men at det på den anden side også er muligt at sløre disse, fx ved hjælp af supplerende prompts. Vi forudser, at dette kapløb mellem sporing og sløring vil blive et vigtigt tema i de kommende år, når uddannelsessektoren skal forsøge at finde en løsning på udfordringerne i de studerendes brug af ChatGPT. I skrivende stund har regeringens nedsatte ekspertgruppe om ChatGPT og andre digitale hjælpemidler allerede peget på, at det vil være nødvendigt at ændre eksamensformerne i fremtiden, dels for at sikre, at chatbotterne ikke bruges til at snyde, dels for at hindre, at der opstår unødvendig mistænkeliggørelse af elever og studerende (Børne- og Undervisningsministeriet 2023b).

Lige nu ser man forskellige forsøg på løsninger på universiteterne, der i højere grad end skolerne og gymnasierne har haft mulighed for at ændre eksamensformerne hurtigt, så de passer til den nye virkelighed.

Nogle fag har valgt at indføre flere monitorerede stedprøver, der sikrer, at de studerende ikke kan bruge chatbotter til eksamen. Dette vanskeliggør dog udprøvningen af mere komplekse læringsmål, der ikke kan demonstreres mundtligt eller 'på stedet', og samtidig kan løsningsstrategien ikke anvendes ved større skriftlige opgaver som fx bacheloropgaver og specialer.

Andre fag giver de studerende lov til at anvende chatbotter under forudsætning af "god akademisk praksis" (SDU 2023). Det vil sige, at de studerende *forventes* at angive, hvis de har anvendt ChatGPT, forklare, hvad de har anvendt chatbotten til, og samtidig sørge for at citere chatbotten og referere til den, hvis den er blevet parafraseret.

Udfordringen ved at anvende chatbotterne er dog fortsat, at god akademisk praksis normalt bygger på, at det i princippet skal være muligt vha. plagiattjek at tjekke, om denne praksis er efterlevet. Men i og med chatbotten som nævnt kan prompts til at sløre sine prog-

lige spor og tekstfodres til at imitere en sprogbrugers idiolekt, bliver vurderingen af, om den gode akademiske praksis nu også har været fulgt, hurtigt til et spørgsmål om tillid. Altså stoler man som underviser og eksaminator på, at den studerende *ikke* har brugt ChatGPT, eller *kun* har brugt ChatGPT i det omfang, som er angivet? Med tanke på, at plagiat har været et stigende problem på de danske universiteter (Information 2018), er det tvivlsomt, om en sådan forudsat tillid er tilstrækkelig til i fremtiden at sikre retvisende bedømmelser af skriftlige opgaver og specialer. Uddannelsesstederne står derfor over for en meget vanskelig udfordring, der muligvis kun kan løses, hvis man på trods af de nævnte udfordringer formår at udvikle en pålidelig sporingsmetode.

Brugen af chatbotter skaber udfordringer, men rummer også potentialer inden for det skrivepædagogiske og danskfaglige område. Forskningslitteraturen viser, at studerende ofte kæmper med at gennemskue den akademiske genrekonvention og med at finde deres akademiske stemme (Gennrich & Dison 2018). Dette angår i særlig grad metadiskursen, hvor den studerende på den ene side har mulighed for at træde frem og positionere sig i forhold til indholdet og engagere og guide læseren, og på den anden side risikerer at overtræde traditionelle genrekonventioner ved at gøre teksten *for* personlig eller *for* involverende. Med de rette prompts gør chatbotterne det muligt for såvel undervisere som studerende at generere tekstvariationer, der viser, hvordan man som skribent kan bruge – eller fravælge – forskellige interaktive og interaktive ressourcer. Man kan for eksempel, som vi har illustreret, prompte chatbotten til at metakommunikere i en personlig og engagerende stil, eller man kan prompte den til at formulere sig mere upersonligt og distanceret. Alt efter den retoriske situation kan den studerende vælge, hvad der egner sig bedst til formålet, og herefter låne og omskrive eller lade sig inspirere mere frit.

Didaktisk set vil det her være afgørende, at de studerende lærer at skærpe deres stilistiske blik – og gehør – for den polyfoni, der hurtigt vil opstå, når de lader deres sprog smelte sammen med chatbottens. Chatbotten rummer i sig selv mange stemmer, og det er i reglen uigenkendskueligt, hvor de stammer fra. Det kan derfor hurtigt resultere i stilistisk kakofoni med alt for mange stemmer på alt for lidt plads, hvis

man som studerende maser tekstbidder fra chatbotten ind i sin egen tekst uden bearbejdelse og kritisk omtanke.

Af samme grund varslers chatbotternes tilkomst næppe døden for den sproglige del af danskfaget. Tværtimod vil der i de kommende år være endnu mere brug for, at de studerende udvikler deres sproglige og analytisk-kritiske kompetencer, så de kan tilpasse chatbotternes sprog til deres eget og til den situation, som teksten er indlejret i.

På den måde vil chatbotterne kunne styrke – frem for at erstatte – de studerendes tekstproduktion.

Jonas Nygaard Blom, lektor
Center for Journalistik, Institut for Statskundskab, Syddansk Universitet
blom@journalism.sdu.dk

Alexandra Holsting, lektor
Institut for Kultur og Sprogvidenskaber, Syddansk Universitet
aho@sdu.dk

Jesper Tinggaard Svendsen, læsevejleder
UC Syd
JTSV@ucsyd.dk

LITTERATUR

- Aarhus Universitet. 2023. Kend reglerne, når du skal til eksamen. <https://studerende.au.dk/proever/eksamenssnyd/kend-reglerne-naar-du-skal-til-eksamen> (tilgået 30.12.2023).
- Akademikerbladet. 2023. Hver 10. studerende har snydt med ChatGPT: "Man kan ikke forbyde det". Akademikerbladet, 30. juni, 2023. <https://dm.dk/akademikerbladet/aktuelt/2023/juni/hver-10-studerende-har-snydt-med-chatgpt-man-kan-ikke-forbyde-det/> (tilgået 30.12.2023).
- AlAfnan, M.A. & S.F. MohdZuki. 2023. Do Artificial Intelligence Chatbots Have a Writing Style? An Investigation into the Stylistic Features of ChatGPT-4. *Journal of Artificial Intelligence and Technology* 23. 85-94. DOI:10.37965/jait.2023.0267

- Anthony, L. 2023. AntConc (version 4.2.3) [computersoftware]. Tokyo, Japan: Waseda University.
- Biber, D. 2006. *University language: A corpus-based study of spoken and written registers*. Amsterdam: John Benjamins.
- Blom, J.N., M. Rathje, B.L.F. Jakobsen, A. Holsting, K.R. Hansen, J.T. Svendsen & A.V. Lindø 2017. Linguistic Deviations in the Written Academic Register of Danish University Students. *Oslo Studies in Language* 9(3).169-190.
- Børne- og Undervisningsministeriet. 2023a. Brug af digitale hjælpemidler og kunstig intelligens ved folkeskolens prøver. <https://www.uvm.dk/folkeskolen/folkeskolens-proever/proeveafholdelse/brug-af-digitale-hjaelpemidler-og-kunstig-intelligens-ved-folkeskolens-proever>
- Børne- og Undervisningsministeriet. 2023b. Ekspertgruppen om ChatGPT og andre digitale hjælpemidlers foreløbige anbefalinger. <https://www.uvm.dk/-/media/filer/uvm/aktuelt/pdf23/dec/231217-ekspertgruppens-foreloebige-anbefalinger-chatgpt.pdf>
- Cai, Z.G., D.A. Haslett, X. Duan, S.Wang & M.J. Pickering. 2023. Does chatgpt resemble humans in language use? arXiv preprint arXiv:2303.08014.
- Chaka, C. 2023. Detecting AI content in responses generated by ChatGPT, YouChat, and Chatsonic: The case of five AI content detection tools. *Journal of Applied Learning and Teaching* 6(2). 1–11. DOI:10.37074/jalt.2023.6.2.12
- Christensen, T. K. 2017. Ordet fanger. Retslingvistik i en dansk kontekst. *NyS, Nydanske Sprogstudier* 1(52-53). 169–206. <https://doi.org/10.7146/nys.v1i52-53.26334>
- Desaire, H., A.E. Chua, M. Isom, R. Jarosova & D. Hua. 2023. Distinguishing academic science writing from humans of ChatGPT with over 99% accuracy using off-the-shelf machine learning tools. *Cell Reports Physical Science* 4(11). DOI: 10.1016/j.xcrp.2023.101672
- Elkhatat, A.M., K. Elsaid & S. Almeer. 2023. Evaluating the efficacy of AI content detection tools in differentiating between human and AI-generated text. *International Journal for Educational Integrity* 19(17). <https://doi.org/10.1007/s40979-023-00140-5>
- Fergus, S., M. Botha & M. Ostovar. 2023. Evaluating Academic Answers Generated Using ChatGPT, *Journal of Chemical Education* 100(4). 1672–1675. DOI:10.1021/acs.jchemed.3c00087
- Gennrich, T. & L. Dison. 2018. Voice matters: Students' struggle to find voice. *Reading & Writing*, 9. DOI: 10.4102/rw.v9i1.173

- Grieve, J. 2023. Register variation explains stylometric authorship analysis. *Corpus Linguistics and Linguistic Theory* 19(1). 47-77. <https://doi.org/10.1515/cllt-2022-0040>
- Herbold, S., A. Hautli-Janisz, U. Heuer, Z. Kikteva & A. Trautsch. 2023. A large-scale comparison of human-written versus ChatGPT-generated essays, *Scientific Reports* 13, 18617, <https://doi.org/10.1038/s41598-023-45644-9>
- Hyland, K. 2004. Stance and Voice in Written Academic Genres. K. Hyland & C.S. Guinda (red.) *Stance and Voice in Written Academic Genres*. Palgrave Macmillan, 134-150.
- Hyland, K. 2005. *Metadiscourse*. London: Continuum.
- Hyland, K. 2012. Undergraduate understandings: Stance and voice in final year reports. K. Hyland & C.S. Guinda (red.) *Stance and Voice in Written Academic Genres*. Palgrave Macmillan, 134-150.
- Hyland, K., & P. Tse. 2004. Metadiscourse in Academic Writing A Reappraisal. *Applied Linguistics*, 25, 156-177.
- Information. 2018. Flere tages i plagiat: Mange studerende ser ikke afskrift som snyd, 12. november, 2018.
- Khalil, M., & E. Er. 2023. Will ChatGPT Get You Caught? Rethinking of Plagiarism Detection. P. Zaphiris, & A. Ioannou, (red.) *Learning and Collaboration Technologies. HCII 2023. Lecture Notes in Computer Science*, vol. 14040. 475-487. https://doi.org/10.1007/978-3-031-34411-4_32
- Liang, W., M. Yuksekgonul, Y. Mao, E. Wu & J. Zou. 2023. GPT detectors are biased against non-native English writers. *Patterns* 4. DOI:10.1016/j.patter.2023.100779
- Lingard, L. 2023. Writing with ChatGPT: An Illustration of its Capacity, Limitations & Implications for Academic Writers. *Perspectives on Medical Education*, 12(1). 261-270. doi: 10.5334/pme.
- Løj, M., & N.E. Wille. 1985. Kan vi undvære passiv? eller: Kan passiv virkelig undværes? En undersøgelse af passivformernes grammatik, stilistik og pragmatik. *NyS – Nydanske Sprogstudier* 15. 5-42. doi:10.7146/nys.v15i15.13333
- Muñoz-Ortiz, A., C. Gómez-Rodríguez & D. Vilares. 2023. Contrasting Linguistic Patterns in Human and LLM-Generated Text. arXiv preprint arXiv: 2308.09067. <https://doi.org/10.48550/arXiv.2308.09067>
- OpenAI. 2023. ChatGPT, version 4.0. <https://openai.com/> (tilgået 30.12.2023).

- Otterbacher, J. (2023). Why technical solutions for detecting AI-generated content in research and education are insufficient. *Patterns* 4(7). <https://doi.org/10.1016/j.patter.2023.100796>
- Ouyang, S., J.M. Zhang, M. Harman & M. Wang. 2023. LLM is Like a Box of Chocolates: the Non-determinism of ChatGPT in Code Generation. arXiv preprint arXiv:2308.02828.
- Perkins, M., J. Roe, D. Postma, J. McGaughran & D. Hickerson 2023. Game of Tonnes: Faculty detection of GPT-4 generated content in university assessments. arXiv preprint arXiv:2305.18081. <https://doi.org/10.48550/arXiv.2305.18081>
- Rathje, M. 2018. Afvigelser fra retskrivningsnormen i universitetsopgaver. Y. Goldstein, I.S. Hansen & T.T. Hougaard (udg.) 17. *Møde om Udforskningen af Dansk Sprog*, Århus. 415-437.
- Rienecker, L. & P.S. Jørgensen. 2010. *Den gode opgave*. 3. udg. Samfundslitteratur.
- Rienecker, L., P.S. Jørgensen & M. Gandil. 2008. *Skriv en artikel: Om videnskabelige, faglige og formidlende artikler*. København: Samfundslitteratur.
- Rolling Stone. 2023. She Was Falsely Accused of Cheating With AI — And She Won't Be the Last. *Rolling Stone*, 6.6.2023 <https://www.rollingstone.com/culture/culture-features/student-accused-ai-cheating-turnitin-1234747351/> (tilgået 30.12.2023).
- Syddansk Universitet. 2023. Brugen af kunstig intelligens. https://sdunet.dk/da/undervisning-og-eksamen/nyheder/2023/0929_aaiie23 (tilgået 30.12.2023).
- Tardy, C.M. 2012. Current Conceptions of Voice. I K. Hyland & C. S. Guinda (red.). *Stance and Voice in Written Academic Genres* (s. 34-48). Palgrave Macmillan.
- Uniavisen. 2023. "Fremtiden er nu: KU bløder op i regler om AI". <https://uniavisen.dk/fremtiden-er-nu-ku-bloeder-op-i-regler-om-ai/>
- Wu, T.Y., S.Z. He, J.P. Liu, S.Q. Sun, K. Liu, Q.-L. Han, & Y. Tang. 2023. A brief overview of ChatGPT: The history, status quo and potential future development, *IEEE/CAA Journal of Automatica Sinica*, 10(5), pp. 1122–1136, May 2023. doi: 10.1109/JAS.2023.123618

Emojier som parasproglige resurser i skrevet onlineinteraktion

TINA THODE HOUGAARD

ABSTRACT

Denne artikel præsenterer en ny forståelse af emojier, nemlig som en kropsforankret semiotisk resurse. Med afsæt i at se interaktion som en kropsligt funderet udfoldelse og med udgangspunkt i tidligere forskning i 'nyskriftlige' resurser der fungerer som kompensation for det manglende kropssprog, undersøger artiklen hvordan primært emojier kan forstås som erstatning for de parasproglige virkemidler der er fraværende i skriftlig interaktion, dvs. kropsbevægelser, ansigtsudtryk og stemmemæssige forhold. Artiklen giver både et indblik i hvordan emojier virker som præfabrikeret teknologi og system, og et overblik over tidligere forskning i emojiers pragmatiske funktion før den via Peirces tre tegntyper, ikon, symbol og indeks, viser hvordan den fysiske krops sprog manifesterer sig i digitale kropsemojier. Med udgangspunkt i forskning i ansigtsudtryk og tidligere analyser af digital gestik undersøges det hvordan deltagere i forskellige online interaktionssituationer bruger det digitale kropssprog. Der argumenteres for at emojier netop egner sig godt til at formidle følelser og til at guide modtageren i at forstå afsenderens hensigt med det skrevne fordi brugen af emojier og forståelsen af dem er kropsligt funderet. Deltagerne i skrevet onlineinteraktion bruger emojier som kompensation for den parasproglige information fordi emojierne i vid udstrækning ligner og føles som kropslig kommunikation.

EMNEORD: emojier; skriftlig onlineinteraktion; parasprog; mimik; gestik

1 INDLEDNING

Emojier er præfabrikerede billedsymboler der vælges på tastaturet på samme måde som forskellige skrifttyper. Emotikoner og smileyer kan ses som forløbere for emojier (Hougaard 2020). Emotikoner formes vha. tastaturtegn (ASCII-tegn) og står dermed på linjen sammen med den øvrige tekst, fx :- (Dresner & Herring 2010), mens betegnelsen *smiley* kan føres tilbage til 1960'ernes Amerika hvor de gule smileyer bl.a. vandt indpas på T-shirts og badges (Moschini 2016). Selv om emojier ofte betegnes som smileyer (Duncker 2019), repræsenterer de

meget andet end ansigter, fx dyr, naturfænomener, mad, bygninger og flag. Da emojierne kom frem i Vesten i starten af tierne, blev de først set som et ungdomssprogligt fænomen der var med til at gøre kommunikationen overfladisk og doven, men efterhånden er de små tegn blevet accepteret som en integreret del af den multimodale meningsskabelse der foregår i de fleste menneskers daglige digitale sociale praksis, dvs. i digitale beskeder via sms, på Messenger, Instagram etc. Emojiers efterhånden store udbredelse er blevet forklaret på flere forskellige måder: Som en visuel komponent skiller emoji'er sig ud fra teksten ved at kommunikere gennem en anden, mere umiddelbar, modus end sprog idet de "add visual annotations to the conceptual content of a message" (Danesi 2017: 10); emoji'er tiltrækker modtagerens opmærksomhed: "they have more of a referential or pointing function: they draw our attention to a specific idea that the emoji, indirectly, evokes" (Evans 2017: 164), og emoji'er anvendes som en form for faktisk kommunikation idet modtageren af en besked kan indikere aktiv lytning og udtrykke en tilknytning til afsenderen og samtalen ved at bruge emoji'er som minimalrespons (Kelly & Watts 2015); emoji'brug faciliterer "a better calibration and expression of your emotions (...) and leads to greater emotional resonance in the recipient" (Evans 2017: 35), og emoji'er synes også adækvate i situationer hvor følelse og affekt opleves som noget der er svært at sætte ord på, fx kærlighed og sorg (Stark & Crawford 2015, Stage & Hougaard 2018).

Ifølge media richness-teorien er det skriftlige medium forbundet med et mere begrænset udtrykkspotentiale end talen fordi skriften normalt ikke gengiver visuelle sociale tegn som gestik og mimik (Daft & Lengel 1984). Skriftsproget har gennem tiden oftest været brugt til den type kommunikation der er asynkron, monologisk, planlagt og varig. En af skriftens største fordele i forhold til talen er at den er et middel til lagring og overførelse af information på tværs af tid og rum, og på mange måder kan man betragte skrift som mere end blot talegengivelse. Skrift er ikke bundet til situationen som tale er, men med internetets udbredelse fremkommer en ny type skriftlig kommunikation der ganske vist hverken kan beskrives som 'samtaler på skrift' eller 'skrevet talesprog', men hvis frembringelsesvilkår minder om dem for talt interaktion og gør at interaktionen fremstår spontan, situationsafhængig

og dialogisk med skiftende ture. Der er dog visse væsentlige forskelle: Deltagerne kan ikke interagere fuldstændig synkront, og de kan ikke se og høre hinanden. Særligt det sidste vilkår var medvirkende til den kompensationsstrategi der voksede frem med skrevet onlineinteraktion (Hougaard 2004, Androutsopoulos 2011). Ud over stemmebrug, udtale og tonefald fungerer også kroppens bevægelser, dvs. gestik, mimik, blikretning, kropsstilling og -bevægelse som nøgleelementer i ansigt til ansigt-interaktioner (Junge 2019, Due 2016). Sammen med stemmebrug, udtale og tonefald kaldes disse kroppens kommunikative bevægelser for parasproglige virkemidler (Martin & Zappavigna 2019) og betragtes som emotionelle, situative og interpersonelle manifestationer af noget der kan være svært at udtrykke i den kropsløse skrift (Salló 2011, Kendon 2004). Det er oplagt også at se emoji-brugen – og særligt de emojis der har med kroppen at gøre – som en del af denne kompensationsstrategi, dvs. som noget der beriger den ellers 'tynde' eller 'fattige' kommunikationsform. Min undersøgelse af webchatdata indsamlet i 2000 viste et stort antal forskellige former for 'nyskriftlige' virkemidler hvis fremkomst i vid udstrækning kunne relateres til den interaktionssituation hvor deltagerne skrev med hinanden kvasisynkront. Webchat var i høj grad præget af relationskommunikation frem for indholdskommunikation (Watzlawick, Bavelas & Jackson 1967), og flere af virkemidlerne kunne med rette kaldes fatisk kommunikation (Malinowski 1923) fordi de "reflekterede det menneskelige behov for at etablere sociale sammenhænge, for gensidig genkendelse og anerkendelse" (Hougaard 2004: 35).

Emojis er kommunikative tegn fordi de signalerer eller repræsenterer en mening, og det er min centrale påstand i denne artikel at størstedelen af emoji-brugen i skrevet onlineinteraktion er kropsligt funderet og repræsenterer en kropslig betydningsgivelse pga. et samtalemæssigt behov for at lade kroppen komme til orde. Mange emojis (primært kropsemojis, dvs. dem der på en eller anden måde henviser til hele eller dele af kroppen) bruges som parasproglige virkemidler, dvs. som en erstatning for det manglende kropssprog, som virtuel mimik, gestik og stemmebrug. Den parasproglige information i skriftlig onlineinteraktion varetages nu i vid udstrækning af emojis. Fokus i denne artikel er dermed ikke så meget på hvordan teknologien ændrer sproget, men på

hvordan mennesket tilpasser teknologien eller teknologiske resurser til sine behov så den inkluderer den parasproglige del af kommunikationen, hvad jeg udfolder nærmere i analysen af de forskellige virkemidler, primært emojierne.

Efter at jeg har forklaret hvordan emoji'er virker som teknologi og system (afsnit 2) og givet et overblik over tidligere forskning i emoji'ers pragmatiske funktion (afsnit 3), vil jeg med udgangspunkt i data hentet fra webchat, Facebook og Instagram dels vise hvordan parasproglige virkemidler tidligere er blevet brugt i computermedieret kommunikation, dels udfolde artiklens centrale argument for at se emoji'er som kropstegn der peger tilbage på afsenderens krop og kropslighed og gør denne virksom i interaktionen (afsnit 4). De anvendte data er for størstedelens vedkommende hentet fra nogle af mine tidligere undersøgelser af sprogbrugen i skriftlige onlineinteraktion. Webchat-dataene stammer fra den interaktion der opstod i foråret og efteråret 2000 i to chatrum som det daværende TeleDanmark (TDC) faciliterede¹. Facebook-dataene kommer fra to danske facebookgrupper der etableredes i 2015 og fulgte og støttede to børn med kræft og deres familier i sygdomsperioden og efter dødsfaldet². Instagram-dataene er hentet fra opslag fra og kommentarer til DR's børneprogram's instagramprofil @dr_ultra i perioden 1.1.2020-23.1.2020.

2 EMOJIER ER MEDIESPROG

Når vi kommunikerer via digitale medier, bliver vores sprog medialisert, og der opstår medialektter, dvs. sproglige træk der er kendetegnende for og særligt udbredt blandt brugere af digitale dialogmedier (Hjarvard 2007). Ifølge Hjarvard giver ”mediernes tekniske krav, sociale organisering og de forskellige genrers æstetiske normer” (2007: 29) særlige betingelser for sprogbrugen som kan ses på alle sproglige niveauer og kommer til udtryk som fx forkortning, lydstavning og brugen af hashtags (Hougaard & Lund 2019, Berthelsen, Hougaard & Christensen 2019). At supplere sin kommunikation med emoji'er er dog at anvende et medialektalt træk der adskiller sig fra andre medialektale

1 Data er yderligere beskrevet i Hougaard (2004).

2 Data er yderligere beskrevet i Stage og Hougaard (2018).

træk fordi emoji³ er små kolorerede billeder der optræder på linje med den øvrige tekst.

Siden tierne er emoji³ hyppigt blevet brugt sammen med sprogtegn i computermedieret kommunikation (CMC) i elektroniske beskeder og på hjemmesider. Emoji³ fremstod oprindeligt i en opløsning på 12 x 12 pixel, men de nuværende emoji³ er meget mere detaljerede. Emoji³ er sociale og semiotiske mikroteknologiske artefakter hvis hyppighed og anvendelighed⁴ kan forklares med Herring og Dainas ord ”they hit a sweet spot – they are neither too large nor too detailed” (Herring & Dainas 2017: 8). Som meddelelssystem fungerer emojisystemet kun rudimentært og med en del begrænsninger, fx mangler kohæensionsmekanismer som ord- og sætningskobling (Cohn et al. 2019). Processer og tidsforløb som fx *i går/i morgen/nu* og begreber som *udvikling* og *forandring* er svære at formulere vha. emoji³. Emoji³ er præfabrikerede og rummer dermed ikke samme fleksibilitet som de fleste sprogtegn der kan bøjes (deklination og konjugation). Man kan således hverken sætte en løbe-emoji i datid, bøje den i bestemthed (en løber/løberen) eller udtrykke et præcist genstands- eller ejerforhold vha. emoji³. Følgende spørgsmål ’?’ rummer dermed mange forståelsesmuligheder og kræver en kontekstualisering for at blive mindre flertydigt. I skrevet onlineinteraktion vil en sådan kontekstualisering oftest udfoldes i de foregående eller efterfølgende ture.

Teknologisk kan emoji³ forstås som en font hvor hvert tegn har sin egen kode eller kodesekvens der er fastlagt af Unicodekonsortiet der dermed sikrer at emoji³ kan vises via teknologiske redskaber på tværs af tusinder af sprog. Denne standardisering ligger til grund for forskellige televirksomheder som fx Apples og Samsungs endelige design af hver enkelt emoji³. Emoji³ er i udpræget grad kulturspecifikke. Der er således en overvægt af emoji³ relateret til amerikansk og japansk kultur

3 Artiklen behandler de Unicodeaccepterede emoji³ der indgår i tekster på lige fod med andre tegn, dvs. ikke memojier, animojier, bitmojier og Facebook-avatarer der er udviklet til specifikke kommunikationssituationer eller medier. Animojier er et Apple-produkt der giver mulighed for animerede dyrehoveder hvis udtryk og mimik skabes ved at bruge telefonens kamera. Memojier er personaliserede animojier der skabes ved at animere ens eget hoved i stedet for dyrehoveder. Bitmojier og Facebook-avatarer er en slags personaliserede tegneseriefigurer eller klistermærker.

4 Selvom nyere forskning, fx Hansen & Stæhr (2021), konkluderer at unge ikke bruger emoji³ så meget som før, er det stadig en udbredt kommunikationsform.

fordi de to lande har bidraget mest til udviklingen. Emojiansigterne er ifølge Moschini et kulturelt blend af munden på de amerikanske gule smileyer og øjnene på de japanske kaomojier – ”the locus of facial expressivity” (Moschini 2016: 16). I modsætning til fx bitmojier og Facebook-avatarer kan emojierne ikke ændres og individualiseres så de kommer til at ligne en bestemt person. Denne homogenitet præger sammen med den ensidige kulturelle repræsentation vores emojikommunikation og begrænser vores udfoldelser. Siden 2015 har man dog kunnet bringe repræsentationen af talerens krop tættere på det virkelighedstro ved at kombinere kropsemojier med fem forskellige hudfarver som alternativ til den gule grundindstilling. Dette gøres ved at supplere koden for den valgte kropsemoji med en kode for hudfarve. Som det ses nedenfor, er ’hvid’ defaultfarven i den anvendte Fitzpatrick-skala, hvad der af mange er blevet opfattet som et racialiseret tiltag med en indirekte bekræftelse af ”white supremacy” (Amin 2022: 16).

EKSEMPEL 1: FITZPATRICKS FARVESKALA



Unicodekonsortiet udvider i de årlige opdateringer systemet med nye emojier, og det er muligt at påvirke den proces ved at udforme et detaljeret ”emoji proposal”; eksempelvis det afviste forslag om en emoji med x’er som øjne der bl.a. kunne vise at man er ”slået ud”⁵. I modsætning til en alfabetisk skrift der bygger på det fonetiske princip at hvert tegn svarer til en eller flere lyde der dermed kan være med til at danne et utal af ord, er emoji-systemet meget uøkonomisk. Selv om Unicodekonsortiet på nuværende tidspunkt har godkendt knap 4000 emojier, er det stadig forsvindende lidt i forhold til antallet af ord i fx den danske Retskrivningsordbog⁶. I forhold til et ideografisk system som emoji-systemet rummer et fonetisk sprogsystem næsten uanede muligheder for sammensætninger og nydannelser og er dermed både meget dynamisk og tilpasningsdygtigt som kommunikationsmiddel, en slags ’open-source’-projekt.

5 Se hele forslaget her: <http://www.unicode.org/L2/L2019/19303-face-with-x-eyes.pdf>.

6 I 2024 var antallet af ord i Retskrivningsordbogen 64.000.

Længere tekster bestående udelukkende af emoji'er forekommer oftest som kunstprodukter⁷. Under coronapandemien kunne følgende budskab dog ses på det sociale medie der dengang hed Twitter:

EKSEMPEL 2: IDEOGRAM-KOMMUNIKATION



Daværende direktør for Sundhedsstyrelsen Søren Brostrøm kommunikerede ofte udelukkende med emoji'er brugt som ideogrammer. Eftersom emoji'er opererer med minimal grammatisk kompleksitet, dvs. mere associativt, og hverken læseretning eller emoji'ers grammatik er universel og entydig (Cohn et al. 2019), har han i det valgte eksempel indsat pile for at tydeliggøre læseretningen og ikke mindst den handlingstilskyndelse der ligger i at se en årsagssammenhæng mellem følelsen af utilpashed eller feber (emoji'er med termometer, forbindelse og overophedning) og det at ringe (telefonrørsemoji) til nogen der kan afklare vha. en test eller undersøgelse (reagensglasemoji) om der er symptomer på covid-19, og om indlæggelse (hospitalemoji) er nødvendig. En sådan emoji-kommunikation rummer dog en del underforstået information og kan derfor kritiseres for at være upræcis og indebære en stor risiko for misforståelser.

3 FORSKNING I EMOJIER

Uagtet emoji'ers relativt nylige fremkomst og udbredelse er der allerede mange forskningsvinkler på fænomenet, fx marketing, psykologi, uddannelse, sundhedskommunikation, jf. review af Bai et al. (2019). I

⁷ Forsøg på kunstnerisk emojiudfoldelse er fx Emoji Dick, Pinocchio og Alice i eventyrland skrevet med emoji'er.

det følgende fremhæver jeg kun relevant kommunikations- og sprogforskning.

3.1 *Emojiers pragmatik og semantik*

Emojier kan forekomme alle steder i digitale tekster der understøtter emoji-formatet⁸, men der er en tendens til at de primært optræder i slutningen af ytringer – både som markør af sætningsgrænser og som afslutning af hele beskeder (Cramer, Juan & Tetreault 2016). I en undersøgelse af 6.461 indlæg på mediet WhatsApp viste Androutsopoulos (2018) at emojier (især ansigtsemojier) er den hyppigste måde at afslutte et indlæg på idet knap en fjerdedel af alle undersøgte indlæg sluttede med en emoji mens kun knap 10 pct. endte med et punktum. Den indlægsfinale position er også den mest almindelige placering for emotikoner (Provine, Spencer & Mandell 2007, Skovholt, Grønning & Kankaanranta 2014). Emojier betragtes dog ikke som neutrale grammatiske markører som fx punktum, men snarere som ”an emotionally charged form of punctuation” (Dürscheid & Siever 2017) der giver betydning til den foregående sætning i stil med udråbstegn og spørgsmålstegn.

Mange emojier udfører en slags ”affective labor” (Stark & Crawford 2015: 4) og forbindes da også ofte med følelser og humør (Hougaard & Rathje 2018, Kelly & Watts 2015), heraf navnet humørikoner. En del af forskningen har afprøvet muligheden for at måle graden af positiv/negativ stemning i tekster (*sentiment analyses*) vha. emojier (An et al. 2018, Quinta et al. 2023, Barbieri et al. 2016, Novak et al. 2015), for emojier formidler affektiv information – ”not enough to alter the valence of the message itself, but enough to alter the intensity of the affect” (Riordan 2017: 355).

Adskillige undersøgelser kommer dog frem til at emojiers overordnede funktion er at anskueliggøre og understrege forståelsen af intentionen i en besked i forhold til den situationsmæssige og interpersonelle betydning (Cramer, Juan & Tetreault 2016, Evans 2017, Herring & Ge-Stadnyk 2023, Kelly & Watts 2015, Pavalanathan & Eisenstein 2016, Provine, Spencer & Mandell 2007, Riordan 2017). Susan Herring har ad flere omgange været i gang med systematisk at kortlægge og fastlægge emojiers diskursive funk-

8 Et eksempel der ikke indeholder de formateringsfunktioner som tillader brug af emojier, er den krypterede asynkrone tovejskommunikationsform ”e-konsultation” som giver patienter mulighed for at kontakte lægen på skrift.

tioner (Herring & Ge-Stadnyk 2023, Herring & Dainas 2020, Herring & Dainas 2017) og er sammen med Dainas kommet frem til at den mest udbredte og basale grund til at bruge emoji'er er ønsket om at indikere "tone modification", dvs. "the use of an emoji to attribute a manner, attitude, or emotion to the text it accompanies" (Dainas & Herring 2021: 9), herunder at modificere og sikre sig mod eventuelle misforståelser af teksten. Brugt som modalitetsmarkører på lige fod med syntaks, ordtryk, tegnsætning og verbers modi fungerer emoji'er som en form for redskab til at forstå en teksts illokutionære, dvs. intenderede, kommunikative funktion (Austin, Urmson & Shisà 1982) og regnes som "illocutionary force indicating devices" (Herring & Ge-Stadnyk 2023). Dainas og Herring går endda så langt som til at påstå at der er sket en form for konventionalisering således at "popular face-representing emoji add tone by default" (Dainas & Herring 2021: 135) og gør at læseren nærmest kan 'høre' afsenderens stemme som fx smilende eller drilsk ud fra den valgte efterstillede emoji. I Hougaard og Rathjes undersøgelse fra 2018 rapporterede informanterne ligeledes at "help the understanding" var den vigtigste eller næstvigtigste årsag til at bruge emoji'er (2018: 779). "Det er så folk forstår i hvilken tone beskeden skal forstås", som en af informanterne forklarede (Hougaard & Rathje 2018: 780).

Afkodningen af den enkelte emoji kan dog kompliceres af semantisk diversitet og tvetydighed (Stark & Crawford 2015) – også på tværs af forskellige platforme, medier og kulturelle kontekster (Miller et al. 2016, Herring & Dainas 2017, Danesi 2017, An et al. 2018, Rodrigues et al. 2017), ligesom også brugernes alder påvirker afkodningen af de enkelte emojis (Herring & Dainas 2020). Ifølge Jaeger & Ares (2017) er særligt de gule ansigtsemojier udsat for en mangfoldig og kompleks aflæsning, men nogle emoji'er er mere tvetydige end andre. Ved at undersøge de 25 mest populære emoji'er fandt Miller et al. (2017) at 95,5 % af disse antropomorfe emoji'er blev forstået meget forskelligt, fx "Grinning Face with Smiling Eyes" hvis forskellige tolkninger (fra 'kampberedt' til 'ovenud lykkelig') særligt afhang af den store grafiske variation i afbildningen af emojiens mund. Den semantiske diversitet giver selvsagt en høj risiko for misforståelser.

3.2 Emoji'er som minimalrespons og social bekræftelse

En af de måder man i fysiske samtaler kan være aktivt lyttende på, er ved at afgive forskellige former for minimalrespons enten i form af bevægelser, fx

nik med hovedet, smil eller løft af øjenbryn, eller ved at sige lyde (*mm*, *mhm*) eller små ord (*ja*, *okay*) (Schegloff 1982) der indikerer at man som lyttende følger med i det der bliver sagt uden at tage ordet (Nielsen & Nielsen 2010). Den type umiddelbar og minimal feedback er i reglen ikke mulig i online skriftlig interaktion fordi denne oftest er kvasisynkron, forstået således at en sendt besked næsten altid er tilgængelig hos modtageren øjeblikkeligt, men uden at modtageren kan følge med i produktionen af beskeden undervejs som man kan i fysiske samtaler (Garcia & Jacobs 1999). I Messenger er det muligt at reagere på andres beskeder ved at 'klistre' en emoji direkte på deres tekst, hvad der minder om en ikketurtagende minimalrespons, jf. de første to ture i eksempel 3, hvor de to deltagere på den måde kvitterer for den andens besked uden at tage turen. En reaktionsemoji der står separat, hvad der ses længere ned i eksempel 3, repræsenterer en selvstændig tur. Begge typer kan have en interaktionel betydning i stil med bekræftere (Nielsen & Nielsen 2010: 102) der viser tilslutning til det kommunikerede, men også social forbundethed. I eksempel 3 viser K. "alignment" (Stivers 2008) ved med det præfererede "Ja!" at svare bekræftende på en anmodning, og ved at indsætte en emoji som selvstændig tur udviser K. tillige "affiliation" (Stivers 2008), dvs. "affektivt samarbejde" (Jørgensen 2018: 29) i form af en stærk tilslutning til den andens (undertegnede) forehavende og indstilling til emnet.

EKSEMPEL 3: EMOJI SOM MINIMALRESPONS



Ifølge Stark & Crawford hjælper emoji'er som en social teknik "people in digital environments [to] cope emotionally with the experience of building and maintaining social ties" (2015: 8); "the social button" (Gerlitz & Helmond 2013) rummer en relationel metabesked om gensidig interesse i at forstå og blive forstået.

4 DEN DIGITALE KROPS ANATOMI

I nedenstående afsnit 4.1 uddybes hvordan deltagerne med fremkomsten af forskellige typer skrevet onlineinteraktion, fx chat og sms, begyndte at digitalisere og virtualisere den parasproglige del af kommunikationen, både i form af 'nyskriftlige' resurser og billedlige elementer som fx emoji'er. Det er min påstand at den digitale krops anatomi har mange fælles træk med den fysiske krops anatomi, og at den i interaktionen spiller en rolle der minder om den fysiske krops, og jeg plæderer således i afsnit 4.2 for at se emoji'erne med Peirces begreber *ikon*, *symbol* og *indeks*. Dette uddybes med eksempler fra materialet i afsnit 4.3 om gestik og 4.4 om mimik. Brugen af emoji'er hvis fysiognomi er urealistisk overdrevet, fx med røde og blå ansigtsfarver, og en udbredt metaforisk brug af visse emoji'er kalder dog sandsynligvis på en udvidelse af forståelsesrammen, hvad jeg kort vil berøre i den efterfølgende diskussion i afsnit 5.

4.1 Digitalt parasprog

Helt tidligt blev nogle af de 'nyskriftlige' fænomener beskrevet som en kompensationsstrategi for tilføjelse af socioemotionelle tegn (Sproull & Kiesler 1986) der som kontekstualiseringsmarkører (Gumperz 1982) giver information om hvordan det skrevne skal forstås. Ud over grafiske elementer som fx skriftfarve og -type (fx Webding) bestod de nyskriftlige virkemidler af grafiske effekter som fx lydstavning, ordforkortning og emfaseangivelse (fx brug af store bogstaver) og mere situations- og handlingsorienterede elementer som fx emotes, der er kommandobaserede handlinger der "allows a user to express actions or internal states of mind" (Cherny 1999: 19), og regibemærkninger om udseende, gestik, mimik (Hougaard 2004, Crystal 2006). Udbredelse af særligt de situations- og handlingsorienterede grafiske virkemidler blev dels tilskrevet deltagernes behov for i tillempet form at tilføre den skriftlige interak-

tionssituation nogle af de parasproglige elementer kendt fra samtalsituationen (Giles, Stommel & Paulus 2017, Meredith & Potter 2013), dels deltagernes ønske om at betone nærværets intimitet og spontanitet (Hougaard 2004, Kelly & Watts 2015).

I den tidlige CMC-forskning blev kroppen fremhævet som fraværende eller maskeret. "[T]he typed text provided the mask" (Danet 1998: 129). Denne anonymitet tillod eksperimenter med identitet og køn (Turkle 1997), men også diskussioner af kropslige simulerede oplevelser (Bolter 1991). Med udgangspunkt i den franske forståelse af 'virtuel' som noget "som har Evne til at virke"⁹, blev kroppen "virtualiseret i teksten, tempoet og nærværet" (Hougaard 2004: 44). "The virtualization of the body is therefore not a form of disembodiment but a recreation, a reincarnation, a multiplication, vectorization, and heterogenesis of the human" (Lévy 1998: 43), og det kom i chat bl.a. til udtryk ved en udbredt brug af emotes eller virtuelle performativer¹⁰, dvs. performative sproghandlinger udtrykt i præsens indikativ og omkranset af asterisker. I skrevet onlineinteraktion tæller den virtuelle krops tegn ikke blot som noget beskrevet eller omtalt, men som noget der virker, dvs. frembringer en tilsigtet virkning¹¹. Forståelsen af disse kropstegn som virtuelle performativer peger på kroppens potentiale og virkekraft i interaktionen. "Den virtuelle interaktion beskæftiger sig ikke med hvad der er virkeligt eller uvirkeligt, sandhed eller løgn, men med hvad der er virksomt og appellerende" (Hougaard 2004: 45). Forestillingen om de virtuelle kropstegn som virksomme og ikke kun beskrivende bibringer afsenderens digitale tilstedeværelse og handling en højere grad af gyl-dighed og troværdighed.

Inden jeg i de følgende afsnit behandler kropssproget, vil jeg her berøre den del af vores parasproglige resurser der har med stemmen at gøre, dvs. stemmekvalitet, stemmestyrke etc. Stemmen har betydning for vores relationskommunikation og følelsestillkendegivelse (Jensen

9 Hentet fra Meyers Fremmedordbog fra 1970 (1837): <https://meyersfremmedordbog.dk/ordbog?query=virtuel>.

10 Jeg forstår performativer som sproghandlinger med sociale konsekvenser, dvs. hvis mening skabes i og med at de ytres.

11 Sherry Turkle (1995) berører kort de virtuelle performativers magt i sin beretning om personen der ikke kunne få en et uønsket virtuelt knus gjort betydningsløst og ugjort.

2009, Steensig 2001), men stemmerelaterede og prosodiske elementer er svære at overføre detaljeret til skriftsproget når man ikke bruger regulær lydskrift. Det er dog efterhånden alment anerkendt at brugen af versaler angiver emfatisk stemmestyrke, mens tegniteration bruges som tegn på lydforlængelse fx 'hvaaaad', og ordforkortning kan være en måde at markere afsnubning af trykssvage endelser på, fx 'ka du ik' (Hougaard 2014, Werry 1996). Desuden kan brugen af lydstavning ud over ordenes udtale indikere talestrømmens rytme, dvs. prosodiske elementer som det ses i følgende eksempel¹², hvor den mandlige chatter Dwølf vha. en virtuel performativ kilder S@ndr@, der reagerer meget talesprogsnært ved at lydstave to interjektioner 'ih' og 'nej'. Interjektionen 'ih' fremstår vokalforlænget pga. de fire efterstillede *h*'er, og interjektionen 'nej' tilføres en ekstra høj vokallydskvalitet fordi *j*'et erstattes iterativt af den høje vokal *i*. Svaret kommer dermed til at konnotere noget stereotyp feminint flirtende.

EKSEMPEL 4: LYDSTAVNING

Dwølf: *kilder sandra*

S@ndr@: *ihhhh neiii Dwølf*GGG*

Et andet sted i mine webchatdata reagerer deltageren SannePigen på nogle overvældende lykønskninger ved at hendes virtuelle krop afgiver tilsyneladende ufrivillige kropstegn som en performativ handling idet hun svarer følgende: "Takker *rødmer af alt den opmærksomhed*" (Hougaard 2004: 210). Generelt er disse webchatdata præget af et stort antal virtuelle performativer lige fra *vinker*, *tøsekrammer* og *hOp-PeR iVrIgT* til de mere specifikke og fiktioniserede som *hiver kitlen på* og *griber blidt*, der alle som virtuelle handlinger ansås som virksomme for deltagerens oplevelse af hinanden. Med Scott Sørensen og Walthers' ord foregik der en bevægelse fra referentiel til operationel kommunikation (2001). Mest udbredt var dog forkortelser for kropslige handlinger og reaktioner såsom *G* for 'griner' og *S* for 'smiler', der oftest optrådte som en form for "post-completion stance markers" (Schegloff 1996: 90). I følgende eksempel anvender afsenderen et *G*

¹² Eksemplerne er hentet fra Hougaard (2004).

for at markere et ironisk grin: ”-ejjiiii, jyder er gode nok...-se bare på Kim fra robinson...*G*”. Disse handlingsforkortelser, også kaldet netkronymer (Hougaard 2004), forsvandt dog med emojiernes entré.

Vi ved ikke præcist hvorfor emoji'er så hurtigt blev allemandsseje, men en hypotese kunne være at den forholdsvis gnidningsløse inkorporering af billedtegnene i vores online interaktion lod sig gøre fordi emoji'erne løste et essentielt problem (fx at visualisere følelser og kompensere for det manglende kropssprog) på en umiddelbart indlysende og anvendelig måde. Mens det kræver et vist kodekendskab at dechiffrere forløbere som emotikoner fx :'(der betyder 'græder', og netkronymer fx *V* for 'vinker' og LOL for 'Laughing Out Loud', er forståelsen og anvendelsen af emoji'erne mere intuitiv fordi deres form ligner det der repræsenteres, så meget. Det kan naturligvis diskuteres om der reelt er tale om en utilstrækkelighed i sproget, og et helt naturligt modargument vil være at skriftsproget har klaret sig glimrende uden emoji'er i mere end tusind år. Det er dog min grundantagelse at en forklaring dels skal hentes i de forhold den skriftlige interaktion produceres under, dels i en peirciansk forståelse af brugen af tegn.

4.2 En peirciansk forståelse af emoji'er

Godt 500 af de omkring 4000 emojis i Unicodekonsortiets seneste opdatering er relateret til den menneskelige krop. Mange emoji'er viser mennesker i bestemte situationer hvor de fx dyrker forskellige sportsgrene, bliver masseret eller befinder sig i forskellige familiekonstellationer eller roller, fx politibetjent eller kok. Hvad enten det drejer sig om emotikoner eller emoji'er, er de kendetegnet af fysisk lighed mellem det betegnede og det betegnende. Selv om emoji'en 🤪 er nemmere at afkode som julemand end emotikonet *<|:-) med samme betydning, er princippet bag begge typer tegn stadig at det skal ligne. Det er derfor min påstand at en 'Peirce-inspireret' forståelsesramme for de kropsligt funderede emoji'er kan hjælpe med at præcisere synspunktet om sammenhængen mellem krop og emoji. Ifølge Peirce har et tegn ”den grundstruktur at noget står for noget for nogen” (Togeby 2019: 21). Tegn refererer til noget andet på tre måder: via lighed (ikonisk), via fælles association og konvention (symbolsk) og via påvirkning (indeksikalsk); standardeksemplet på den indeksikalske reference er at røg in-

dekserer ild, eller at fodspor indekserer at her har nogen gået. Det er indlysende at emoji'er refererer til noget fordi de ligner (ikonisk) dette; en smilende emoji 😊 betyder det samme som et smilende ansigt, hvad der normalt vil forstås som en venlig og imødekommende mimik. Men en del af emoji'erne er snarere at regne som symboler hvis betydning vi i kraft af konventioner er enige om således at vi får samme mere eller mindre arbitrære association når vi ser bestemte symbolske tegn fx ❤️, der står for romantisk kærlighed eller positive følelser. Symbolers betydning er ofte almene og kulturelt specifikke og kan derfor forhandles. Emoji'er kan dog også referere til det som de betegner fordi de befinder sig i en årsags- eller nærhedsrelation til det betegnede (indeksikalsk), jf. Peirces definition: "Et *index* er et tegn, som refererer til det Objekt, som det betegner, i kraft af at det virkelig er påvirket af det Objekt." (Peirce 1995: 100). Det indeksikalske uddybes nedenfor.

Når det kommer specifikt til de kropslige emoji'er, er den ikoniske reference mest tydelig idet mange af emoji'erne ligner en nemt genkendelig kropsspositur eller -bevægelse eller mimik. En hånd med to spredte fingre viser V-tegnet (victory hand) ✌️, men samlet tommel- og pegefinger indikerer noget småt 🤏, og mange af de gule smiley'ers ansigts-træk kan nemt afkodes selv om den grafiske afbildning er let fortegnet og karikeret i sin forenkling, fx 🙄 der via øjnene, der er større end normalt for emoji'er, viser et ansigt hvis lille mund og sænkede øjenbryn udtrykker bekymring, bønfoldelse el. lign. En del af de kropslige emoji'er indeholder også symbolske elementer når øjnene erstattes af hjerter, stjerner eller dollartegn: 🥰, 🥳 og 🥳, munden af en lynlås: 🗨️, og næsen af en pinocchio-lang næse: 🤪 Dette giver smiley'erne symbolske betydningsnuancer i retning af henholdsvis kærlighed, begejstring, penge, tavshed og løgn fordi disse tegn i mange kulturer bruges med denne konventionaliserede betydning. Det kan naturligvis diskuteres hvor grænsen mellem ikonisk og symbolsk går. Til eksempel kan det at få en lang næse på dansk også betyde at blive snydt uden at nogen nødvendigvis har løjet for en. Dette idiom kan også antydes med en håndbevægelse gående fra næsen og nedad, så hvis det er den håndbevægelse der forsøges mimet, kan man argumentere for at det er et ikonisk tegn.

Endelig er det efter min mening oplagt også at undersøge emoji-brugen som en indeksikalsk handling fordi brugen af kropsemoji'er kan ses

som et udtryk for en tilskyndelse til at lade kroppen komme til orde. Den reference der er mellem en kropsemoji og det som emoji'en refererer til, virker i kraft af kausalitet eller nærhed. Det at bruge en opkastende eller kvalme emoji, 🤮 eller 🤢, indekserer som udgangspunkt det at føle kvalme uagtet om denne kvalme stammer fra sygdom eller væmmelse. Det at føle kvalme er en kropslig reaktion der ikke umiddelbart lader sig styre, og derfor er der ikke tale om en intentionel handling. At indsætte en emoji i en skrevet tekst er naturligvis en intentionel handling da tekster ikke skriver sig selv, men det at den kropslige reaktion mimes, kan ses som et forsøg på at bringe teksten tættere på den kropslige virkelighed.

Selv om tabel 1 nedenfor opdeler tegnbruken i de tre overordnede kategorier, er det efter min mening mere hensigtsmæssigt at tale om aspekter af betydning snarere end forskellige kategorier af betydning. Den opkastende emoji vil fx også kunne forstås gennem lighed idet en person der udtrykker afsky, kan ligne en der har kvalme eller er ved at kaste op. På lignende vis kan emoji'en der bøjer armen 🦵, både referere til en stærk arm (ikonisk tegnbrug) og forstås som indeksikalsk tegnbrug fordi der i de fleste kulturer er en kausal eller erfaringsbaseret sammenhæng mellem styrke og respekt. Det indeksikalske aspekt fun-derer sproget i kroppen – som en form for ekko af kropslige oplevelser. Der skabes forbindelse mellem den fysiske krop og den medierede og medialiserede krop, forstået således at det at være ”grøn i hovedet” in-dekserer kvalme, og ”grøn i hovedet”-emoji'en bruges ikonisk for det at se ud som om man har kvalme.

TABEL 1: PEIRCE-INSPIRERET OVERSIGT OVER EMOJIBRUG

	Forklaring	Mine eksempler og forklaringer
Ikonisk	Ligner kropsbevægelser eller ansigtsudtryk	👍 (jeg anerkender) 😁 (græder af grin) 😊 (lille smule / præcision)
Symbolsk	Konventionaliseret bestemmelse	💀 (jeg dør af grin) 🦵 (jeg anerkender, stærkt) 👎 (nej/ikke godt) 🤦 (rundt på gulvet)
Indeksikalsk	Årsags- eller nærhedsrelation	👉 (henlede opmærksomhed) 🤮🤢 (symptom på foragt)

Særligt når det kommer til kropstegnene, vil jeg fremhæve indeksikaliteten, hvor det ikke så meget handler om tegnet (emojien) i sig selv, men i højere grad om det som emojien peger tilbage på mhp. at nærme sig en oplevelse eller følelse af noget fysisk. Med den forståelse kommer de kommunikerendes kroppe i kraft af oversættelsen til emojis tilstede i interaktionen gennem afsenderens pegen (gennem indeks) på deres vigtighed og virksomhed (jf. betydningen af virtuel). Den kropslige kommunikation der er essentiel i mundtlige samtaler, føres over i skriften som en augmenteret virtuel kropslighed. I følgende eksempel ”Spørg en anden 🙋” er den emoji der trækker på skuldrene, brugt som forklaring på hvorfor afsenderen sender den spørgende videre til en anden person. Emojien henviser til handlingen at trække på skuldrene hvad der både kan understrege og træde i stedet for udsagnet ”for jeg ved det ikke”. Emojien skaber en forklarende forbindelse mellem det verbale udsagn og kroppens udsagn. Emojierne kan ses som et forsøg på at overskride oplevelsen af mediering ved at øge en interpersonel transparens og gøre kroppen og handlingerne virksomme (jf. forståelsen af virtuel) og reducere oplevelsen af mediering (Bolter & Grusin 1999). På den måde bliver emojis en troværdighedshandling hvor de højner oplevelsen af u-middelbarhed og nærvær og dermed afsenderens autenticitet.

4.3 Digital gestik

Mange forskere har berørt det parasproglige i forbindelse med forståelsen af emojis funktioner (Novak et al. 2015, Na’aman, Provenza & Montoya 2017, Herring & Ge-Stadnyk 2023), men kun få går som Michael Schandorf og Lauren Gawne og Gretchen McCulloch i dybden med hvordan emojis og andre parasproglige resurser hjælper os med at skabe det betydningslag i interaktionen der formidles vha. gestik. Efter en kort gennemgang af de to forskningsbidrag sammenligner jeg deres tilgange med eksempler fra mit materiale og når frem til at det giver mening at operere med i alt seks kategorier for emoji-gestik.

I sin artikel om ”Mediated gesture” (2012) behandler Schandorf ikke emojis da de endnu ikke var udbredt på det tidspunkt, men han når til gengæld omkring adskillige andre typer tegn i sit for-

søg på at redegøre for disse som parasproglig kommunikation på Twitter. Schandorf betragter gestik som "the emotional grounding of social interaction" (Schandorf 2012: 335) og argumenterer for at se analogier mellem de medierede praksisser og forskellige fysiske gestik, og han påstår at både fysisk og medieret gestik "arise from the same embodied cognitive roots and needs to serve equivalent social and communicative functions" (Schandorf 2012: 321). Med udgangspunkt i Kendon der ser gestik som "as much a part of language as speaking is" (Kendon 2017: 168), fokuserer Schandorf på "co-speech gesture" (Schandorf 2012: 324) og arbejder desuden ud fra McNeills definition af gestik som "a multiplicity of communicative movements, primarily but not always of the hands and arms" (McNeill 2006: 58). Schandorf analyserer fx tegnet @ som en pegende gestus der henviser læseren til andre deltageres brugernavne og omformer disse til links, og det at skrive med en kombination af versaler og lydskrift ser han som en form for "transcriptions of vocal gestures" (Schandorf 2012: 330). Han konkluderer at "the compressed, extensive, paralinguistic emotional connections of phatic gesture are embodied in new forms afforded by new ways of 'keeping in touch'" (Schandorf 2012: 338).

I deres artikel "Emoji as digital gesture" (2019) konstaterer Gawne og McCulloch at selv om Schandorf (2012) har trukket paralleller mellem gestik og forskellige typer af digital skrivning, herunder emotikoner, så forbliver "the precise nature of the parallels between emoji and gesture" underudforsket (Gawne & McCulloch 2019). Uden at påstå at alle emojis fungerer som intentionel gestik, ønsker Gawne og McCulloch at vise at denne emojibrug folder sig ud over nogle kategorier der til dels overlapper med Schandorfs, som det ses i tabel 2, hvor jeg har suppleret gennemgangen med egne eksempler:

TABEL 2: EKSEMPLIFICERET FORKLARING AF SCHANDORF OG GAWNE & MCCULLOCHS KATEGORIER AF MEDIERET GESTIK

	Schandorf	Gawne & McCulloch	Betydning og brug	Mine eksempler
1	Deictic	Pointing	Spatiotemporal indikation og opmærksomhed	👉 Follow @xxxx for more
2	Rhythmic	Beat	Strukturering og emfase	Jeg 🎵 elsker 🎵 det 🎵
3	Imagistic	Illustrative, metaphoric ¹³	Illustrerende eller metaforisk imitation	Stop 🙅 👍
4	Expressive	Illocutionary	Intention, tonefald	Spørg en anden 🗣️ Kom nu 🙌
5	Emblem		Kulturspecifikke og konventionaliserede symboler	vi krydser fingre for din drøm går i opfyldelse 🙌🙌 👍
6		Backchannelling	Feedback og minimal-respons	👍 ❤️

Ifølge Gawne og McCulloch er der udbredt enighed om de første tre kategorier: Deictic/Pointing, Rhythmic/Beat og Imagistic/Illustrative/Metaphoric. Disse typer gestik kan siges at handle om at styre modtagerens opmærksomhed på forskellige måder, nemlig ved at pege, strukturere det (rytmiske) ordflow eller illustrere handlinger, tanker og abstrakte ideer. Deiktiske emoji'er kan både udpege noget betydningsfuldt og pege i en bestemt retning hvad der ofte ses på Instagram, hvor den pegende hånd fungerer handlingsanvisende fx i forhold til at opfordre til at begynde at 'følge' nogen, fx '👉 Follow (@xxxx for more'. Med rytmisk gestik kan man slå takten i en passage og understrege det sagte eller særligt betydningsfulde ord, eller med Kendons ord "beat the tempo of his locomotion" (Kendon 2004: 333). Ifølge Schandorf kan denne tegnbrug "create pauses, suspense/expectation/anticipation" (Schandorf 2012: 324) og angive en form for læserytme. Emoji'erne i et indlæg som 'Jeg 🎵 elsker 🎵 det 🎵' markerer en rytme som læses (eller formodes at blive læst) med fornemmelsen af pause mellem hvert ord og fungerer dermed som en fremhævelse eller betoning af hvert ord i udsagnet

13 Jeg har tilladt mig at slå "Illustrative" sammen med "Metaphoric" fordi Gawne og McCulloch forklarer begge kategorier som illustrerende noget konkret eller abstrakt, hvad der næsten svarer til Schandorfs "Imagistic".

på en anden måde end et afsluttende udråbstegn ville have gjort. Den imiterende gestik påpeger en form for lighed mellem emoji og gestus. I eksemplet 'Stop 🙌' er signalet fra den løftede hånd ret klart, men i fysisk interaktion ville denne type gestik virke temmelig uhøflig og ansigtstruende, måske fordi vi primært forbinder håndbevægelsen med ordensmagten eller en kraftig advarsel. Derfor kan emoji'en i stedet ses som en slags oversættelse af den måde man forestiller sig "stop" udtalt på, nemlig med særligt eftertryk og understregning af vigtigheden. Selv om vi genkender stop-hånden og let afkoder den pga. ligheden, er der altså ikke et 1:1-forhold mellem brugen af emoji'en og håndbevægelsen. Et andet eksempel på en imitation som ofte ses når man gerne vil have nogen til at ringe eller bruge telefonen, er 🙌. På dr_ultras spørgsmål om hvordan man imponerer et crush, svarer en yngre følger på Instagram til sidst at de da bare kan ringe (i form af en emoji) hvis de vil vide mere:

EKSEMPEL 5: EMOJIGESTIK

[rahmasahan7](#)

@dr_ultra man kan vise det man er god til som er specialt

fx. At ku' karate. 🙌

[dr_ultra](#)

@rahmasahan7 det var en idé ja 🙌

[rahmasahan7](#)

@dr_ultra 🙌

De tre øvrige kategorier (nr. 4-6) behandles forskelligt af forskerne idet Schandorf fokuserer på det konventionaliserede (Emblem) og på afsenderens sindstilstand (Expressive): "Actions that express or indicate emotion or state of mind are expressive gestures" (Schandorf 2012: 329), mens Gawne og McCulloch fremhæver det interaktionelle (Backchanneling) og afsenderens hensigt: "the intention of the speaker in saying a particular utterance" (Gawne & McCulloch 2019). Fælles for forskerne er en insisteren på at de behandlede tegn optræder med samme kommunikative og sociale funktion som gestik i samtaler. De digitale tegn er afledt af de kropslige tegn og supplerer den verbale kommunikation på samme måde.

De forskellige måder at forstå emoji'er på varierer afhængigt af synsvinkel, dvs. om man fokuserer på afsenderen (ekspressiv funktion) eller modtageren (konativ funktion) (Jakobson 1960). Dette kan illustreres

ved et eksempel: 'come on 🖐️' som enten kan forstås som en hilsende håndbevægelse mellem venner (fist-bump) der fortæller modtageren at han godt kan regne med den andens opbakning, eller forstås som en afsender der illustrerer sin frustration (punch) over den andens forsinkelse eller sløvhed. Konteksten vil være helt afgørende for forståelsen af en sådan emoji.

Schandorfs emblematiske kategori arbejder med kulturelt etablerede betydninger som fx emoji'en 🙌 der henviser til den udbredte skik med at krydse fingre som et håb eller ønske om noget godt som dr_ultra gør i følgende eksempel:

[dr_ultra](#)

[@vildemollegarden](#) vi krydser fingre for din drøm går i opfyldelse



Emojier der bruges som en form for minimalrespons, er som tidligere nævnt en almindelig praksis i sociale medier, også kendt som *likes* og *reaktioner*, som er et værktøj til indtryksstyring (Gerlitz 2013, Ozanne et al. 2017) eller minimale tegn på forståelse og mikroanerkendelse. Da samtalsens små feedbacktegn ikke er tilgængelige i skriftlig onlineinteraktion, overføres deres funktionalitet i nogen grad til brugen af bestemte emojier fx 🙌👍. Ifølge teorierne om minimalrespons er denne form for respons vigtig for at opretholde interaktionen og relationen, og man kan derfor mene at "the urge to produce a minimal response could explain the growing role of likes and other reactions on social media" (Stage & Hougaard 2018: 33).

4.4 Digital mimik

I dette afsnit viser jeg hvordan ansigtsemojier bruges som digital mimik der understreger, supplerer eller erstatter det skrevne indhold. Ud af de godt 500 kropsrelaterede emojier refererer 60 til kropsdele (primært hænder), hvoraf de fleste udfører en genkendelig gestik, og omkring 100 ansigtsemojier mimer forskellige ansigtsudtryk¹⁴ og kan

¹⁴ Unicodekonsortiet inddeler på deres Full Emoji List, v15.1 (<https://unicode.org/emoji/charts/full-emoji-list.html>) ansigtsemojierne i følgende underkategorier: smiling, affection, tongue, hand, neutral-sceptical, sleepy, unwell, hat, glasses, concerned og negative.

dermed rubriceres under Ekman og Friesens kategori "affective displays" (Ekman & Friesen 1969). Enkelte af ansigtsemojierne mimer hovedbevægelser, fx det at nikke med hovedet, lige som der også findes 'overkropsemojier' der mimer typiske kropsbevægelser som Ekman og Friesen indregner under nonverbal adfærd, fx en person der rækker hånden op.

Når dr_ultra på Instagram stiller spørgsmålet '@dxxxx hvor var du til fest henne? 😊😏', viser den første smilende emoji afsenderens imødekomende mimik idet denne spørger 'med et smil', mens emoji nummer to nok snarere understreger temaet i den supplerede tekst, nemlig det festlige. Ofte understreger dr_ultra det spørgende vha. en emoji med et undrende ansigt som i følgende eksempel: '@bxxxx hvad ville du forandre? 🤔'. Eftersom der allerede står et spørgsmålstegn der viser den spørgende modus, henleder den tænkende emoji med de sænkede øjenbryn i stedet tanken på eftertænkksomhed og kan være med til at forstærke oprigtigheden i ønsket om at få et svar idet afsenderens valg af emoji understreger at dr_ultra ikke bare spørger af ren høflighed. Denne eftertænkssomme emoji kan dog også forstås som udtryk for skepsis eller forbehold – i dette tilfælde i forhold til hvorvidt den adspurgte faktisk kan eller vil forandre noget.

Ofte er emojiernes billedsprog overdrevet i forhold til det subtile kropssprog hvad der fx ses i den udbredte brug af den emoji der skal forestille det at rulle rundt på gulvet af grin 😂 eller den emoji der har tårerne bredt strømmende ned ad kinderne 😭. Når dr_ultra ønsker en forklaring på hvordan det føles af være forelsket, lyder et af svarene: '@dr_ultra Ved det ikke syntes bare det klamt og mærkeligt at blive forelsket 🙄'. Emojien der står på hovedet, er højest sandsynligt ikke et nøjagtigt udtryk for den fysiske virkelighed, men snarere en mere metaforisk beskrivelse af den kropslige følelse af at være rundt på gulvet eller blot en indikation af at det er et svært og måske lidt skørt spørgsmål at skulle besvare.

I et svar på et andet af dr_ultras spørgsmål om hvorvidt det er ok at drille en man er forelsket i, viser emojierne både noget tilstræbt mimisk og lydligt: '@dr_ultra Øhm nej 🙄 det synes jeg heller ikke er spor ok okay hvis der er nogle der driller 😏 mig på en ond måde også kan jeg finde på at blive rigtigt bange 😱 for det'. De tre emoji angiver vha. an-

sigtsudtryk at afsenderen dels føler sig så trist at vedkommende kunne græde, dels at afsenderen kniber øjnene sammen pga. drillerier og bliver bleg af frygt, men de sidste to emoji'er er placeret temmelig usædvanligt inde i sætningerne og foranlediger derfor at de foranstillede ord forestilles udtalt med henholdsvis en bekymret og en beklemt stemmekvalitet. Her giver emoji'erne en fornemmelse af den aktuelle kropsfølelse og kan være med til at fremkalde en intonationsfornemmelse.

I de undersøgte facebookgrupper spillede emoji'erne en stor rolle i måden de mange medlemmer af grupperne kommunikerede sorg og medfølelse på. 91 % af alle kommentarerne indeholdt en eller flere emoji'er, og hele 29 % af kommentarerne bestod udelukkende af emoji'er hvor hjertet var mest udbredt. Grædende emoji'er forekom både som den triste med en enkelt tåre i øjenkrogen og som en hulkende emoji hvor gråden løber ned ad emojiens ansigt:

Hvor er det dog uretfærdigt 😞 det gør mig så ondt. R.I.P lille engel 😊💜💜💜
💜💜💜💜💜

Åh sikke en forfærdelig nyhed 😞 rip Miv. 😞😞😞

Da faren fortalte den forfærdelige nyhed om barnets død i en af facebookgrupperne, reagerede mange af deltagerne primært med emoji'er som i følgende eksempler

What!? 💜💜💜💜💜💜💜💜
:-/ ??? 😞😞😞😞😞😞😞
😞😞😞😞💜💜

Den umiddelbarhed der kan ligge i at kommunikere alene med emoji'er, stemmer fint overens med mediets muligheder og situationens behov for nærvær og sensitivitet. Kroppens spontane følelsesudtryk opleves mere autentiske og ærlige og dermed mere passende end reflekterede tilkendegivelser der ofte kan være præget af abstraktionens og eftertankens distance. Den verbalsproglige semantiske præcision er desuden krævende både at ind- og afkode, mens det at ty til de ekspressive, men også flertydige emoji'er på en gang rummer øjeblikkets impuls og krop-pens intensitet.

5 DISKUSSION

Emojier er ofte blevet fremhævet som følelses- eller intentionsinformation, og det har i denne artikel ikke været mit mål at tage afstand fra denne forståelse af hvad man bruger emojier til i skrevet onlineinteraktion. Tværtimod er det gennemgående argument i artiklen at emojier netop egner sig så godt til at formidle følelser og til at guide modtageren i at forstå afsenderens hensigt med det skrevne, dvs. den illokutionære funktion, fordi brugen af emojier og forståelsen af dem er kropsligt funderet. Deltagerne i skrevet onlineinteraktion bruger emojier som kompensation for den parasproglige information fordi emojierne i vid udstrækning ligner og føles som kropslig kommunikation.

Lebduska fremhæver at "the production of writing is not a disembodied activity of pure cognitive processes but is instead a physical activity" (2014) og ser emojiers visuelle komponent som afgørende for at de åbner "a gateway to non-discursive language" (Lebduska 2014: 1). Faktisk imødekommer emojier ifølge Lebduska de 'krav' der blev fremsat tilbage i slutningen af 1980'erne af forskerne Murr og Williams:

Since language shapes our perception of reality it is clear that new concepts of language must be developed with the emergence of the visual culture. Just as the practice of language in Western civilization, as served by speech and memory, was altered by the invention of an alphabet and later its encoding in books through monotype, so must it be altered again as a result of the invention of visual tools. (Murr & Williams 1988: 415)

Det er nok en overdrivelse at se emojiernes forandringspotentiale som noget der minder om trykkekunstens, men deltagerne i skriftlige online interaktioner tilpasser tydeligvis de teknologiske resurser til deres behov så den parasproglige del af kommunikationen bliver integreret. Kroppens supplement til kommunikation får efter min mening særligt i emojierne en mere rammende analog referent end tidligere set i skriftlig interaktion. Sammenligningen af emojier med gestik, mimik og stemmeføring rummer dog nogle svagheder. Der er mange ting der adskiller emojier fra parasproglige tegn: Antallet af emojier er begrænset, de vælges fra et tastatur, deres gengivelse særligt af mimik er stærkt forenklet, og emo-

jierne bruges ikke tæt synkroniseret med talestrømmen som gestik gør i en samtale (McNeill 1992). Nogle gange bruges kropsemojier på måder som man ikke ser i fysiske samtaler. Selv om '🦵' eksempelvis beskrives som 'Flexed Biceps' på emojipedia.com og forklares som repræsenterende begrebet *styrke* (symbolsk) eller begrebet *træner* (ikonisk), anvendes den stærke overarmsemoji ofte som en vurdering i stil med udtrykket 'Godt klaret', dvs. som en opmuntrende respons der minder om andre 'sociale knapper' brugt i sociale medier, dog ser man meget sjældent nogen vise deres bøjede overarm som respons i fysiske samtaler. Sammen med andre metaforiske anvendelser taler dette for ikke udelukkende at forstå brugen af kropsemojier på samme måde som fysiske kropstegn. Selv om der udføres eksperimenter med at udvikle "recognition systems for mapping symbolic hand gestures to analogous emojis" (Koh et al. 2019: 1), er en direkte oversættelse af den fysiske krop til digitale tegn som emoji'er således både vanskelig og kun til dels rammende.

En anden udfordring ved at se emoji'er som kropssprog er at vi ikke kan stole på en universel enighed om forståelsen af hver enkelt emoji. Selv om Ekman antager at der er syv universelle følelser, og at disse kan genkendes i menneskers mimik (Ekman & Friesen 1975), er de betydningsnuancer der formidles af mimik, gestik og stemmeføring langt flere og mere komplekse at fortolke. Det skyldes at det her gælder, som med enhver anden semiotisk resurser, at receptionsprocessen forløber som en fælles systematisk fortolkningsprocedure hvor den individuelle afsender forudser en afkodning, og den individuelle modtager foretager faktiske afkodninger af repræsentationer (Ravelli & Van Leeuwen 2018) – begge præget af kontekst i bredest mulige forstand:

'[U]nderstanding' itself is subject to the principle of contemporality, i.e. is limited by what, at any given time, participants are aware of and how they contextualize this in relation to past and (projected) future experience. (Harris 1998: 105, i Duncker 2019: 179)

Ligesom emotikoner kan emoji'er ses som "a highly sensitive, and versatile, semiological instrument that enables participants to calibrate the indeterminacy-determinacy balance" i deres udtalelser (Duncker 2019: 179), men måske er en vis grad af semantisk diversitet ikke uønskelig da

det dels åbner for en mere flydende, antydende og legende kommunikationsform, dels nemmere kan tilpasses forskellige kulturer og kontekster (Pohl, Domin & Rohs 2017, Herring & Dainas 2017).

6 KONKLUSION

Som kontekstmarkør og inferenspunkt hjælper emoji'er læseren med at identificere den sociale diskurs de er involveret i. De er videreudviklinger af og supplerer til tidligere kompenserende funktioner i online interaktion så som emotikoner, versaler og andre former for lydefterligning (Pavalanathan & Eisenstein 2016, Hougaard 2004), men de kan også fungere som 'illocutionary force indicating devices' (Dresner & Herring 2010, 2013). Emoji'er bruges desuden som følelsesangivelse og som samtalebidrag i sig selv, fx når tommel op-emoji'en eller den lattemilde emoji bruges som minimalrespons (Herring & Dainas 2017).

Med denne artikel har jeg ønsket at supplere forståelsen af emoji'er med en kropsfunderet begrundelse for den udbredte brug af emoji'er. Via Peirces semiotiske tegnforståelse har jeg vist hvordan brugen af emoji'er både er ikonisk, symbolsk og indeksikalsk: Ikonisk fordi emoji'er ligner det som de betegner, symbolsk fordi betydningen af mange emoji'er er bestemt via konvention, og indeksikalsk fordi særligt kropsemoji'erne skaber forbindelse mellem den fysiske krop og den medierede og medialiserede krop. Emoji'er bruges som en kropsforankret semiotisk resurse der kompenserer for det manglende kropssprog idet kropsemoji'ernes mimik og gestik let fortegnet ligner den fysiske krops bevægelser og ansigtsudtryk. Kropsemoji'er bruges som digital mimik og gestik og kan tilmed tolkes som tillempede angivelser af tonefald og prosodi. Som kontekstualiseringstegn hjælper emoji'er fortolkningsprocessen, men bringer også kroppens kommunikation og nærvær ind i interaktionen og fremstår som manifestationer af noget der kan være svært at udtrykke skriftligt (Salló 2011: 8, Herring, Stein & Virtanen 2013). Som digitale kropstegn viser emoji'er tilbage til den fysiske krop og skaber en forbindelse til den selv om skriftlige interaktionen egentlig er kropsløs.

Den digitale krop materialiserer sig på forskellig vis, og selve det at taste på et tastatur og finde den rigtige emoji er også en kropslig aktivitet. Hvad enten man som Sundén (2003) kalder det et overlap mellem den digitale og den fysiske krop, eller man som i denne artikel ser emo-

jier som digitale kropstegn der gennem en form for indeksikalisering henleder opmærksomheden på den fysiske krop og lader den komme til orde, er jeg med min undersøgelse nået frem til at konstatere at kroppen har en central betydning særligt for brugen af kropsemojier. Som kropstegn understreger og supplerer emojiernes verbalsprog og gør den ofte hastigt producerede interaktion tydeligere og mere nuanceret, appellerende og nærværende.

TAK

Jeg vil gerne takke en anonym fagfælle samt Jacob Thøgersen og Michael Nguyen fra NyS-redaktionen for grundige og konstruktive kommentarer til en tidligere version af artiklen.

Tina Thode Hougaard, lektor
Afdeling for Nordiske Studier, Aarhus Universitet
north@cc.au.dk

LITTERATUR

- Amin, H. 2022. Emojis and the neoliberal coding of diversity: A monocultural multiculturalism. *on_culture (online)* 14. 1-26. DOI: 10.22029/oc.2022.1302.
- An, J., T. Li, Y. Teng & P. Zhang. 2018. Factors influencing emoji usage in smartphone mediated communications. G. Chowdhury, J. McLeod, V. Gillet & P. Willett (red.), *Transforming digital worlds*. iConference 2018. Lecture Notes in Computer Science vol 10766. 423-428. Springer, Cham. DOI: 10.1007/978-3-319-78105-1_46.
- Androutsopoulos, J. 2011. Language change and digital media: A review of conceptions and evidence. T. Kristiansen & N. Coupland (red.), *Standard languages and language standards in a changing Europe*. 145-159. Oslo: Novus. DOI: 10.1515/9783110472226-033.
- Androutsopoulos, J.K. 2018. Digitale Interpunktion: Stilistische Ressourcen und soziolinguistischer Wandel in der informellen digitalen Schriftlichkeit von Jugendlichen. A. Ziegler (red.), *Jugendsprachen/Youth languages: Aktuelle Perspektiven internationaler Forschung/Current perspectives of international research*, 722-748. Berlin: De Gruyter.

- Austin, J.L. 1982. *How to do things with words: the William James lectures delivered at Harvard University in 1955*. 2. udg. Oxford: Oxford University Press.
- Bai, Q., D. Qi, M. Zhe & Y. Maokun. 2019. A systematic review of emoji: Current research and future perspectives. *Frontiers in Psychology* 10. DOI: 10.3389/fpsyg.2019.02221.
- Barbieri, F., K. German, R. Francesco & S. Horacio. 2016. How cosmopolitan are emojis?: Exploring emojis usage and meaning over different languages with distributional semantics. *Proceedings of the 2016 ACM on Multimedia Conference*, 531-533. DOI: 10.1145/2964284.2967278.
- Berthelsen, U.D., T.T. Hougaard & M.V. Christensen. 2019. Medialektur i et grammatisk og didaktisk perspektiv. Y. Goldshtein, I.S. Hansen & T.T. Hougaard (red.), *17. Møde om Udforskningen af Dansk Sprog* 17, 95-111. Aarhus Universitet.
- Bolter, J.D. & R. Grusin. 1999. *Remediation: understanding new media*. Cambridge, MA: MIT Press. DOI: 10.1108/ccij.1999.4.4.208.1.
- Bolter, J.D. 1991. *Writing space: the computer, hypertext, and the history of writing*. Hillsdale, N.J.: L. Erlbaum Associates.
- Cherny, L. 1999. *Conversation and community. Chat in a virtual world*. Lealand Stanford Junior University: CSLI Publications.
- Cohn, N., J. Engelen & J. Schilperoord. 2019. The grammar of emoji? Constraints on communicative pictorial sequencing. *Cognitive Research: Principles and Implications* 4(1). 2-18. DOI: 10.1186/s41235-019-0177-0.
- Cramer, H., P.d. Juan & J. Tetreault. 2016. Sender-intended functions of emojis in US messaging. *Proceedings of the 18th International Conference on Human-Computer Interaction with Mobile Devices and Services*, 504-509. New York: ACM. DOI: 10.1145/2935334.2935370.
- Crystal, D. 2006. *Language and the Internet*. 2. udg. Cambridge: Cambridge University Press.
- Daft, R.L. & R.H. Lengel. 1984. *Information richness: a new approach to managerial behavior & organizational design*. L.L. Cummings & B. M. Staw (red.), *Research in organizational behavior*, Vol. 6, 191-233. Homewood, IL: JAI Press. DOI: 10.21236/ADA128980.
- Dainas, A. & S.C. Herring. 2021. Interpreting emoji pragmatics. C. Xie, F. Yus & H. Haberland (red.), *Approaches to internet pragmatics: Theory and practice*, 107-144. Amsterdam: John Benjamins. DOI: 10.1075/pbns.318.04dai.

- Danesi, M. 2017. *The semiotics of emoji, Bloomsbury advances in semiotics*. London: Bloomsbury Academic. DOI: 10.5040/9781474282024.
- Danet, B. 1998. Text as mask: Gender, play, and performance on the Internet. S.G. Jones (red.), *Cybersociety 2.0: Revisiting computer-mediated communication and community*, 129-158. Thousand Oaks, CA: SAGE. DOI: 10.4135/9781452243689.n5.
- Dresner, E. & S.C. Herring. 2010. Functions of the nonverbal in CMC: Emoticons and illocutionary force. *Communication theory* 20(3). 249-268. DOI: 10.1111/j.1468-2885.2010.01362.x.
- Dresner, E. & S.C. Herring. 2013. *Emoticons and illocutionary force*. Dordrecht: Springer Netherlands. DOI: 10.1007/978-94-007-7131-4_8.
- Due, B.L. 2016. Fælles orientering som ressource for idéudvikling: En single case-analyse baseret på Distributed Cognition (DC) & Conversation Analysis (CA). *NyS, Nydanske sprogstudier* 1(50). 86-119. DOI: 10.7146/nys.v1i50.23799.
- Duncker, D. 2019. *The reflexivity of language and linguistic inquiry: Integrational linguistics in practice*. London: Routledge. DOI: 10.4324/9781351060394.
- Dürscheid, C. & C.M. Siever. 2017. Jenseits des Alphabets – Kommunikation mit Emojis. *Zeitschrift für germanistische Linguistik* 45(2). 256-285. DOI: 10.1515/zgl-2017-0013.
- Ekman, P. & W.V. Friesen. 1969. The repertoire of nonverbal behavioral categories. *Semiotica* 1. 49-98. DOI: 10.1515/semi.1969.1.1.49.
- Ekman, P. & W.V. Friesen. 1975. *Unmasking the face: A guide to recognizing emotions from facial clues*. Prentice Hall: Englewood Cliffs.
- Evans, V. 2017. *The emoji code: the linguistics behind smiley faces and scaredy cats*. New York: Picador.
- Garcia, A.C. & J.B. Jacobs. 1999. The eyes of the beholder: Understanding the turn-taking system in quasi-cynchronous computer-mediated communication. *Research on Language and Social Interaction* 32(4). 337-367. DOI: 10.1207/S15327973rls3204_2.
- Gawne, L. & G. McCulloch. 2019. Emoji as digital gestures. *Language@Internet* 17(2).
- Gerlitz, C. & A. Helmond. 2013. The like economy. *New Media and Society* 15(8). 1348-1365. DOI: 10.1177/1461444812472322.
- Giles, D., W. Stommel & T.M. Paulus. 2017. The microanalysis of online data: The next stage. *Journal of Pragmatics* 115. 37-41. DOI: 10.1016/j.pragma.2017.02.007.
- Gumperz, J.J. 1982. *Discourse strategies*. Cambridge: Cambridge University Press. DOI: 10.1017/CBO9780511611834.

- Hansen, M.H. & A.C. Stæhr. 2021. Sproglige generationsforskelle på de sociale medier. *NyS, Nydanske Sprogstudier* 1(59). 113-156. DOI: 10.7146/nys.v1i59.121811.
- Herring, S.C. & J. Ge-Stadnyk. 2023. Emoji and illocutionarity: Acting On, and acting as, language. M. Gill, A. Malmivirta & B. Warvik (red.), *Structures in discourse: Studies in interaction, adaptability, and pragmatic functions*, 123-154. Amsterdam: John Benjamins.
- Herring, S.C. & A.R. Dainas. 2020. Gender and age influences on interpretation of emoji functions. *ACM Transactions on Social Computing* 10 (Special Issue on Emoji Understanding and Applications in Social Media). 1-26. doi: DOI:10.1145/3375629.
- Herring, S.C. & A.R. Dainas. 2017. "Nice picture comment!" Graphicons in Facebook comment threads. *Proceedings of the Fiftieth Hawai'i International Conference on System Sciences*. DOI:10.24251/HICSS.2017.264.
- Herring, S.C., D. Stein & T. Virtanen. 2013. Introduction to the pragmatics of computer-mediated communication. S.C. Herring, D. Stein & T. Virtanen (red.), *Handbook of pragmatics of computer-mediated communication*, 3-31. Berlin: Mouton. DOI: 10.1515/9783110214468.
- Hjarvard, S. 2007. Sprogets medialisering. *Språk i Norden*. 29-45.
- Hougaard, T.T. 2004. *Nærmest en leg. En undersøgelse af sprog og interaktion i dansk webchat*. Ph.d.-afhandling. Aarhus Universitet.
- Hougaard, T.T. 2014. Sproglige forandringer i de nye medier : fra chatstil til hashtag-poesi. *NyS, Nydanske Sprogstudier* 46. 39-66. DOI:10.7146/nys.v46i46.17524.
- Hougaard, T.T. 2020. Emojis. *Tænkepauser* (84). Aarhus: Aarhus Universitetsforlag. DOI: 10.2307/j.ctv34wmr49.
- Hougaard, T.T. & L.B.C. Lund. 2019. Mediale træk i digitale dialoger. *Språk i Norden* 49(1). 7-23.
- Hougaard, T.T. & M. Rathje. 2018. Emojis in the digital writings of young Danes. A. Ziegler (red.), *Jugendsprachen/Youth Languages. Aktuelle Perspektiven internationaler Forschung/Current Perspectives of International Research*, 773-806. Berlin: De Gruyter. DOI: 10.1515/9783110472226-035.
- Jaeger, S.R. & G. Ares. 2017. Dominant meanings of facial emoji: Insights from Chinese consumers and comparison with meanings from internet resources. *Food Quality and Preference* 62. 275-283. DOI: 10.1016/j.foodqual.2017.04.009.
- Jakobson, R. 1960. Closing statement: Linguistics and poetics. T.A. Sebeok, *Style in language*, 350-377. New York og London: The Technology Press of Massachusetts Institute of Technology og John Wiley & Sons.

- Jensen, T.W. 2009. Følelser og sprog - den verbale manifestation af følelser i sproglig interaktion. *NyS, Nydanske Sprogstudier* 37. 34-60. DOI: 10.7146/nys.v37i37.13502.
- Junge, T.A. 2019. Hører du overhovedet efter? En undersøgelse af talerens mimik i en situation med manglende respons. *NyS, Nydanske Sprogstudier* 56. 133-158. DOI: 10.7146/nys.v1i56.111288.
- Jørgensen, M. 2018. Samtalegrammatik i skriftlig online interaktion: En analyse af de tre kommentarindledende partikler *altså*, *øhm* og *hm* i to danske Facebookgrupper. *Skandinaviske sprogstudier* 9(1). 24-51. DOI: 10.7146/sss.v9i1.109462.
- Kelly, R. & L. Watts. 2015. Characterising the inventive appropriation of emoji as relationally meaningful in mediated close personal relationships. *Experiences of technology appropriation: Unanticipated users, usage, circumstances, and design*. University of Bath.
- Kendon, A. 2004. *Gesture: Visible action as utterance*. Cambridge: Cambridge University Press.
- Kendon, A. 2017. Reflections of the "gesture-first" hypothesis of language origins. *Psychon Bull Rev* 24. 163-170. DOI: 10.1017/CBO9780511807572.
- Koh, J., J. Cherian, P. Taele & T. Hammond. 2019. Developing a hand gesture recognition system for mapping symbolic hand gestures to analogous emojis in computer-mediated communication. *ACM transactions on interactive intelligent systems* 9(1). 1-35. DOI: 10.1145/3297277.
- Lebduska, L. 2014. Emoji, emoji, what for art thou? *Harlot* 12. DOI: 10.15760/harlot.2014.12.6.
- Lévy, P. 1998. *Becoming virtual: reality in the digital age*. New York: Plenum.
- Malinowski, B. 1923. The problem of meaning in primitive languages. C.K. Ogden & I.A. Richards (red.), *The meaning of meaning*, 451-510. London: Routledge & Kegan Paul.
- Martin, J. R. & M. Zappavigna. 2019. Embodied meaning: a systemic functional perspective on paralinguistic. *Functional linguistics* 6(1): 1-33. DOI: 10.1186/s40554-018-0065-9.
- McNeill, D. 1992. *Hand and mind - What gestures reveal about thought*. Chicago: University of Chicago Press.
- McNeill, D. 2006. Gesture: A psycholinguistic approach. R.E. Asher & J.M.Y Simpson (red.), *The Encyclopedia of Language and Linguistics*, 58-66. DOI: /10.1016/B0-08-044854-2/00798-7.

- Meredith, J. & J. Potter. 2013. Conversation analysis and electronic interactions: Methodological, analytic and technical considerations, H. Lim & F. Sudweeks (red.), *Innovative methods and technologies for electronic discourse analysis*, 370-393. Hershey, PA: IGI Global.
- Miller, H., D. Kluver, L. Thebault-Spieker, L. Terveen & B. Hecht. 2017. Understanding emoji ambiguity in context: The role of text in emoji-related miscommunication. *Proceedings of the International AAAI Conference on Web and Social Media*, 152-161 DOI: 10.1609/icwsm.v11i1.14901.
- Miller, H., J. Thebault-Spieker, S. Chang, I. Johnson, L. Terveen & B. Hecht. 2016. "Blissfully happy" or "ready to fight": Varying interpretations of emoji, 259-268. *International AAAI Conference on Web and Social Media*.
- Moschini, I. 2016. The "face with tears of joy" emoji. A socio-semiotic and multimodal insight into a Japan-America mash-up. *HERMES-Journal of Language and Communication in Business* 55. 11-25. DOI: 10.7146/hjlb.v0i55.24286.
- Murr, L.E. & J.B. Williams. 1988. Half-brained ideas about education: Thinking and learning with both the left and right brain in a visual culture. *Leonardo* 21(4). 413-419. DOI: 10.2307/1578704.
- Na'aman, N., H. Provenza & O. Montoya. 2017. MojiSem: Varying linguistic purposes of emoji in (Twitter) context. *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics- Student Research Workshop*. 136-141. Vancouver: Association for Computational Linguistics. DOI: 10.18653/v1/P17-3022.
- Nielsen, M.F. & S.B. Nielsen. 2010. *Samtaleanalyse*. Fredriksberg: Samfundslitteratur.
- Novak, P.K., J. Smailović, B. Sluban & I. Mozetič. 2015. Sentiment of emojis. *PLoS One* 10(12). 1-21. DOI: 10.1371/journal.pone.0144296.
- Ozanne, M., A.C. Navas, A.S. Mattila & H.B. Van Hoof. 2017. An investigation into Facebook "liking" behavior - an exploratory study. *Social Media + Society* 3(2). DOI: 10.1177/2056305117706785.
- Pavalanathan, U. & J. Eisenstein. 2016. Emoticons vs. emojis on Twitter: A causal inference Approach. *ArXiv abs/1510.08480*.
- Peirce, C.S. 1995. *Semiotik og pragmatisme. Moderne tanker*. København: Samlerens Bogklub.
- Pohl, H., C. Domin & M. Rohs. 2017. Beyond just text: Semantic emoji similarity modeling to support expressive communication. *ACM transactions on computer-human interaction* 24(1). 1-42. DOI: 10.1145/3039685.

- Provine, R.R., R.J. Spencer & D.L. Mandell. 2007. Emotional expression online: Emoticons punctuate website text messages. *Journal of Language and Social Psychology* 26(3). 299-307. DOI: 10.1177/0261927X06303481.
- Quinta, N. da , E.S. Cruz, Y. Rios, B. Alfaro & I.M.d. Maraón. 2023. What is behind a facial emoji? The effects of context, age, and gender on children's understanding of emoji. *Food Quality and Preference* 105(3). DOI: 10.1016/j.foodqual.2022.104761.
- Ravelli, L.J. & T.V. Leeuwen. 2018. Modality in the digital age. *Visual communication*. 17(3). 277-297. DOI: 10.1177/1470357218764436.
- Riordan, M.A. 2017. Emojis as tools for emotion work: Communicating affect in text messages. *Journal of Language and Social Psychology* 36(5). 549-567. DOI: 10.1177/0261927X17704238.
- Rodrigues, D., D. Lopes, M. Prada, D. Thompson & M.V. Garrido. 2017. A frown emoji can be worth a thousand words: Perceptions of emoji use in text messages exchanged between romantic partners. *Telematics and Informatics* 34(8). 1532-1543. DOI: 10.1016/j.tele.2017.07.001.
- Salló, S. 2011. The faces of messenger emoticons in the virtual communication. *Acta Universitatis Sapientiae. Social Analysis* 1(2). 307-332.
- Schandorf, M. 2012. Mediated gesture: Paralinguistic communication and phatic text. *Convergence: The international Journal of Research into New Media Technologies* 19(3). 319-344. DOI: 10.1177/1354856512439501.
- Schegloff, E.A. 1982. Discourse as an interactional achievement: some uses of 'uh huh' and other things that come between sentences. D. Tannen (red.), *Analyzing discourse: Text and talk*, 71-93. Georgetown: Georgetown University Press.
- Schegloff, E.A. 1996. Turn organization: one intersection of grammar and interaction. E. Ochs, E.A. Schegloff & S.A. Thompson (red.), *Interaction and grammar*, 52-133. United Kingdom: Cambridge University Press. DOI:10.1017/CBO9780511620874.002.
- Scott Sørensen, A. & B.K. Walther. 2001. *Cyberflux: digital kunst og kommunikation, Cyberflux; 1*. Odense: Odense Universitetsforlag.
- Skovholt, K., A. Grønning & A. Kankaanranta. 2014. The communicative functions of emoticons in workplace e-mails: :-). *Journal of Computer-Mediated Communication* 19. 780-797. DOI:10.1111/jcc4.12063.
- Sproull, L. & S. Kiesler. 1986. Reducing social context cues: Electronic mail in organizational communications. *Management Science* 32(11). 1492-1512. DOI: 10.1287/mnsc.32.11.1492.

- Stage, C. & T.T. Hougaard. 2018. *The language of illness and death on social media. An affective approach*. D.R. Christensen & K. Sandvik (red.), *Sharing death online*, Vol. 1. Bingley: Emerald Publishing.
- Stark, L. & K. Crawford. 2015. The conservatism of emoji: Work, affect, and communication. *Social Media + Society* 1(2). 87-96. DOI:10.1177/2056305115604853.
- Steensig, J. 2001. *Sprog i virkeligheden: bidrag til en interaktionel lingvistik*. Århus: Aarhus Universitetsforlag.
- Stivers, T. 2008. Stance, alignment, and affiliation during storytelling: When nodding is a token of affiliation. *Research on language and social interaction* 41(1). 31-57. DOI: 10.1080/08351810701691123.
- Sundén, J. 2003. *Material virtualities*. New York: Peter Lang.
- Togeby, O. 2019. *Om tegn analogi, konvention, situationalitet*. Aarhus: Aarhus Universitetsforlag. DOI:10.2307/j.ctv34wmv9b.
- Turkle, S. 1997. *Life on the screen: identity in the age of the internet*. New York: Touchstone.
- Watzlawick, P., J.B. Bavelas & D.D. Jackson. 1967. *Pragmatics of human communication: a study of interactional patterns, pathologies, and paradoxes*. New York: W. W. Norton.
- Werry, C.C. 1996. Linguistic and interactional features of Internet Relay Chat. S.C. Herring (red.), *Computer-mediated communication. Linguistic, social and cross-cultural perspectives*, 47-63. Amsterdam: John Benjamins. DOI: 10.1075/pbns.39.06wer.

Anmeldelse af Peter Widell:

Sproget i verden. Tekstforståelsens filosofiske grundlag

Aarhus Universitetsforlag. 2023. (355 sider)

IB ULBÆK

INDLEDNING

Peter Widell har gennem mange år været lektor i nordisk sprog ved Aarhus Universitets Institut for Kommunikation og Kultur og har siden 2014 være lektor emeritus. Widell har en stor produktion af forskningsartikler bag sig, men har i de senere år fået mulighed for at udgive monografier. I 2018 udkom *Hvad er et billede? Fire billedbegreber – et opgør med illusionismen og den semiotiske billedteori*. Og nu denne bog, der ønsker at gå dybere i funderingen af en tekstteori, nemlig ved at give en egentlig filosofisk fundering.

Jeg skal i det følgende give en oversigt over indholdet i bogen, og jeg vil forsøge at indkredse, om denne bestræbelse lykkes, eller i hvilket omfang den lykkes.

BOGENS OPBYGNING OG INDHOLD

Ud over et forord består bogen af ti kapitler, men også af tre tillægskapitler: et appendiks om modallogik; definitioner af grundbegreber og den formallogiske notation. Litteraturlisten udmærker sig ved også at angive hovedlitteraturen for de overordnede emner, bogen behandler, fx for logik, talehandlinger og metaforer. Hvert kapitel afsluttes med ”spørgsmål, refleksioner og uddybninger”, hvor læseren kan afprøve sin viden. Desværre uden ”facitliste”, så knap så velegnet til selvstudium. Nogle eller fx hver anden opgave burde have et løsningsforslag bag i bogen.

FØRSTE KAPITEL: INDLEDNING

Dette korte indledende kapitel opridser bogens afgørende valg af begre-

ber til at undersøge tekstforståelse, nemlig de tre begreber *instrumentel handling*, *talehandling* og *implikatur*. Ud af den instrumentelle handling kommer den kommunikative handling. Det kendte skel mellem saglitteratur og skønlitteratur angiver Widell, at man kan forstå klart i lyset af de foregående begreber. Forholdet mellem mislykket og vellykket kommunikation skal belyses gennem den argumentative diskurs.

Men under det hele ligger logikken som fundament og forudsætning. De fremsættende sætninger belyses gennem logisk analyse, hvor propositioner/udsagn relateres til hinanden gennem slutninger, dvs. inferenser mellem præmisser frem mod en konklusion. Og herunder særligt de kontekstbundne slutninger kendt fra pragmatisk sprogfilosofi, implikaturerne. Begrebet *hjemmel* kendt fra Toulmins *The Uses of Argument* (1958) bruges til at lagdele virkeligheden ud fra Widells tredeling: ting (kausalitet), agent (instrumentel handling), person (kommunikativ handling). Herefter følger så tre tekstdomæner kendt fra både litteraturteori og sproglig tekstanalyse (tekstlingvistik): saglitteratur, skønlitteratur og genrer.

Kapitlet afsluttes med forklaring af litteraturlisterne, noter og mulige læsestrategier. Nu skulle læseren være klædt på til at læse selve bogen.

ANDET KAPITEL: LOGIK

Kapitlet begynder med at introducere en kendt distinktion inden for logisk teori, nemlig om logikken skal forstås enten semantisk eller syntaktisk, dvs. som repræsentation af kendsgerninger, modelteoretisk logik, eller som et slutningsbaseret system, en ren formalisme. Forskellen kan beskrives som forskellen mellem en teori der opererer med fuldt tegn med indhold (semantik) og udtryk (syntaks), kontra en teori som opererer med ren syntaks som de gamle genskrivningsregler hos fx Chomsky. Widell vælger side til fordel for den modelteoretiske logik, hvor atomerne i udsagnslogikken er propositioner, der kan være enten sande eller falske.

Hvis propositionen angår en kendsgerning, der indgår i den model eller mængde, den fortolker, så er propositionen sand, og ellers er den falsk. Den propositionelle logik er den enkleste logik, og på den bygger den såkaldte førsteordens prædikatlogik (se nedenfor). Forskellen er, at propositionslogikken brydes op, og det viser sig, at atomet indeholder

elementer, ligesom virkelighedens atomer, der indeholder elektroner, protoner og neutroner. Propositionen kan indeholde argumenter og prædikater, som i udsagnet "Postkassen er rød", der formelt noteres som $R(p)$. Dette udsagn er sandt, hvis og kun hvis den verden, den mængde, det fortolker, indeholder en rød postkasse (og den er den eneste postkasse).

Udtrykket "førsteordens-" relaterer sig til det særlige ved denne nye formelle logik. Den tyske matematiker og filosof Gottlob Frege introducerede i 1879 to kvantorer, der som funktioner kan tage argumenter og binde dem til prædikater og returnere en sandhedsværdi. "Der eksisterer noget, som ..." (eksistenskvantor), "Alle noget (sic!) er ..." (alkvantor). "Der eksisterer noget, som er rødt" er sandt, hvis mængden indeholder en rød entitet. Så det er sandt, hvis vores mængde fx indeholder den røde postkasse. "Alle ting er røde" er falsk, hvis vores mængde indeholder elementer, der ikke er røde, fx en blå bil. Andenordenslogik er, hvis man kvantificerer over prædikaterne. Modallogik er, hvis man indfører operatorer, der virker på propositioner, fx $N(P)$ (nødvendigtvis P) eller $M(P)$ (muligvis P). Den tilknyttede semantik til modallogikken er ofte den såkaldte muligverdenssemantik. Tolkningen af operatorerne er så, at hvis noget er nødvendigvis sandt, så er det sandt i alle mulige verdner, og hvis noget er muligt sandt, så er der mindst én verden, hvor udsagnet er sandt. Der er en hær af andre logikker, som ikke behandles i bogen, for Widells centralpåstand er, at den eneste nødvendige logik er førsteordens prædikatlogik (s. 47).

Logik bruges ikke bare til at repræsentere individuelle sagsforhold, men de kan bindes sammen på en måde, så man kan slutte noget ud fra deres indbyrdes forhold, der er udtrykt via konnektorer. Det er de konnektorer, der er kendt fra logikkens grundbøger, og som relaterer udsagnslogikkens ikke-analyserede propositioner til hinanden: p og q ; p eller q ; hvis p , så q ; ikke- p (p 's negation) m.fl. Relationen mellem propositionerne kan opstilles som præmis(ser) for en konklusion og tolkes semantisk med sandhedsværdierne s og f (sandt og falsk). Ved hjælp af sandhedstavler for de forbundne udsagn kan man afgøre om konklusionen følger logisk af præmisserne. Widell viser dog ikke, hvordan ægte førsteordens sætninger med eksistens- og alkvantorer forholder sig til hinanden.

Widell vælger at bruge Toulmins slutningsskemaer, for selv om de langt hen ad vejen blot er klassiske syllogismer, kan de bruges i tekstanalysen. Det at rekonstruere en fuldstændig syllogisme og eksplicite det underforståede er en del af det praktiske arbejde med Toulmin-skemaerne. Det er det, Aristoteles kalder *entymeme*: nemlig den del af syllogismen, som udtrykkes manifest, og som kan fuldstændiggøres af modtager af teksten. Det underforståede i syllogismen tages op i kapitel fem om implikaturer.

De mere overordnede pointer i kapitlet er, (i) at logikken ikke er sprogbundet, logikken gælder nemlig for handling og perception; (ii) at tanken heller ikke er sprogbundet i anden forstand end at den følger logikkens love; (iii) at modalitet kan behandles uden modallogik og muligverdenssemantik, og blot med forståelsen for handlingens instrumentalitet (s. 56). Det sidste er emnet for næste kapitel.

TREDJE KAPITEL: DEN INSTRUMENTELLE HANDLING

Widells ontologi består af en tredelt stratificering af verden: Der er en verden af ting, der er en verden af agenter, som handler på og med tingene, og der er en verden af personer, som handler indbyrdes med hinanden i et instrumentalt rum. Dette kapitel tager sig af ting og agent, det næste kapitel tager sig af person-kategorien.

Agenter handler målrationelt, bruger midler til at opnå deres mål. Dvs. agenten har en intention, som kan opfyldes instrumentelt, ved at agenten virker i verden. Målet er at slå græs, midlet er en græsslåmaskine, og agenten opnår sit mål ved at anvende græsslåmaskinen til at slå græsset. Når målet er nået, ophører aktiviteten. Ifølge Widell er instrumentel handlen normbaseret, dvs. at målet opnås vellykket ved at følge en bestemt regel. Man laver te ved at følge en bestemt procedure for tefremstilling. Mislykket handlen analyseres kontrafaktisk: Hvis man havde handlet korrekt, så var målet nået. Denne form for muligverden-analyse accepteres af Widell, fordi den er forankret i den faktiske verden og ikke alle mulige verdner.

Handling kræver viden, og den viden indhentes via perception, men selve handleskemaet er ikke perciperbart. Widell angiver, at skemaet har en transcendental karakter i Kants forstand (s. 70). Og viden angår den fysiske og biologiske verden, der omgiver os. Naturen har kausal

karakter, men igen ser vi den ikke, vi forudsætter den transcendentalt som kausale skemaer. For Kant er hændelsers kausalitet en konsekvens af menneskelig fortolkning. Kausalitet er ikke i verden, men tillægges af den menneskelige beskuer. Dette fører Widell ind i en omfattende diskussion af forholdet mellem kausalitet og intentionalitet. Han fastholder, at de to domæner er uafhængige af hinanden og ikke kan reduceres til hinanden. Det ene ville være en kausal determinisme, det andet en platonisk idealisme.

”I det øjeblik vi identificerer os som handlende væsener, må vi nødvendigvis forudsætte frihed, defineret som mulighed for at vælge mellem alternativer. Men derfor kan intentionaliteten netop ikke reduceres eller elimineres til fordel for kausaliteten (s. 77).”

Hvis instrumentel handlen foregår mellem mennesker, kan det forstås som at behandle hinanden som ting, men handlen kan også være interaktionel og føre frem til kollektiv handling og bevidsthed, der kan munde ud i institutionelle handlinger, herunder talehandling, der foregår i en kollektiv ramme af forpligtelser.

FJERDE KAPITEL: TALEHANDLINGEN

Med introduktionen af talehandlingsbegrebet er vi nu i den tredje sfære efter ting og agent. Her er der tale om, at mennesker samvirker på en samfundsmæssig baggrund, hvor regler og institutioner gælder indbyrdes, ved at deres eksistens anerkendes. Penge er en værdi, fordi samfundets borgere anerkender dem som sådan. En pengeseddel eller nogle tal på en konto tæller som et givet beløb, der giver ejeren mulighed for at købe varer og tjenester. Denne tælle som-funktion ligger også bag det at ytre en sætning. Den ytre sætning er ikke blot en lyd eller betydning, men er en social handling, der tæller som noget mellem afsender og modtager, fx en anmodning, en hilsen eller en oplysning.

Widell følger til en vis grad den traditionelle pragmatiske sprogteoris forståelse af talehandlingen som en tredelt størrelse efter John Austin og en typologi med et antal typer af talehandling, men han flytter typeinddelingen fra den illokutionære funktion til den perlokutionære. Der er ifølge Widell kun en type illokution, nemlig assertion, den hævde eller fremsættende talehandling. Den kan så få forskellige udformninger, alt efter hvad afsender har til hensigt i forhold til modtageren, om

det er at informere, beordre, love etc. For at tydeliggøre intentionen kan afsender betjene sig af metaassertiver, fx ”jeg advarer dig, det bliver koldt”. Så bliver der ikke kun asserteret, at deres tur foregår i kolde omgivelser, men at modtager bør tage sine forholdsregler. Af perlokutionære handlinger skelner Widell mellem informativ, kommissiv og direktiv, dvs. meddele, forpligte og beordre.

FEMTE KAPITEL: IMPLIKATURER

Dette kapitel er også klassisk pragmatisk sprogforskning. Begrebet stammer dog ikke fra Austin og Searle, men fra Paul Grice. Implikaturen er en form for indirekte sproghandling, hvor man bruger en ytring til at underforstå et andet indhold. Implikaturen er grundet i samarbejdsprincippet, at afsender og modtager ikke fører hinanden bag lyset. Implikaturen, den intenderede, men ikke udtrykte mening, opstår, fordi afsender tilsyneladende ikke samarbejder og derfor siger afsender enten noget forkert, for lidt eller for meget, noget irrelevant eller udtrykker sig forkert i konteksten. Men givet samarbejdsprincippet, så kan modtager lede efter den intenderede mening. Klassikeren ”Kan du nå saltet?” spørger ikke til modtagers mulighed for at gøre det, men ved at bryde relevansmaksimen ønsker afsender at få modtager til at række denne saltet.

Da Widell hævder assertivet som fundamental talehandling, så fremsætter han den – kontroversielle? – påstand, at de andre talehandlinger skyldes implikaturer. Searle (2013 [1964]) har et skema, hvor ethvert udtrykt udsagn både bærer en proposition (p) og en ”force-indikator (F)”, så talehandlingen kan skrives F(p). F er så hans talehandlingstyper, fx repræsentativer, direktiver, ekspressiver etc. Hos Widell bliver det til kun en force-indikator, den assertive, som så er overflødig, fordi den altid er der, dvs. ingen force-indikator er nødvendig. Forskellen kan måske illustreres således: ”Stop!” kunne forstås direkte, dvs. ud fra formen, som et direktiv. Widell mener, at det må være en ellipse for et assertiv: ”Du skal stoppe” forstået som en ren beskrivelse (af en fremtidig hændelse). Den bliver først til en direktiv via en implikatur, der passer til situationen, hvor modtager indser, at der er et indirekte budskab: ”Jeg beordrer at stoppe”.

Præsuppositioner, den forudsatte information, behandles også som implikaturer, hvilket er i overensstemmelse med traditionen. ”Min bil

er på værksted”: Her er subjektet forudsat, nemlig at jeg har en bil, og at modtager ved det. Men om det er en implikatur helt generelt, er mere usikkert. Hvis man bruger udtrykket, vel vidende at modtager ikke ved det, så har man overtrådt maksimet om sandhed, men da man samarbejder, vil modtager slutte, at man ikke præsupponerer, at man har en bil, men at man bruger udtrykket som ny viden. Til gengæld vil de klassiske påduttende præsuppositioner som ”Er du holdt op med at drikke?” ikke være baseret på implikaturer, fordi afsender ikke samarbejder, men modarbejder modtager.

Widell ser også på to andre former for implikaturer, nemlig (i) eksplikatur, tydeliggørelse eller berigelse, og (ii) metaassertiv tydeliggørelse, dvs. når man i stedet for at anvende et assertiv benytter et metaassertiv til at tydeliggøre typen af talehandling. ”Jens åbnede vinduet” kan tydeliggøres ved at angive stedet: ”Jens åbnede vinduet i sit hus”. Og tydeliggørelsen kan foregå via metaassertivet: ”Jeg fortæller dig: Jens åbnede vinduet.”

SJETTE KAPITEL: DE GRUNDLÆGGENDE HJEMMELTYPER

Af dette kapitel fremgår det, at Widell ikke kun bruger Toulmin-ske-maet som pædagogisk opstilling af argumenter, men at han også bruger det til at samle det filosofiske system sammen, som han har opstillet i de foregående fem kapitler. Hjemlerne bruges af Toulmin (og Widell) til at binde påstand og belæg sammen ved hjælp af almene love og regulariteter. Hjemlerne er de domæner, som man kan have viden om og derfor kan ræsonnere over. Der er syv hjemmeltyper, selv om overskriften til 6.8 angiver seks: logik, metafysisk lovmæssighed, materiel-analytisk lovmæssighed, empirisk lov, instrumentel handleregel, definition og æstetisk generalisation.

SYVENDE KAPITEL: DEN SAGLITTERÆRE TEKST

Efter 165 sider er vi nået frem til de tekster, de foregående seks kapitler har forberedt. Kapitlet om saglitterære tekster er bogens længste, men alligevel blot 70 sider langt for så stort et emne. Der er mange emner i gang for at komme rundt om tekstbegrebet. Den overordnede kommunikative handlen deler sig i samtale og tekster henholdsvis karakteriseret ved mangel på struktur versus strukturering. Der er kort om

tekster i tekster, referat og citat, herunder skjulte citater, der peger på intertekstualitet.

Sagteksters egenart består i, at de overholder de griceske maksimer. Tekster opnår en klar struktur ved at have ét overordnet sigte og ét overordnet emne. De er sekvenser af assertiver bundet sammen af logiske konnektorer, og som perlokutionære handlinger kan de inddeles i kommissive, direktive eller informative tekster.

Herefter kommer de traditionelle tekstlingvistiske begreber i spil: tekstreferenter, anaforicitet, kohæsion og kohærens. Implikaturerne er med til at skabe kohærens. Makrostrukturteorien introduceres og eksemplificeres.

Tekstlingvistikken forlades igen for at gennemgå fire typer af argumentative diskurser. Det diskursive register opstår, når brud på forståelsen optræder, og så forståelse derefter retableres af den diskursive forhandling. De fire områder for diskurs er: filosofisk-analytisk diskurs, teoretisk-empirisk diskurs, praktisk diskurs og æstetisk diskurs. De fire diskurstyper gennemgås over 30 sider af høj sværhedsgrad.

OTTENDE KAPITEL: DEN SKØNLITTERÆRE TEKST

Widell bestemmer den skønlitterære tekst i to retninger, dels som en parasitær tekst, dvs. en uegentlig tekst, der ligner ”rigtige” tekster, og dels som en tekst, hvor der er et særligt fokus på det sproglige udtryk med henblik på at skabe en æstetisk effekt.

Det parasitære – forstået korrekt – har at gøre med fikcionalitet, at disse tekster foregiver ægte reference og sandhed. De lader, som om de handler om virkelige verdner, men har i virkeligheden ingen reference til en virkelig verden. Widell afviser at bruge muligverdener i analysen af det fiktive.

Det æstetiske har at gøre med det, Roman Jakobson kaldte den poetiske funktion i tekster. Teksterne har en afvigende sprogbrug i forhold til normaltekster, hvorfor disse særtræk kommer i fokus. Det gælder retoriske figurer som rim og rytme og gentagelse og de retoriske troper metafor, metonymi og symbol. Her har Widell sin egen metafor-teori, der gør op med den allestedsnærværende teori fra Lakoff og Johnson først fremsat i *Metaphors we live by* (1980) og viser fornuften i den klassiske metafor-teori, der går tilbage til Aristoteles.

Tekstlig realisme i skønlitterære tekster er ikke en mere eller mindre lighed med realiteten, som i malerkunsten, men i forholdet til de normalsproglige tekster, som de mimer og afviger fra. Den skønlitterære realisme har ikke noget at gøre med filosofisk realisme.

NIENDE KAPITEL: GENRER

Genrer forstås som forventningsmønstre, som letter både afsender og modtager i henholdsvis skrive- og læsarbejdet.

Tekstgenrerne bindes sammen med talehandlingstyperne informative, direktiver og kommissiver. Teksterne indeholder genremarkører, som muliggør deres identifikation. De inddeles overordnet i fiktive og ikke-fiktive genrer og efter talehandling (s. 281). Der foretages yderligere opdeling efter indhold og efter, om tekster er argumentative eller ikke. Kapitlet afsluttes med tre tematiske genreanalyser, om teksterne vedgår årsag-virkning, det instrumentelle eller det institutionelle, igen ud fra den ontologiske forskel på ting, agent og person.

DISKUSSION AF BOGENS HOVEDÆRINDE

På mange måde ligner bogen en doktorafhandling – og burde måske have været indleveret som sådan. Peter Widells viden er legendarisk, og hans kamp for at bygge modsigelsesfrie systemer er det også. Det har jeg stor respekt for. På den måde er bogen en storslået kæmpe i sit ønske om at skabe en metafysik ud af ”ting, agent, person” og ende som grundlag for tekstteori. Et bedømmelsesudvalg på dokterniveau er tekstens rette modtager.

Men – og der er et men – det er ikke en doktorafhandling, men struktureret som en lærebog til universitetsstuderende med pædagogiske gennemgange og øvelser til at afprøve, om man har forstået de enkelte kapitler. Jeg har for år tilbage prøvet at bruge et tidligere manuskript af nærværende bog på et kandidatkursus i tekst og tekstlingvistik. Jeg tænkte, at det kunne være en fin afslutning på kurset, hvor de studerende kurset igennem havde tilegnet sig forudsætninger for at kunne læse Widells manuskript. Men det gik ikke. Den første halvdel af manuskriptet svarende til de 165 sider satte dem skakmat. Jeg gennemgik logik etc. og lovede så, at sidste kursusgang så blev væsentlig nemmere, da vi så var tættere på tekstlingvistikken, som vi

havde gennemgået. Da dagen så oprandt, havde ingen læst den sidste del, så også dér måtte jeg gennemgå pointerne og desværre stort set nyttesløst, da de jo ikke engang kendte stoffet. Jeg har efterfølgende tænkt en del over, hvad der kunne være gjort. Der er to løsninger. Den ene er at korte de første 165 sider ned til ca. 20 sider, hvor de absolut nødvendigeste forudsætninger kunne stå, og så gå til kapitel syv og ud. Den anden var at begynde med kapitel syv og så indarbejde forudsætningerne, efterhånden som de behøves. Men desværre tror jeg ikke, at nogen af forslagene ville være til at føre ud i livet. Den lange indånding, før Widell når frem til tekstkapitlerne, betyder således, at der er ganske lidt plads til at udfolde tekstbegreberne på. Som bemærket ovenfor er kapitlet om saglitterære tekster bogens længste, men alligevel mærkelig kort. Også fordi gennemgangen af de fire former for diskurs stjæler 30 sider, så når spørgsmål, refleksion og uddybninger tælles med her, er der kun 34 sider tilbage til at analysere tekster og uddrage konklusioner på basis heraf. Ole Tøgeby's *Praxt* fylder i modsætning hertil immervæk over 800 sider med pragmatisk tekstteori fra først til sidst (Tøgeby 1993).

WIDELLS TEKSTBEGREB

Som det fremgår af gennemgangen af *Sproget i verden* ovenfor, så begynder den egentlige forholden sig til tekster i kapitel syv og otte, om den saglitterære tekst og den skønlitterære tekst. Det er for mig overraskende, at interessen for tekster begynder i to underbegreber uden at forklare, hvorfor der ikke kan siges noget om overbegrebet, nemlig tekst. Hvorfor er der ikke et indledende kapitel om tekster i almindelighed? Nu er de almene tekstovervejelser klemmt ned i sagprosa-kapitlet. Har saglitterære og skønlitterære tekster ikke en fællesmængde, som man kan sige noget meningsfuldt og dybt om? Tilsyneladende ikke, for Widell behandler tekstbegrebet som sådan meget stedmoderligt. Der er to steder i bogen, som kort definerer *tekst an sich* – og de er ikke engang enige:

”Idet vi i det følgende vil skelne mellem tekst og samtale, vil vi ved en tekst forstå en meningsfuld sekvens af korrekt formulerede sætninger, såkaldte **talehandlinger**, produceret af kun én person ... (s. 8).”

”Tekster er talehandlingssekvenser udført af én og samme afsender (som godt vil kunne omfatte flere personer) inden for én og samme taletur i en samtale eller i forbindelse med kommunikativ handlen generelt (s. 339, det samme siges s. 166).”

Det sidste citat er fra kapitel tolv, Definition af grundbegreber, indgang 15) og ikke 14), som angivet i den alfabetiske nøgle s. 323.

Men det er en kursorisk behandling af tekster. Jeg opregner selv i min lille bog om tekstlingvistik (Ulbæk 2005) 10 kendetegn på, hvad en tekst er, Beaugrande og Dressler (1981) opregner 7 standarder for, hvad der udgør en tekst.

Tekster udgør ikke bare talehandlingssekvenser, tekster udgør *organiserede* sekvenser af sætninger, propositioner eller talehandlinger. En tekst er jo ikke et løbende bånd af sætninger, den starter et sted og slutter et andet sted. Og begge de steder er ikke tilfældige. Teksters organisering behandles ikke i bogen. Den tekstlige orden er væsentlig. Brydes den, brydes tekstens væv, dens sammenhæng, jf:

- a. Der sidder en mand ved Netto. Han drikker øl.
- b. Han drikker øl. Der sidder en mand ved Netto.

Måske skyldes denne forglemmelse, at Widell trækker på udsagnslogikken, når han forklarer, hvordan tekster er bundet sammen:

”Først og fremmest er tekster bundet sammen ved logiske konnektorer, for det meste som udtryk for (logiske) konjunktioner, og normalt blot angivet som punktummer (s. 178).”

Men logisk set er der ingen forskel på sekvenserne: a & b & c; c & b & a; b & a & c osv. For det gælder for alle sekvenserne, at de kun udgør en sand enhed, hvis de alle er sande. Her spiller rækkefølge ingen rolle. Men det gør det for både tekstproducenter og tekstkonsumenter. Vi vurderer ikke de enkelte udsagn isoleret, men i lyset af hinanden. Hvis b kommer efter a, så aflæses b's bidrag til teksten i lyset af a. Tekster opdateres sekventielt i modsætning til fx billeder, der ikke aflæses som en tekststreng,

men følger andre organiseringsprincipper. Derfor kan det også undre, at det kun er udsagnslogikken, der gennemgås i kapitel 2. Den er en delmængde af prædikatlogikken, men behandler som angivet ovenfor de enkelte propositioner som atomer. Prædikatlogikken er mere avanceret, fordi vi kan se ind i propositionen. Så hvor udsagnslogikken behandler udsagnet "John er høj" som et hele, fx symboliseret med p , så kan prædikatlogikken symbolisere John og $er_høj$ som argument og prædikat: $H(j)$. Og det er nødvendigt, hvis man fx vil sige noget om koreference. Henvisninger til fx Hans Kamps diskursrepræsentationssemantik (Kamp & Reyle 1993) er fraværende trods det, at hvis man er interesseret i den logiske opbygning af tekster, så skal man kunne behandle anaforiske udtryk. Kamps semantik er en del af en bredere *update semantics*, hvor det betyder noget, hvor meget information man er præsenteret for, for næste information ses i lyset af de forudgående propositioners bidrag. Men når man, som Widell, kun arbejder med udsagnslogikken, så er der grænser for, hvor meget logikken kan belyse tekstbegrebet.

Noget andet, som er underbelyst i Widells redegørelse, er den tekstlige organisering. Hvis vi kun ser teksten som en lang række af forbundne sætninger eller talehandlinger, så kan vi ikke se den organisering, tekster har. Bøger har kapitler, kapitler har underinddelinger, som selv har afsnit. Her gælder den regel, at et afsnit er mere kohærent internt end mellem to afsnit. Men de er alligevel anbragt efter hinanden, fordi de er mere kohærente indbyrdes end afsnit længere væk i teksten. Alle Widells teksteksempler (som ikke er mange givet den lille plads) er korte og uden afsnit, og derfor kan han i sagens natur ikke sige noget om afsnits funktion i tekster. Eller det kunne han selvfølgelig, men det falder ikke naturligt. Og det falder åbenbart uden for hans interessehorisont. Bøger begynder et sted og slutter et andet. Hvorfor gør de det? Hvis man vil fortælle noget om et problem, så begynder man med at introducere problemet. Derefter foreslår man en løsning. Og slutter med at fortælle, at nu er problemet løst. Det begyndelse, midte, slutnings-skema har været kendt siden oldtidens retorik. Eller hvis man har en narrativ tekst, så er der en naturlig fortælleorden, der spejler det fortalte tidslighed. Nogle gange afsluttes en tekst, ved at man gentager formuleringer, der er nævnt i begyndelsen, som for at sige, at nu er der sagt det, der skulle siges om den sag.

FIKTION OG MULIGE VERDNER

Kapitlet om den skønlitterære tekst er også kort, men der er en god redegørelse for metaforer, som viser, at den aristoteliske metafor-teori stadig er valid. Som sagt er muligverdenssemantik ikke Widells kopie, her følger han den amerikanske filosof Quine i dennes afvisning (se s. 53). Han siger således om fiktive tekster, at forfatteren allerede i parateksten, fx ved at kalde bogen for roman, har afmonteret et krav om sandhed:

”Om en tekst er en fiktionstekst eller ej, afgøres i virkeligheden én gang for alle og *normalt for teksten som helhed* ved, at forfatteren enten para- eller intratekstligt over for læseren signalerer, at *sandhedsspørgsmålet i teksten er irrelevant* (s. 269).”

Det er jo ikke hele sandheden om fiktion og fiktive tekster. Martin A. Hansens *Løgneren* kan tjene som eksempel. Sandø, som historien om degnen og skolelæreren Johannes Vig og de andre personer foregår på, er opdigtet, og personerne ligeså. Men det er så også det eneste, som afviger fra Danmark i 1940'ernes slutning. Forfatteren er forpligtet på, at hovedstaden er København, at kongen er Frederik d. 9., og alt muligt andet. Den mulige verden, som Sandø nødvendiggør, er stort set som den aktuelle verden, blot med denne lille ”udposning”. Men der er meget af det, som foregår på Sandø, som også er sandt: at der findes snepper, hunde, træer, jagtgeværer, kirker etc. At der er menneskelige reaktionsmønstre, som svarer til virkelige mennesker: forelskelse, bedrag, jalousi, planlægning, overvejelser, handlingslamelse ... Teksten forekommer ”realistisk” af den simple grund, at Sandø som mulig verden ligger så nær den aktuelle verden, som den gør. Omvendt er en science fiction-roman som *Rumrejsen 2001* (Clarke 1983 [1968]) meget fjern fra den aktuelle verden, fordi den indeholder entiteter, som ikke har eksisteret, fx den monolit, som sætter gang i menneskets udvikling væk fra aberne, og som stammer fra en ekstraterrestrisk civilisation. Omvendt skabte Karl Ove Knausgård seksbindsværket *Min kamp* (2009-2011), som han deklarerede roman – og derfor skulle signalere, at ”sandhedsspørgsmålet i teksten er irrelevant”. Men det er den jo langt fra, hvad virkningshistorien for

længst har bevidnet. Men den autofiktive roman kan i øvrigt sagtens læses fiktivt om en person, der tilfældigvis hedder det samme som forfatteren.

Der er et særligt aspekt ved skønlitterære tekster, som Widell ikke berører: at der er to afsendere og to modtagere. Der er de virkelige afsendere og modtagere: Martin A. Hansen, der skrev *Løgneren* (1950), og de virkelige personer, der læste og læser hans værk. Og så er der fortælleren Johannes Vig, der beretter sin historie via sin dagbog og adresserer Nathaniel (manden uden svig fra Det ny testamente). Men han kunne også fortælle til en for ham virkelig person som Clamence i Camus' *Faldet* (1957). Forfatteren (den virkelige) giver læseren adgang til en verden, hvor eksistentielle eller romantiske dramaer udspiller sig, som en person, der låner en anden person sine skuldre for at kunne skue ind i en ellers utilgængelig verden, helt banalt en fodboldkamp, som ingen af dem har råd til. Ligesom læseren ikke interagerer i den handling, der udspiller sig, gør forfatteren heller ikke. Kun Svend Åge Madsen drister sig til den slags.

ET KRÆVENDE OG VIDTFORGRENET VÆRK

Jeg spurgte indledningsvis om Peter Widell lykkes med at fundere en tekstteori filosofisk. Svaret synes at være, at det gør han kun delvis, fordi bogen to dele (kap. 1-6, kap. 7-9) taler forbi hinanden. Pointerne fra de seks første kapitler suges ikke op i de sidste tre, og begreber, som beskriver de saglitterære og skønlitterære tekster, er ikke begrundet filosofisk, det gælder især kohærensbegrebet, der er et kernebegreb for at forstå, at en tekst er én tekst.

Bortset fra fejl i enkelte af figurerne (s. 60 og s. 149 f.) er bogen stort set fejlfri. Bogen er også skrevet i et klart sprog, men med et meget krævende forudsætningsniveau. Det betyder, at bogen, skønt klart skrevet, er meget vanskelig at tilegne. Jeg skal ikke påstå, at jeg har forstået alle redegørelserne, heller ikke dem i de lange fodnoter. Bogen har mange diskussioner i mange retninger, og mange af dem er meget abstrakte og svære at følge, fx hvorfor muligverdenssemantik ikke er nogen god ide. Det er heller ikke indlysende, at kausalitet hele tiden omtales som hume'sk, al den stund Widell ikke har problemer med hverken realisme angående verdens realitet eller kant'sk transcendens.

Der kæmpes en del kampe mod teoretiske positioner, hvor relevansen ikke er indlysende.

Uanset kritikken ytret her er der tale om en stor intellektuel bedrift.

Ib Ulbæk, lektor emeritus
Institut for Nordiske Studier og Sprogvidenskab
Københavns Universitet
ibulbaek@gmail.com

LITTERATUR

- Beaugrande, R.-A. de & W. Dressler 1981. *Introduction to text linguistics*. London: Longman. DOI: 10.4324/9781315835839.
- Camus, A. 1957. *Faldet* (oversat af Frank Jæger). København: Gyldendal.
- Clarke, A.C. 1983 [1968]. 2001. København: Gyldendal. Oprindeligt udgivet 1969 på Forlaget Spektrum. Engelsk udgave 1968.
- Hansen, M.A. 1950. *Løgneren*. København: Gyldendal.
- Kamp, H. & U. Reyle. 1993. *From discourse to logic*. Dordrecht: Kluwer. DOI: 10.1007/978-94-017-1616-1.
- Knausgård, K.O. 2009-2011. *Min kamp 1-6* (oversat af Sara Koch). København: Lindhardt & Ringhof.
- Lakoff, G. & M. Johnsson. 1980. *Metaphors we live by*. Chicago: University of Chicago Press.
- Searle, J. 2013 [1964]. What is a speech act? M. Black (red.), *Philosophy in America*, 221–239. London: Allen & Unwin.
- Togeby, O. 1993. *Praxt. Pragmatisk tekstteori*. Aarhus: Aarhus Universitetsforlag.
- Toulmin, S. 1958. *The uses of argument*. Cambridge: Cambridge University Press.
- Ulbæk, I. 2005. *Sproglig tekstanalyse – Introduktion til pragmatisk tekstteori*. Aarhus: Academica.

Bogstav-lyd-forbindelser i dansk: The Ultimate Collection

Anmeldelse af Christian Becker-Christensen: *Dansk Retskrivning: Bogstav – Lyd – Bogstav. En grafemisk beskrivelse af dansk ortografi*. Syddansk Universitetsforlag. 2021. (416 sider)

HOLGER JUUL

”Findes det, findes det her”, lød sloganet for De Gule Sider i sin tid. Det samme kan man sige om bogen *Dansk Retskrivning: Bogstav – Lyd – Bogstav*, hvis forfatter, Christian Becker-Christensen (CBC), her samler stort set alt hvad der er at vide om forholdet mellem bogstaver og lyde i dansk. Bogen er dermed en meget værdifuld resurse for alle der vil have indsigt i skriftsprogets indretning. Jeg tør godt spå at dette værk vil blive stående i mange år, medmindre vi får en ortografisk revolution. Og selv da vil bogen være meget nyttig. Bogens afsnit 2.6 giver nemlig en forbilledlig fremstilling af de mange rent sproglige forhold man må tage i betragtning hvis man vil ændre ortografien i større eller mindre grad.

I det følgende vil jeg kort præsentere bogens dele for derefter at gå lidt mere i detaljer med dens tilgang til bogstav-lyd-forbindelser, med dens opdeling af ordstoffet i hjemligt og fremmed og med dens egnethed som opslagsværk.

BOGENS DELE

Efter et forord og en indledning der gør rede for bogens formål, disposition og terminologi, består bogen af tre sektioner.

Den første sektion, *Retskrivning og udtale*, giver fyldig baggrundsviden om den danske retskrivningsnorm (herunder blandt andet en redegørelse for de forskellige typer af relationer der kan være mellem skriftens bogstaver og talesprogets lydinventar); om lydsystemet i dansk (både fonetisk og fonologisk); og om lydskrift. Hvor CBC's tidligere værk *Bogstav og Lyd* (Becker-Christensen 1988) brugte Dania, bruger den nye bog samme IPA-baserede lydskrift som *Den Danske Ordbog*

(DDO) og websitet *Bogstavlyd* (bogstavlyd.ku.dk). Den aktuelle bog er således umiddelbart sammenlignelig med disse. Tak for det.

Den næste sektion, *Bogstav-lyd-relationer*, er bogens centrale del. Her giver separate kapitler om vokaler og konsonanter en grundig gennemgang af relationerne mellem bogstaver og lyde i dansk med en lang række anskueliggørende figurer og tabeller. Herefter følger tre kapitler der udskiller de relationer man finder hhv. i fremmedord, i trykssvag stavelse og i morfologisk komplekse ord. Sektionen indledes med et kapitel med titlen *Relationelle forbemærkninger* og endnu et kapitel, *Vokaler og konsonanter*, som også til dels formidler baggrundsviden, idet der her gøres rede for mulige stavelsesstrukturer i dansk, men dette kapitel tager også fat på bogens kernestof med et meget informativt afsnit om forholdet mellem vokallængde og konsonantfordobling. Sektionen afsluttes med en alfabetisk oversigt over bogstav-lyd-relationer.

Bogens sidste sektion udgøres af en række ordlister der viser relationerne mellem vokalbogstaver og -udtaler i 3.770 trykstærke stavelser som forfatteren har uddraget af opslagsordene i *Politikens Skoleordbog*.

Som det bemærkes i forordet, er det ”vanskeligt at tale om komplicerede fænomener på en anskuelig måde” (s. 13), og bogstav-lyd-relationerne i dansk er unægtelig et kompliceret fænomen. CBC har gjort sig betydelige anstrengelser for at strukturere og anskueliggøre, men man kommer ikke uden om at man som læser skal have mange grundbegreber og et rimeligt kendskab til lydskrift på plads for at få det fulde udbytte. Hertil kommer at man skal lære at finde rundt i bogen.

Når jeg i denne anmeldelse nævner enkelte forhold der kunne have været anderledes, skal det ses på baggrund at emnet *er* meget komplekst. Overordnet lykkes bogen med at formidle al denne kompleksitet, som på mange måder også er det fascinerende ved dansk ortografi.

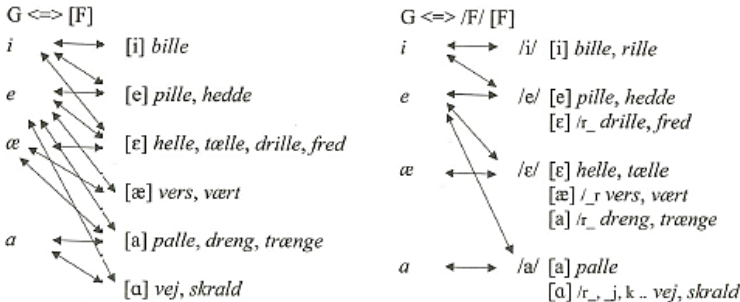
BOGSTAVER, FONEMER OG FONER

Alfabetiske skriftsprog som det danske bygger på udtalen, men afspejler ikke alle udtalenuancer. Fx staves t-lyde med bogstavet *t* uanset om den skrivende udtaler lyden ”tørt” (aspireret) eller med en markant s-fase (affrikeret). Det alfabetiske princip er teknisk sagt ikke et fonetisk princip, men et fonematisk princip: Det handler om at repræsentere de udtaleforskelle der er vigtige når ordene skal holdes ude fra hinanden,

mens ikke-betydningsadskillende nuancer såsom aspiration vs. affrikation kan ignoreres. I overensstemmelse med dette relateres bogstaverne i bogens redegørelser i første række til *fonemer*, som så først i næste række relateres til de lyde (med en teknisk term: *foner*) vi rent faktisk siger.

Et nogenlunde simpelt eksempel på fonemniveauets relevans er at bogstavet *r* på fonniveau må relateres til både en konsonantisk udtale [ʀ] som i starten af ordet *rom* og en vokalisk udtale [ɐ] som i slutningen af ordet *mor*. På fonemniveau kan *r* derimod blot relateres til fonemet /r/, hvis udtale så kan afledes af fonemets stilling i stavelsen (før eller efter vokalen), jf. bogens afsnit 8.5.10. Sprogbrugere ville næppe have megen glæde af en fonetisk orienteret ortografi hvor de to *r*-udtaler i *rom* og *mor* blev stavet forskelligt. Indskydelsen af et fonemniveau mellem bogstaver og foner gør således relationerne mere overskuelige end når bogstaver og foner relateres direkte. Dette viser CBC side 103 med de urundede fortunevokaler som eksempel – se gengivelsen i figur 1, hvor der, som det ses, er brug for langt flere pile når relationerne vises i forhold til foner (i skarpe klammer, figurens venstre side) end i forhold til fonemer (i skrå klammer, figurens højre side). (G og F er her forkortelser for grafem (bogstav) og fonem/fon).

FIGUR 1. UDSNIT AF BOGENS SIDE 103.



Hvilke fonemer der findes i et sprog, er imidlertid ikke nogen given sag. Et foneminventar er resultatet af en analyse der kan være mere eller mindre abstrakt (se fx Grønnum under udg. og Horslund m.fl. 2022). CBC vælger en analyse i den abstrakte ende, omtrent som hos Grønnum (2005). Eksempelvis anlægges bifonematiske tolkninger af lang

åben a-lyd [ɑ:] som /ar/ og af ng-lyd [ŋ] som /ng/ – altså sådan at én lyd i udtalen analyseres som en manifestation to fonemer. Abstraktioner som disse giver færre komplikationer i relationen mellem bogstav og lyd end hvis man opstiller /ɑ:/ og /ŋ/ som selvstændige fonemer, der så ortografisk repræsenteres af komplekse grafemer, *ar* og *ng*.

Man kan sige at i CBC's fremstilling fremstår dansk ortografi så regelmæssig som det overhovedet er muligt. Det har den fordel at bogen derved skelner mellem kompleksiteter der indiskutabelt kendetegner forholdet mellem skrift og tale, og kompleksiteter der i hvert fald i princippet kan beskrives rent fonologisk, dvs. i en redegørelse for forholdet mellem fonemer og foner. Bogen viser at den ortografiske side af sagen er kompliceret nok endda.

Spørgsmålet er så bare om vi som sprogbrugere faktisk har adgang til de abstrakte fonemer, eller om det primært er ortografien der får os til at opfatte det som logisk at der skal være *r* i et ord som *farve* og *g* i et ord som *fange*. Læsere der er skeptiske over for CBC's fonemanalyse, kan imidlertid blot springe dette niveau over i bogens beskrivelser.

RELATIONERNE FOR *T* OG *D* – ET EKSEMPEL PÅ BOGENS INDHOLD

For at give et nærmere indtryk af bogens stof og den mangefacetterede behandling af det, vil jeg tage afsnit 8.3 som eksempel. Afsnittet handler om konsonantbogstaverne *t* og *d* og indledes med den oversigt der er gengivet i figur 2. Her ses det at *t* på fonemniveau forbindes med to nummererede udtalemuligheder: 1. /t/ og 2. /d/, mens *d* forbindes med tre: 1. /d/, 2. /t/ og 3. /0/. Den tredje mulighed for *d*, nulfonemet, betyder at *d* kan være stumt.

I højre side af oversigten får vi så fonemerne omsat til foner med angivelse af de kontekster hvor disse foner forekommer. For /t/ drejer det sig initialt i stavelsen om [t] – hvilket i næste linje er eksemplificeret med ordene *tå* og *motto*, mens det efter *s* og finalt i stavelsen drejer sig om [d] – eksemplificeret med ordene *stå* og *fætter*. For /d/ har vi initialt udtalen [d] (ordeksempler: *due*, *dråbe*) og finalt har vi [ð], altså "blødt d" (ordeksempler: *fod*, *hedde*).

Til højre i oversigten gives også et eksempel på at *d* kan korrespondere med /t/, nemlig ordet *sekund*, hvor /t/ har [d]-udtale på grund af den

stavelsesfinale stilling. Det bemærkes at denne forbindelse – noteret *d* [d] = /t/ – forekommer i visse fremmedord. Ligeledes gives et eksempel på at *t* kan korrespondere med /d/, nemlig ordet *meget* – med [ð]-udtale på grund af /d/-fonemets finale stilling. Endelig gives der eksempler på ord med stumt *d* (*kind, bord, krudt*), og det anføres at denne nuludtale forekommer stavelsesfinalt efter *l, n* og *r* og før *s* og *t*.

FIGUR 2. RELATIONER FOR BOGSTAVERNE T OG D. FRA SIDE 187.

<i>t, d</i>	
<i>t</i> ← 1 →	/t/ [t] /fin.; [d] /fin. s_, fin.
	<i>t:</i> <i>tå, motto; stå, sætter</i>
2 ↗	<i>d:</i> <i>sekund</i> (fin. <i>d</i> [d] = /t/ i visse fremmedord)
↘ 2	
<i>d</i> ← 1 →	/d/ [d] /fin.; [ð] /fin.
	<i>d:</i> <i>due, dråbe; fod, hedde</i>
3 ↘	<i>t:</i> /fin. og _s <i>meget</i>
	/0/
	<i>d:</i> /fin. <i>l, n, r_ og _s, t</i> <i>kind, bord, krudt</i>

Jo, man skal holde tungen lige i munden, og forholdene for *t* og *d* hører endda til i den forholdsvis regelmæssige ende. Det tager lidt tid at forstå oversigtens struktur, men belønningen er at man får overblik over forhold som *er* svære at overskue.

Efter oversigten følger separate afsnit om *t* (afsnit 8.3.1, to sider) og *d* (afsnit 8.3.2-7, godt seks sider). Hvert af disse begynder med tabeller med yderligere ordeksempler for hver af bogstavets korrespondenser, og disse er så igen opdelt på foner og disses kontekster. Listerne suppleres med frekvensopgørelser fra websitet Bogstavlyd, som giver et indtryk af hyppigheden af de respektive udtalemuligheder. Bogstavlyd relaterer bogstaver direkte til foner, men CBC har bearbejdet forekomsttallene så vi får en kvantitativ opgørelse af udtalemulighederne på fonemniveau. For bogstavet *d* fremgår det således hos CBC at Bogstavlyds forekomsttal (der bygger på data fra DDO) fordeler sig på 7.329 /d/-udtaler, 663 /t/-udtaler og 5.171 stumme *d*'er. Bogstavlyds analyse af bogstav-lyd-re-

lationer bruger ikke stumme bogstaver, så også oplysningen om hyppigheden af stumt *d* skyldes CBC, der har beregnet det samlede forekomst-tal for Bogstavlyds komplekse grafemer med *d*, hhv. *ds*, *dt*, *ld*, *nd* og *rd*.

Man kunne have ønsket sig at hyppigheden af de respektive fonem-udtaler også blev givet i procenter, men dem kan den interesserede selv finde ud fra CBC's genberegninger, som må have krævet et betydeligt arbejde.

Ud over de nævnte tabeller og frekvensoplysninger gives der så kommentarer til hver af de nummererede korrespondensmuligheder. For *t* har vi således korte afsnit om hhv. $t^1 \Rightarrow /t/$ og $t^2 \Rightarrow /d/$. På blot en enkelt trykside (side 189) får vi blandt andet omtale af

- stavelsesgrænsers betydning for om $/t/$ -fonemet udtales $[t]$ som i *eks-trakt* eller $[d]$ som i *ek-stra*; yderligere bemærkes det at der kan være $[t]$ i fremmedord som *mon-tre*, vs. $[d]$ i hjemlige ord som *munt-re*.
- mulig svækkelse af $[t]$ til $[d]$ i fx *foto*, men ikke i fx *motto*.
- at finalt *-et* oftest er en aflednings- eller bøjningsendelse – men med undtagelser i ord som *broget*, *skaldet* og i enkelte fremmedord som *basket*, *cricket* – de to sidstnævnte med $[d]$ -udtale.
- at endelsen *-et* kan have $[d]$ som følge af dissimilation i ord hvor stammen ender på $[\partial]$, fx *rodet* udtalt $[\text{ʁo:}\partial\text{əd}]$.
- at udtalen af endelsen *-et* varierer regionalt, fx med $[d]$ -udtale i midt- og sønderjysk og fynsk.

Herefter noteres det, stadig på samme side, om *t* i *morfologisk komplekse ord*

- at der kan være alternation mellem t $[d]$ og t $[t]$, fx i *asfalt-as-faltre*.
- at $[d]$ kan falde bort i trekonsonantgrupper, fx i *frugtbar*, *postkasse*.
- at *Retskrivningsordbogen* 2012 har regler for hvornår man må bruge finalt *-t* som bøjningsendelse, sådan at *-t* er en mulig endelse efter *-sk* som i *frisk-t*, men ikke efter *-isk*, **kritisk-t*.

Endelig anføres det at *t* også indgår i *andre relationer* i ord som *det*, *filet*, *lektie* og *restaurant*.

Ovenstående eksempler på blot en enkelt sides informationer om et enkelt bogstav illustrerer hvor righoldig en bog der er tale om. Om *d*'s udtalerelationer får vi endnu mere at vide. Jeg kan ikke nævne det hele her, men vil fremhæve

- oversigtstabellerne (afsnit 8.3.3-4) der viser fordelingen af blødt og stumt *d* i forbindelserne *ds* (fx *rædsel* vs. *gidsel*) og *dt* (fx *kådt* vs. *vådt*) og tilsvarende for forbindelserne *ld*, *nd*, *rd*.
- afsnittet om kortformer med stumt *d*, fx verberne *bede*, *rede*, og præpositionerne *med*, *ved*. Det bemærkes at kortformen kan indgå i visse komplekse ord, men ikke i andre, fx *med-mindre* vs. *medbringe*.
- de uddybende afsnit om ord der tidligere har haft stumt *d* (fx substantiverne *Svedske* og *Pindse* og datidsformerne *kunde*, *skulde*, *vilde*) og om stavemåderne *ld* og *nd*'s funktion og historie.

Efter behandlingen af *t* og *d* hver for sig, giver afsnit 8.3.8 en oversigt over forholdene i staveretningen, altså i retningen fra lyd til bogstav. Her er udgangspunktet de foner der kan staves med de bogstaver der er hovedafsnittets emne – dvs. [t], [d] og [ð]. Endelig behandler afsnit 8.3.9 relationernes entydighed/flertydighed i retning fra bogstav til fonem og vice versa.

Hele vejen igennem afsnit 8.3 (og igennem hele bogen) gives der henvisninger til yderligere litteratur om de behandlede emner. Det er således ikke så lidt man får for pengene hos CBC. *Findes det*, *findes det her*.

FREMMED OG HJEMLIGT

Et vigtigt strukturerende greb i bogen er at fremmedord udskilles fra hjemlige ord. Ikke sådan at forstå at fremmedord slet ikke nævnes i de centrale kapitler om vokal- og konsonant-relationer (kapitel 7 og 8), for det gør de. Men fokus i disse kapitlers oversigter, jf. eksemplet i figur 2, er relationerne i hjemlige ord. For *t* medtager oversigten således ikke sj-udtalen i *lektie* og nulrelationen i *filet*. De relationer der medtages,

er dem man først og fremmest må kende for kunne stave hjemlige ord.

I kapitel 9, *Fremmedord*, får vi så en grundigere behandling af de ”andre relationer”, som i kapitel 7-8 oftest blot er noteret kort. Her giver afsnittet om *t* (9.3.18) yderligere eksempler på ord med sj-udtalen (fx *aktie*, *aktion*, *ambitiøs*). Det noteres at denne udtale af bogstavkombinationen *ti* findes i ord fra latin og fransk, at den ikke forekommer i et ords første stavelse – og at vi også har fremmedord hvor *ti* udtales med [t] (fx *gratie*, *gratiale*, *otium*). Tilsvarende gives der yderligere eksempler på ord hvor *t* er stumt. Karakteristisk for bogens grundighed nævnes det her at *Retskrivningsordbogen* til og med 2001-udgaven tillod apostrof før bøjningsendelse, fx *bidet’et*, *filet’en*. Endelig omtales det komplekse grafem *th* med udtalemulighederne [ð] som i *smoothie* og [θ] som i *thriller*.

Kapitel 9 er, som det fremgår, en slags raritetskabinet der samler de mærkværdige relationer som man finder i importerede ord. Opdelingen i fremmed og hjemligt er fornuftig fordi de specielle importordsrelationer, som det anføres i afsnit 9.1.11, oftest blot findes i nogle få og relativt lavfrekvente ord.

Det nævnte afsnit 9.1.11, der handler om importord som delsystem, påpeger at der er en tilsvarende skelnen i Nina Grønnums *Fonetik og Fonologi* (Grønnum 2005), hvor behandlingen af de særlige lydlig forhold i låne- og fremmedord adskilles fra det danske ordstof. CBC går dog ikke så vidt at han opstiller en separat fremmedordsfonologi – en opgave der vel også ville have ført for vidt i bogens sammenhæng. På fonemniveau analyseres fremmedord således som om de var danske, sådan at der ifølge CBC eksempelvis er /d/ i *smoothie*, /aʁ/ i *afrikaans* og /ŋg/ i *sæson*. Disse analyser kan diskuteres. Det er fx en forudsætning for udtalen af /d/ som [ð] at fonemet står finalt i en stavelse, men gælder dette faktisk for [ð]’et i *smoothie*? Det er jeg ikke sikker på.

Afsnit 9.1.11 omtaler også, med et citat af Pia Jarvad, at man kan betragte stavesystemet for engelske låneord som adskilt fra stavesystemet for resten af ordforrådet, fordi de engelske ord næsten altid beholder deres engelske stavemåde. Det fremgår ikke om CBC tilslutter sig dette synspunkt, men noget er der om snakken. Hvis man som jeg føler et vist ubehag ved ordformen *beautyboks*, er det nok *Retskrivningsordbogens* vaklen mellem de to systemer der er grunden: *beauty* som på engelsk + *box* fordansket til *boks*.

STIL OG STRUKTUR

CBC gør en dyd ud af ikke at skrive mere ordrigt end højst nødvendigt. Dette er vitterligt en dyd, da bogen uden de mange koncise formuleringer ville have været endnu tykkere end sine godt 400 sider. Visse steder er den forudsætningsløse læser dog nok i fare for at blive hægtet af – som når det side 44 hedder: ”Et sjældnere tilfælde er relationer, hvor man som komplekse enheder må regne både bogstaver og lyde, **GG** ↔ [**FF**], det forekommer med bl.a. *gn* [nj] i *cognac*”. Jeg er enig i udsagnet (se Juul 2010), men jeg tror det havde støttet forståelsen hvis læserne var blevet mindet om grunden til at bogstavet *n* i *cognac* ikke kan forbindes direkte med udtalens [n] – nemlig at rækkefølgen af bogstaver skal svare til rækkefølgen af lyde i udtalen (som nævnt i bogens afsnit 2.1.1).

Den koncise stil giver sig også udslag i at der ikke ofres megen plads på diskussion af forfatterens valg af metoder og empirisk grundlag. Vi får fx at vide at ordlisterne i bogens tredje sektion bygger på *Politikens Skoleordbog*, og at vokalrelationerne er ordnet efter første efterfølgende konsonant – men vi får ikke forklaret hvorfor netop denne ordbog er valgt som grundlag, eller hvorfor den nævnte ordning af ordstoffet er foretrukket.

Overordnet giver bogens struktur god mening, men det er ikke altid helt nemt at finde rundt. Der er en vis støtte i nummereringen af afsnit – men hvis man er særligt interesseret i lyd til bogstav-afsnittene, kan de ikke findes ud fra underafsnittets nummer, da dette afhænger af hvor mange foregående underafsnit der har været i det pågældende hovedafsnit. Det er desuden et vilkår at disse underafsnit er fordelt på mange forskellige hovedafsnit. Hvis man vil vide hvordan sj-lyden [c] staves, finder man således eksempler både i afsnit 8.5.3 og 8.5.13 om hhv. *j* og *s* og i en oversigt der viser konsonantlydes stavemåder i fremmedord (afsnit 9.3.23), samt under de enkelte konsonantbogstaver i fremmedordskapitlet.

Hvis man er interesseret i at finde omtalerne af grammatisk bestemte stavemåder, må man også lede lidt. Det vanskelige nutids-*r* er omtalt i kapitlet om trykssvage stavelser, nemlig i afsnit 10.2.4 og 10.2.12 om hhv. om schwa-assimilation og stavemåder for trykssvagt [ʌ] – men ikke, som man også kunne mistænke, i kapitel 11's afsnit om bøjede ord.

Lang tillægsform på *-ende* indgår i afsnit 8.3.4's oversigt over kombinationer med stumt *d* – men nævnes, så vidt jeg har kunne konstatere, ikke andre steder i bogen.

Læserens orientering støttes dog i nogen grad af bogens layout, der fremstår gennemarbejdet. Uddybende redegørelser er sat med lille skriftstørrelse, som signal om at vi her forlader hovedsporet. Ofte finder man dog ganske væsentlige oplysninger i disse afsnit. Således er det i det tidligere citerede afsnit 9.1.11 man får begrundelsen for bogens vigtige og fornuftige skelnen mellem ord af hjemlig og fremmed oprindelse. Også andre steder har de sproglige tydeliggørelser af tekstens struktur en mindre fremtrædende plads end man kunne have ønsket. Eksempelvis præsenteres strukturen i bogens anden sektion ikke ved starten af denne, men i afsnit 6.1.2. Dette kunne have været løst hvis det foregående kapitel 5 havde været integreret i bogens indledende sektion, hvilket emnemæssigt ville have givet mening.

Når man vil bruge bogen som opslagsværk, er der ofte god hjælp i bogens indeks. Dette måtte dog gerne have været fyldigere og mere gennemarbejdet. Det er fx ikke registreret at en vigtig term som *korrespondens* er forklaret side 22. Termerne *onset* og *coda* mangler selv om de er introduceret i afsnit 5.3.8. Og indekset medtager forkortelsen SDU (*Den Store Danske Udtaleordbog*), men ikke SDE (*Den Store Danske Encyklopædi*).

TAK FOR AT DELE

CBC præsenterer et komplekst stof som det ikke er let at ordne og fremstille overskueligt. De ovenanførte kritiske bemærkninger til trods, er mit hovedbudskab at der er tale om et meget værdifuldt og gennemarbejdet værk. Bogen er en gave til alle der vil i dybden med dansk ortografi. Der er stof til mange sproglige causerier hvis man med bogen i hånden sammenligner udtalen af *ea* i *sweater* vs. *teater* og de hundredvis af andre tilfælde hvor bogens formidling af flertydigheder og kringelkroge åbner ens øjne for hvor stor en opgave det i grunden er at lære at læse og stave på dansk.

Jeg har ikke kunnet yde alle bogens aspekter retfærdighed i denne anmeldelse, men afslutningsvis må jeg igen fremhæve det meget indsigtsfulde afsnit om "(u)mulige generelle reguleringer" af dansk retskriv-

ning. Det kan man håbe at redaktionen af den kommende udgave af *Retskrivningsordbogen* har studeret.

Der er næppe nogen der ved så meget om bogstav-lyd-relationer i dansk som CBC. Han fortjener en stor tak for at have samlet og struktureret bogens meget omfattende stof sådan at vi andre kan få del i denne viden – hvis vi opper os lidt.

Holger Juul, ph.d., lektor
Institut for Nordiske Studier og Sprogvidenskab
Københavns Universitet
juul@hum.ku.dk

REFERENCER

- Becker-Christensen, C. 1988. *Bogstav og Lyd. Dansk retskrivning og rigsmålsudtale*. København: Gyldendal.
- Becker-Christensen, C. 2021. *Dansk Retskrivning. Bogstav – Lyd – Bogstav. En grafemisk beskrivelse af dansk ortografi*. Odense: Syddansk Universitetsforlag.
- Grønnum, N. (under udgivelse). Fonologisk analyse af dansk. *Danske Studier* 2024.
- Grønnum, N. 2005. *Fonetik og fonologi. Almen og dansk*, 3. udg. København: Akademisk Forlag.
- Horslund, C.S., R. Puggaard-Rode & H. Jørgensen. 2022. A phonetically-based phoneme analysis of the Danish consonant system. *Acta Linguistica Hafniensia* 54(1). 73–105. DOI: 10.1080/03740463.2021.2022866.
- Juul, H. 2010. K-a-tt-e-p-i-n-er. Om komplekse bogstav-lyd-forbindelser i danske ord. *NyS – Nydanske sprogstudier* 39. 10–32. DOI: 10.7146/nys.v39i39.16848.

Evaluating language understanding in Danish LLMs based on semantic dictionaries

BOLETTE SANDFORD PEDERSEN, NATHALIE C. HAU
SØRENSEN, SUSSI OLSEN & SANNI NIMB

In this paper, we describe how we have generated a number of evaluation datasets – a so-called benchmark – in order to evaluate certain reasoning and understanding capacities of Danish language models. We hypothesize that the semantic knowledge already given in existing Danish dictionaries can be conceived as a ‘ground truth’ for the semantics of the Danish vocabulary. Our method therefore regards turning the semantic dictionaries into a number of evaluation datasets that can be used for testing how well the models understand Danish. More specifically, we examine how well the models i) understand synonymy and semantic association between concepts, ii) make inferences in relation to conceptual knowledge and inheritance structures from super-concept to sub-concepts, iii) make correct inferences in relation to acts and events, iv) disambiguate correctly when words have multiple meanings, and v) treat ‘sentiment’, i.e. positive and negative connotation, in running text. We test our datasets on ChatGPT 3.5 turbo and ChatGPT 4.0 and find that the models are challenged in an adequate manner even if ChatGPT 4.0 does perform very well on several of the datasets.

KEYWORDS: Danish language models; evaluation; language understanding; semantic lexicons; benchmark; ChatGPT

Can AI reproduce disciplinary voice?

A comparative corpus-based study of GPT-4's ability to reproduce disciplinary voice in AI-generated linguistic prose

EA LINDHARDT OVERGAARD & ULF DALVAD BERTHESEN

The purpose of this article is to uncover the extent to which generative AI models – particularly GPT-4 – can reproduce disciplinary voice in Danish academic prose. These new large language models come with a promise of transforming our writing practices, including academic writing. However, the quality of the autogenerated contributions, especially when the models are used for smaller languages like Danish, is still unclear. We focus on disciplinary voice because it is a relatively well-described phenomenon that can also be examined quantitatively through the analysis of the surface structure of texts. We specifically focus on three aspects of disciplinary voice, *stance*, *engagement*, and *discipline-specific vocabulary*, investigating them quantitatively through a corpus-based comparative study. We compare a corpus of Danish linguistic articles with a corpus of AI-generated academic prose with linguistic content. Our study indicates that AI-generated texts significantly differ from the authentic texts in some areas. In the discipline-specific vocabulary category, the differences are notably large, whereas the stance and engagement categories exhibit smaller differences. In these two latter categories, the differences are so marginal that AI-generated texts appear to reproduce the disciplinary voice in a manner that, from a quantitative perspective, mirrors what we observe in authentic texts.

KEYWORDS: generative artificial intelligence; chatbots; text quality; disciplinary voice; corpus linguistics

In search of the chatbot's linguistic fingerprint

JONAS NYGAARD BLOM, ALEXANDRA HOLSTING &
JESPER TINGGAARD SVENDSEN

Chatbots have recently entered the educational system, creating opportunities but also challenges. A particularly pressing issue is the lack of a fully reliable method to trace if a text is written by a chatbot. This means, on the one hand, that students might use chatbots to write assignments undetected, and on the other hand, that students might be suspected of using chatbots even when they have not done so. In this study, we contribute to current research on the linguistic characteristics of chatbots by examining whether it is possible to distinguish linguistically between texts written by Danish students and ChatGPT. The results show that, generally – but not in every individual case – there are differences in the writing styles of the students and the chatbot. However, at the same time, it is possible to obscure the chatbot's linguistic fingerprint using supplementary prompts. This highlights the difficulty of tracing texts written by artificial intelligence. In conclusion, we discuss the implications of this for Danish educational institutions.

KEYWORDS: chatbots; ChatGPT; artificial intelligence; linguistic authorship analysis

Emojis as paralinguistic resources in written online interaction

TINA THODE HOUGAARD

This article introduces a new perspective on emojis, viewing them as an embodied semiotic resource. Building on previous research regarding ways to compensate for the absence of physical cues in written online interaction, the article examines how emojis can serve as substitutes for nonverbal cues such as body movements, facial expressions and vocal tones. Emphasizing emojis as a prefabricated system of communication technology, the article offers a summary of prior studies on the pragmatic function of emojis. Drawing on Peirce's semiotics, the article illustrates how body language manifests itself in digital body-emojis functioning as icons, symbols and indices. The analysis investigates how participants utilize digital body language in various online interactions, suggesting that viewing emojis as embodied communication tools clarifies their effectiveness in conveying emotions and guiding recipients in understanding the sender's intentions. By mimicking the appearance and sensations of physical gestures, emojis are used by participants in online communication to compensate for the lack of nonverbal cues.

KEYWORDS: emojis; written online interaction; paralanguage; mime; gestures

NyS 67 – Temanummer

Sproglige perspektiver på klima og vejr

”Alle taler om vejret, men ingen gør noget ved det”, hedder det i et citat som tilskrives Storm P. En moderne version kunne være: ”Alle taler om klimaet, men vi er ikke enige om hvad vi skal gøre ved det”. Klimakrisen fylder i både offentlige og private diskurser, debatter og samtaler. Men taler og skriver vi om klimaet, hvilke positioneringer og perspektiver bringer vi i spil i debatter om klimaet og i vores bud på løsninger (eller ikke-løsninger) på klimakrisen? Hvis man da overhovedet anerkender at der er en krise?

Med temanummeret *NyS 67 – Sproglige perspektiver på klima og vejr* inviterer vi hermed kommunikations- og sprogforskere til at give deres forskningsperspektiver på klimaet og klimaforandringer, på klimadebatten og på vores diskurser om klima. Og vores diskurser om vejret, for klima og vejr er jo uløseligt forbundne. Det gælder ikke mindst i lægmandstilgangen til klimadiskussionen, der som oftest opstår ud fra vejrmæssige oplevelser og observationer. Med andre ord: Måden som vi omtaler og opfatter vejret på, spiller ind på måden vi omtaler og opfatter klimaet på. Det er en kliche at sige at man taler om vejret når man ikke har andet at tale om. Det mener vi ikke er sandt. Vi mener at der er overordentligt mange interessante ting at sige om vejret – og om vores sprog om vejret. Vi inviterer derfor bidrag som knytter sig til emnet klima og vejr. Bidragene kan falde under følgende punkter, som på ingen måde er udtømmende.

- **Retorikken om klimaet:** Hvordan debatteres der om klimaforandringer? Hvilke positioner indtages, hvilke topoi (eller perspektiver) og hvilke typer af argumenter bringer parterne i spil? Hvornår begyndte man egentlig at debattere klimaforandringer?
- **Kommunikation om klimaet:** Hvordan formuleres, frames og spinnes klimapolitikker af virksomheder, organisa-

tioner og politiske partier? Hvilke kommunikationsformer tages i brug i klimaopråb og 'dommedagsvisioner' (fx Greta Thunberg eller Den Grønne Ungdomsbevægelse) og i forhalingsstrategier og teknologioptimisme (fx Bjørn Lomborg, industrien og olieselskaber), og hvad kendetegner kommunikationen bag greenwashing?

- **Klimaet og vejret i medierne:** I medierne formidles klimaproblemer og vejret af fagfolk. Men hvordan? Hvilke genrer og registre tages i brug? Hvordan kobles skrift, tale, billeder og video multimodalt i formidlingen, og med hvilken effekt?
- **Samtaler om klimaet og vejret:** Hvordan forløber samtaler om klima og klimabekymringer, og hvordan evalueres vejret i samtaler? Er snak om vejret virkelig kun fatisk kommunikation? Hvilke positioneringer finder sted i dagligdagssamtaler om klima og vejr, og hvilke synspunkter affilierer man sig med?
- **Klimaet og vejrets leksikon:** Hvilke vejr- og klimaord er dansk beriget med, og hvilken verdensopfattelse afspejler disse? Viser vejrord i dialekter en forbundethed med naturen som vi har tabt? Hvornår – og hvordan – bliver meteorologiske fagord lægmandsord?
- **Klimaet og vejrets grammatik:** Vejrfænomener konstrueres ofte agentløse, "det sner", "det regner". Er klimaændringerne også agentløse, eller tilskriver vi fx havvandsstigninger og hedeølger agens? Hvilke andre grammatiske konstruktioner er kendetegnende for vejr og klima? Hvordan udnyttes principper for informationsstruktur i tale og skrift når synspunkter på klima og vejr sprogliggøres? Kommer årsager eller effekter i fokusposition?

Forslag, dvs. abstracts, til artikler på mellem 300 og 500 ord indsendes til nys.red@dsn.dk senest d. 1. september 2024. Forslagene kan være forfattet på dansk, norsk og svensk. Filerne bedes navngivet som følger: NyS67_Temanummer[forfatternavn]. En mailadresse bedes tilføjet nederst i forslaget. Forslagene bliver bedømt i redaktionen, og alle forfattere bliver kontaktet. De forslag der udvælges til eventuel publicering

i temanummeret, gennemgår den sædvanlige redaktionsprocedure for artikler til NyS, herunder bedømmelse af en anonym fagfælle. Artiklerne kan have et omfang på op til 45.000 anslag (ca. 20 sider), inklusive mellemrum, men eksklusive litteraturliste og noter. Deadline for artikler er 1. december 2024.

NyS udkommer to gange om året på www.nys.dk.
Du kan bestille tidligere trykte enkeltnumre i boghandlen
eller henvende dig til:

Dansk Sprognævn
adm@dsn.dk

Gamle numre er tilgængelige på www.nys.dk.

Vi opfordrer læsere til at registrere sig:
<https://www.nys.dk/user/register>
Brug linket Registrer øverst på www.nys.dk.
Denne registrering medfører at du via e-mail
får besked når der udkommer en ny udgave af NyS.

Bidrag til NyS

Har du lyst til at indsende en artikel,
en anmeldelse, et debatindlæg eller andet,
så kontakt redaktionen:

nys.red@dsn.dk
NyS-redaktionen
Institut for Nordiske Studier og Sprogvidenskab
Københavns Universitet
Emil Holms Kanal 2
2300 København S

Oplysninger om formater mv. kan findes på
www.nys.dk/information/authors.