

MASTERPROJEKT

Når generativ AI fejler: En undersøgelse af blindes oplevelser, strategier og tillid

RASMUS PAASCH

JUNI 2025

KØBENHAVN

VEJLEDER

**PROF. ANN KATHRIN MEILANDT BYGHOLM
AAU**

Resumé

Denne masteropgave undersøger en grundlæggende udfordring i anvendelsen af generativ AI for blinde brugere: hvordan navigerer de i en teknologi, der leverer information, de ikke selv kan verificere visuelt? Med udgangspunkt i problemformuleringen om blinde brugeres oplevelser og strategiske mønstre med Be My AI belyser undersøgelsen den informationsasymmetri, der opstår, når assisterende teknologi skaber nye former for afhængighed. Gennem en contextual inquiry-inspireret tilgang med seks danske blinde brugere (n=6) afdækker undersøgelsen tre centrale strategiske mønstre. Brugerne udvikler kalibreret tillid – en situeret, kontekstafhængig vurdering af teknologiens pålidelighed, der varierer systematisk baseret på funktionelle, kontekstuelle og konsekvensbaserede faktorer. Som respons på Be My AI's "billede først – intention bagefter"-design, der kan skabe paradoksalt ekskludering, udvikler brugerne kreative intentionskommunikationsstrategier gennem billedkomposition og promptingteknikker. Endelig fungerer brugerne som teknologihackere, der aktivt orkestrerer deres personlige teknologiske økosystem for at kompensere for systemets begrænsninger. Undersøgelsen bidrager ved at udforske, hvordan kalibreret tillid fungerer i kontekster med informationsasymmetri og positionere brugerinnovation som aktiv vidensskabelse frem for passiv tilpasning. Praktisk peger analysen på behovet for designprincipper, der understøtter brugerautonomi og intentionskommunikation i assisterende AI-teknologi.

Nøgleord: Assisterende teknologi, generativ AI, blindhed, brugerstrategier, tillid, tilgængelighedsdesign

Abstract

This master's thesis investigates a fundamental challenge in the application of generative AI for blind users: how do they navigate technology that delivers information they cannot visually verify themselves? Based on the research question about blind users' experiences and strategic patterns with Be My AI, the study examines the information asymmetry that arises when assistive technology creates new forms of dependency. Through a contextual inquiry-inspired approach with six Danish blind users (n=6), the study reveals three central strategic patterns. Users develop calibrated trust – a situated, context-dependent assessment of the technology's reliability that varies systematically based on functional, contextual, and consequential factors. In response to Be My AI's "image first – intention after" design, which can create paradoxical exclusion, users develop creative intention communication strategies through image composition and prompting techniques. Finally, users function as technology hackers, actively orchestrating their personal technological ecosystem to compensate for system limitations. The study contributes by exploring how calibrated trust operates in contexts with information asymmetry and suggesting that user innovation represents active knowledge creation rather than passive adaptation. The findings indicate potential design considerations for supporting user autonomy and intention communication in assistive AI technology.

Keywords: Assistive technology, generative AI, blindness, user strategies, trust, accessibility design

Indhold

1	Introduktion	4
1.1	Baggrund og motivation	4
1.2	Problemfelt	5
1.3	Problemformulering	5
1.4	Afgrænsning	5
2	Begreber, definitioner og overvejelser	7
2.1	Be My AI som teknologi	7
2.2	Fejltyper i brugerpraksis	8
2.3	Brugerstrategier og adaptioner	8
2.4	Tillid i fravær af visuel verifikation	9
3	Litteraturgennemgang	10
3.1	Generativ AI i assisterende teknologi for blinde	12
3.2	Fejlmekanismer og brugerstrategier	13
3.3	Sociotekniske og pragmatiske perspektiver på AI	14
3.4	Mangler i den eksisterende forskning	14
4	Videnskabsteoretisk afsæt	16
4.1	Deweys pragmatisme og inquiry-baseret læring	16
4.2	Sociotekniske perspektiver og situeret handling	18
4.3	Disability studies og relationel handicapforståelse	18
4.4	Teoretisk integration	19
5	Metode	20
5.1	Metodisk design	20
5.2	Udvælgelse af informanter	21
5.3	Dataindsamling	22
5.4	Analysestrategi	28
5.5	Metodiske begrænsninger	30
5.6	Opsamling	32
6	Analyse	33
6.1	Tillidsforhold og risikohåndtering	33
6.2	Brugerstrategier for intentionskommunikation	37
6.3	Muligheder og begrænsninger i praksis	40

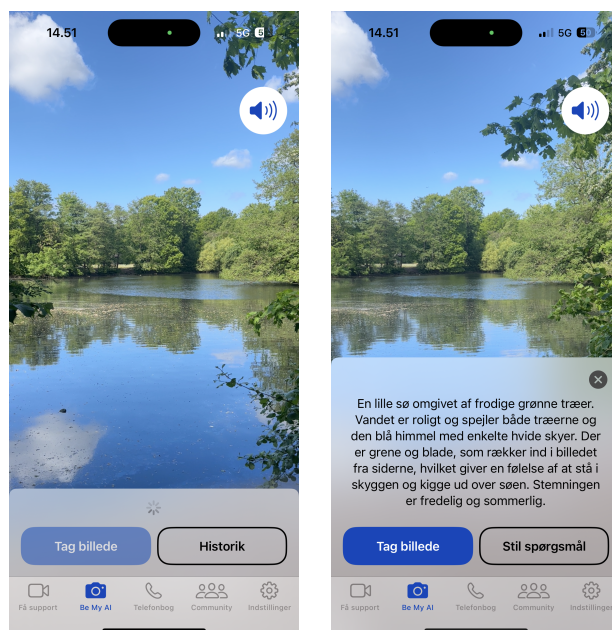
6.4 Opsamling på analyse	42
7 Diskussion	44
7.1 Kalibreret tillid som udfordring til eksisterende tillidsforståelser	44
7.2 "Billede først - intention bagefter" som designbarrierer	46
7.3 Teknologihacking som bidrag til adaptiv brugeradfærd	47
7.4 Tværgående mønstre	48
7.5 Hvordan danske fund supplerer og udfordrer internationale studier	50
7.6 Metodiske overvejelser	51
7.7 Undersøgelsens analytiske bidrag og implikationer	51
7.8 Implikationer for synsprofessionelle i Danmark	52
8 Konklusion	53
9 Bibliografi	55
Bilag A: Søgestrategi	58
Bilag B: Skærmlæsere	59
Bilag C: Interviewguide	60
Bilag D: Kodningsoversigt	62
Bilag E: Samtykkeerklæring	71

Introduktion

I mit daglige arbejde som IKT- og digital tilgængelighedskonsulent ved Institut for Blinde og Svagsynede (IBOS) er jeg blevet optaget af, hvordan generativ AI åbner nye muligheder for personer med synshandicap, samtidig med at teknologien skaber udfordringer, som både fagpersoner og brugere først nu er i fuld gang med at forstå det fulde omfang af.

1.1 Baggrund og motivation

Gennem de seneste år har jeg observeret, hvordan billedbeskrivende teknologier som appen Seeing AI og funktionen Be My AI i appen Be My Eyes har vundet indpas blandt mennesker med blindhed og svagsyn. Disse teknologier tilbyder muligheden for at "se" omverdenen gennem AI-genererede billedbeskrivelser, hvilket øger selvstændighed i en primært visuelt orienteret verden. Min faglige interesse er centreret omkring, hvordan teknologiens lovede muligheder ofte adskiller sig fra, hvad brugerne faktisk oplever i deres hverdag. En central udfordring er konsekvenserne af fejl. Når AI leverer upræcise beskrivelser, befinder brugerne sig i en sårbar situation uden mulighed for visuel verifikation. De må derfor udvikle alternative strategier for at vurdere troværdigheden eller søge bekræftelse fra seende personer. Jeg finder det interessant at undersøge, hvordan brugere aktivt former deres anvendelse af teknologien og udvikler strategier til fejlhåndtering. Denne tilpasningsproces mellem bruger og teknologi repræsenterer et væsentligt aspekt af teknologiimplementering, som fortjener dybere udforskning. Min motivation for dette projekt er dobbelt: at kvalificere min egen rådgivningspraksis samt at bidrage med viden til fremtidig udvikling af tilgængelighedsteknologier baseret på generativ AI.



Figur 1.1: Screenshots af Be My AI-funktionen i Be My Eyes som det ser ud i foråret 2025. V: Interaktionen indledes. H: Det genererede svar vises og læses højt. Kilde: Egen optagelse.

1.2 Problemfelt

Generativ AI i form af billede-til-tekst-apps som Seeing AI, Google Lookout og Be My AI lover øget selvstændighed for mennesker med synshandicap, men teknologien er stadig fejlbarlig (Adnin & Das, 2024). Særligt Be My AI adskiller sig ved et „billede først - intention bagefter“-design, hvor brugeren ikke kan angive formål, før billedet er taget. Det rejser spørgsmål om, hvordan brugerne udvikler tillid og praktiske strategier til fejlhåndtering. International forskning fokuserer primært på teknisk præcision (Gonzalez Penuela m.fl., 2024; Hong & Kacorri, 2024). Kvalitative studier af brugernes oplevelser – især danske – mangler. Derfor undersøger jeg seks danske brugere gennem interviews og observation (n=6), jf. kapitel 5.

1.3 Problemformulering

På baggrund af ovenstående problemfelt formuleres følgende problemformulering:

Hvordan oplever blinde brugere generativ AI i billede-til-tekst-tjenester som Be My AI, og hvilke strategiske mønstre udvikler de for at navigere i teknologiens funktionalitet og fejl?

For at operationalisere problemformuleringen undersøges følgende forskningsspørgsmål:

- a) Hvordan bruger blinde personer generativ AI i deres dagligdag?
- b) Hvilke typer af fejl opstår, og hvordan påvirker de brugeroplevelsen?
- c) Hvordan håndterer brugerne fejl og begrænsninger i systemet?
- d) Hvad gør, at brugerne stoler eller ikke stoler på AI?
- e) Hvilke aspekter af danske brugeres oplevelser med Be My AI, om nogen, ser ud til at supplere eller udfordre fund fra internationale studier på området?

Disse forskningsspørgsmål danner grundlag for den metodiske tilgang, dataindsamlingen og den efterfølgende analyse. Gennem en kontekstual inquiry-inspireret tilgang, hvor brugere observeres i deres naturlige kontekst, søger undersøgelsen at afdække både de eksplicitte og implicitte aspekter af interaktionen med Be My AI med særligt fokus på fejlhåndtering og tillidsopbygning.

Forskningsspørgsmålene kræver præcise analytiske begreber. Særligt Be My AI's unikke interaktionsdesign og de specifikke former for fejl og strategier, som blinde brugere udvikler, skal begrebsmæssigt afklares, før de kan analyseres i forhold til eksisterende forskning.

1.4 Afgrænsning

Min undersøgelse koncentrerer sig om blinde brugeres erfaringer med Be My AI. Denne funktion til billedbeskrivelse er en del af Be My Eyes-platformen. Jeg har bevidst valgt ikke at undersøge

andre lignende apps som Microsoft Seeing AI og Google Lookout, selvom jeg kort nævner dem i litteraturgennemgangen for at skabe kontekst. Jeg fokuserer primært på hvordan brugerne oplever teknologien og hvilke strategier de udvikler, da der mangler forskning på dette område. Valget af netop Be My AI er desuden motiveret af dens ”billede først - intention bagefter” interaktionsmodel, som, jf. mit problemfelt, rejser specifikke spørgsmål omkring brugeroplevelse og -strategier, der er centrale for denne undersøgelse. Be My AI’s interaktionsmodel adskiller sig fra fx Seeing AI, da man i sidstnævnte app kan foretage bestemte valg på forhånd for i højere grad at kunne få korrekte og valide svar. Jeg går ikke i dybden med de tekniske detaljer bag Be My AI (der bygger på GPT-4v), men nævner kun tekniske aspekter når det hjælper til at forstå brugernes oplevelser.

For at sikre en sammenhængende kulturel og sproglig kontekst har jeg valgt at begrænse min undersøgelse til danske brugere. Ved at fokusere på danske brugere kan jeg gennemføre grundige observationer og interviews inden for den tid, jeg har til rådighed. Jeg er bevidst om, at dette valg afgrænser undersøgelsens generaliserbarhed, men prioriterer dybde frem for bredde i dataindsamlingen.

Jeg har valgt en kvalitativ undersøgelsesmetode med seks informanter. Dette giver ikke grundlag for statistiske generaliseringer, men tillader mig i stedet at gå i dybden med hver enkelt informants erfaringer og dermed bedre forstå de komplekse måder, teknologien bruges på.

Tidsmæssigt fokuserer min undersøgelse på brugernes aktuelle erfaringer med den version af Be My AI, der var tilgængelig i foråret 2025.

Før jeg præsenterer det teoretiske landskab, vil jeg i kapitel 2 afklare de centrale begreber, der danner grundlag for undersøgelsen.

Begreber, definitioner og overvejelser

I dette kapitel præsenterer jeg centrale begreber og definitioner, der danner grundlag for undersøgelsen. Disse begrebsafklaringer er baseret på litteraturgennemgangen, som uddybes i kapitel 3. For at skabe et klart fundament for læseren præsenteres begreberne først, mens den bredere forskningskontekst og placering i litteraturen følger i det efterfølgende kapitel.

2.1 Be My AI som teknologi

Be My AI er en billedbeskrivende funktion integreret i Be My Eyes-appen, baseret på GPT-4's multimodale modeller. Teknologien repræsenterer et væsentligt skift fra tidligere billedgenkendelsesteknologier, idet den ikke blot kan identificere og kategorisere objekter, men kan generere kontekstuelle beskrivelser ved at kombinere visuel afkodning med omfattende sproglige modeller. Min forståelse af dette bygger på Adnin og Das (2024) og Gonzalez Penuela m.fl. (2024).

For min undersøgelse af Be My AI's særlige interaktionsdesign af central betydning. I modsætning til teknologier som Microsoft Seeing AI, hvor brugeren på forhånd kan vælge funktioner (tekstlæsning, objektgenkendelse, personidentifikation), følger Be My AI et "billede først - intention bagefter"-princip. Dette betyder, at brugeren først må tage et billede, før formålet med interaktionen kan specificeres. Se figur 2.1. Brugeren skal til en start tage et billede eller uploade et, før det bliver muligt at præcisere, hvilken information der søges. Se figur 1.1 til sammenligning. Dette design skaber en unik kontekst for at undersøge, hvordan brugerne udvikler strategier for at kommunikere deres intentioner til systemet, når de ikke kan styre konteksten på forhånd.

Teknologien fungerer ved at lade GPT-4v analysere og generere en indledende tekstbeskrivelse, der automatisk læses højt af af den indbyggede tekst-til-tale-funktion. Herefter kan brugeren stille opfølgende spørgsmål til det samme billede, hvilket muliggør en dialogisk interaktion med det



Figur 2.1: Screenshot af Seeing AI. Bemærk muligheder for at vælge funktion. Kilde: Egen optagelse.

visuelle indhold. Netop dialogen adskiller sig markant fra klassiske billedgenkendelsesteknologier og skaber grundlag for de komplekse interaktionsstrategier, som analyseres i min undersøgelse.

I projektets sammenhæng betegner ”Be My Eyes” hele app-plattformen, mens ”Be My AI” specifikt refererer til den generative AI-funktion, der er undersøgelsens fokus.

2.2 Fejltyper i brugerpraksis

I analysen af brugernes interaktion med Be My AI opererer jeg med fire kategorier af fejl, der fokuserer på fejlenes praktiske betydning for brugeroplevelsen frem for teknisk præcision. Her er jeg primært informeret af Alharbi m.fl. (2024), Bisante m.fl. (2023) samt Gonzalez Penuela m.fl. (2024).

Hallucinationer defineres som situationer, hvor Be My AI beskriver elementer, der ikke er til stede i billedet. Disse fejl er særligt problematiske for blinde brugere, da de ikke har mulighed for visuel verifikation af påstandene.

Kontekstuelle fejl omfatter misforståelser af sammenhænge eller relationer i billedet. Dette inkluderer både spatial forvirring (fejlagtig beskrivelse af objekters placering) og funktionel misforståelse (fejlfortolkning af objekters formål i den givne kontekst).

Informationsudeladelser refererer til væsentlige detaljer, som Be My AI overser eller ikke nævner i sin beskrivelse. Disse mangler kan få praktiske konsekvenser, særligt når kritisk information for brugerens formål udelades.

Relevansbaserede fejl beskriver AI’ens tendens til at fokusere på detaljer, der er irrelevante for brugerens sandsynlige informationsbehov, samtidig med at væsentlige elementer nedprioriteres eller ignoreres.

Disse fejlkategorier kan optræde samtidigt i en enkelt interaktion og forstås primært ud fra deres konsekvenser for brugerens handlemuligheder frem for som objektive målinger af systemets tekniske formåen.

2.3 Brugerstrategier og adaptationer

Brugerstrategier forstår jeg som adaptive praksisser, der udvikles gennem gentagen interaktion med AI-teknologi. Min kategorisering af disse strategier er informeret af forskning i situeret læring og fejlhåndteringsstrategier (Adnin & Das, 2024; Alharbi m.fl., 2024; Herskovitz m.fl., 2023; Hong & Kacorri, 2024; Suchman, 1987).

Proaktive strategier omfatter handlinger, brugeren foretager forud for eller under billedoptagelse med henblik på at optimere AI’ens beskrivelse. Dette inkluderer teknikker til billedkomposition, objektplacering og miljøkontrol, der udvikles gennem erfaring med systemets responsemønstre.

Reaktive strategier aktiveres efter modtagelse af AI’ens beskrivelse og omfatter verifikation, korrektion og strategiskift. Disse strategier fungerer som kvalitetskontrol og fejlhåndtering, hvor brugeren vurderer, korrigerer eller supplerer den modtagne information.

Metakognitive strategier involverer brugerens refleksion over egen interaktion og udvikling af mentale modeller for teknologiens funktionalitet. Dette omfatter læring om systemets styrker og

svagheder, udvikling af forventninger til output-kvalitet, og strategisk valg af hvornår teknologien anvendes.

Disse strategier er dynamiske og udvikles kontinuerligt gennem praktisk erfaring. De fungerer som brugerens måde at navigere i teknologiens muligheder og begrænsninger i konkrete anvendelsessituationer.

2.4 Tillid i fravær af visuel verifikation

Blinde brugeres tillidsforhold til AI-teknologi rummer en grundlæggende udfordring: hvordan kan man stole på information, man ikke selv kan verificere? Når brugeren ikke kan se billedet for at tjekke AI'ens beskrivelse, opstår der en særlig form for tillidsopbygning, som adskiller sig markant fra almindelig teknologibrug. Min forståelse af dette tillidsforhold bygger på forskning i AI-tillid og verifikationsstrategier (Alharbi m.fl., 2024; Bendel, 2024; Gonzalez Penuela m.fl., 2024; Stangl m.fl., 2023).

I analysen arbejder jeg med tre former for tillid, som blinde brugere navigerer mellem: Funktionel tillid handler om, hvorvidt brugeren stoler på systemets tekniske kvalitet – leverer det sammenhængende og troværdige beskrivelser? Kontekstuel tillid drejer sig om relevans – giver AI'en den information, brugeren faktisk har brug for i situationen? Konsekvenstillid involverer risikovurdering – hvad sker der, hvis AI'en tager fejl i denne konkrete situation?

Forskning viser, at blinde brugeres tillid til AI-beskrivelser generelt er ret lav, hvilket understreger, hvor vigtige alternative strategier til at tjekke informationen er. Tillidsforholdet påvirkes også af etiske overvejelser som privatlivsbeskyttelse og spørgsmål om, hvorvidt teknologien respekterer brugerens værdier.

Tillid udvikles gradvist gennem erfaring. Jo oftere brugeren oplever, at AI'en leverer pålidelig information i forskellige situationer, jo bedre bliver vedkommende til at vurdere, hvornår teknologien kan stoles på – og hvornår der er grund til skepsis.

Disse begreber danner grundlag for min analyse, hvor jeg forstår Be My AI-interaktion som situeret problemløsning. Brugernes strategier udvikles gennem praktisk erfaring med teknologien, hvor fejl ikke blot er tekniske mangler, men forstyrrelser i samspillet mellem bruger, AI og kontekst, som kræver kreative tilpasninger.

Når blinde brugere navigerer mellem disse tre former for tillid baseret på konkrete erfaringer med teknologien, udvikler de det, som J. D. Lee og See (2004) betegner som kalibreret tillid – en dynamisk afstemning af tillid efter systemets faktiske kapaciteter i forskellige kontekster. I modsætning til binær tillid eller mistillid justeres denne tillidsrelation kontinuerligt gennem praktisk erfaring med teknologiens styrker og begrænsninger.

Litteraturgennemgang

Denne litteraturgennemgang belyser, hvordan mit studie forholder sig til eksisterende forskning. Gennemgangen identificerer tre centrale forskningsområder: Anvendelse af generativ AI for blinde, fejlmekanismer og brugerstrategier, samt tillidsaspekter. Gennemgangen sætter projektet i forhold til eksisterende forskning, identificerer centrale begreber og afgrænser fokusområder for analysen af brugernes oplevelser, strategier og tillidsforhold til Be My AI.

Litteraturgennemgangen anlægger et pragmatisk perspektiv, hvor teknologien undersøges gennem dens konkrete anvendelse, og hvor brugernes praktiske strategier og erfaringer står centralt.

For at sikre en systematisk tilgang er litteraturen identificeret gennem strukturerede søgninger i flere databaser: ACM Digital Library, Scopus, Google Scholar, ERIC og PsycINFO. Søgningen er udført i februar-marts 2025 med en struktureret og systematisk tilgang.

Den oprindelige søgestrategi var organiseret omkring tre tematiske dimensioner med tilhørende centrale søgeord, der afspejler de termer, der er anvendt i databaserne:

1. Målgruppe: Personer med synshandicap (fx blind*, "visual* impair*", "low vision", accessib*)
2. Teknologi: Anvendelsen af Generativ AI i specifikke billedassisterende teknologier (fx "Be My Eyes", "Be My AI", samt overordnede teknologibegreber som "generativ AI", "artificial intelligence")
3. Kvalitet og interaktion: Fejl, verifikation og tillid i teknologianvendelsen (fx error*, fail*, verif*, trust*, reliab*)

Baseret på disse overordnede termer udformede jeg specifikke søgestrengene tilpasset de enkelte databasers søgefunktionalitet, som detaljeret fremgår af tabel 8 på bilag Bilag A: Søgestrategi. Særligt for ACM Digital Library og ERIC var søgningen specifikt målrettet "Be My Eyes" i kombination med synshandicap, mens søgningerne i Scopus og Google Scholar var bredere formuleret omkring AI-teknologi og fejlhåndtering.

Jeg har udvalgt både artikler fra anerkendte konferencer som CHI og ASSETS, som er relevante for HCI og tilgængelighed, og enkelte preprints fra 2024-2025. Jeg har valgt at tage preprints med, fordi udviklingen på området går så hurtigt – især efter de nye multimodale AI-modeller blev introduceret i 2023 – at nogle af de nyeste resultater endnu ikke er blevet

publiceret. De giver et opdateret billede af de seneste teknologiske fremskridt og brugerstrategier, men jeg har vurderet dem ekstra kritisk.

Et metodisk valg i denne litteraturgennemgang er fokuseringen på "Be My Eyes" som primær case for generativ AI i billede-til-tekst-teknologi for målgruppen. Dette valg afspejler en hermeneutisk tilgang, hvor teknologien undersøges i sin specifikke kontekst frem for gennem en komparativ analyse på tværs af platforme. Andre relevante teknologier såsom "Seeing AI" er ikke inkluderet som specifikke søgeord, hvilket repræsenterer en bevidst afgrænsning baseret på følgende overvejelser:

- Be My Eyes har gennemgået en særlig interessant transformation fra en menneske-til-menneske hjælpetjeneste til integration af generativ AI, hvilket udgør en unik kontekst for at studere brugerstrategier i overgangen mellem forskellige teknologiske tilgange.
- Et dybtgående fokus på én specifik teknologisk platform muliggør en mere nuanceret analyse af kontekstspecifikke brugspraksisser og fejlhåndteringsstrategier, i tråd med projektets sociotekniske perspektiv og fokus på situeret læring.
- Den opdeling mellem tekniske og sociotekniske tilgange, der præger forskningsfeltet, kalder på dybdegående casestudier, der kan bygge bro mellem teknologiske perspektiver og brugernes praksis og erfaringer.
- Be My AI skiller sig ud, fordi det tilbyder en hurtig og enkel interaktion uden krav om brugerens aktive valg af funktioner. Denne umiddelbare tilgang giver en særlig kontekst for at undersøge, hvordan brugerne håndterer fejl og udvikler tillid til teknologien. I modsætning til fx Seeing AI, hvor brugeren selv skal vælge funktioner og styre processen, gør Be My Eyes det lettere at få hurtige svar, hvilket kan påvirke både fejlhåndtering og tillid.

Dette fokus skal ikke ses som en vurdering af Be My AI's teknologiske overlegenhed, men snarere et metodisk valg om at prioritere dybde frem for bredde i undersøgelsen af fejlhåndteringsstrategier i deres specifikke kontekst.

ACM Digital Library-søgningen blev gennemført med en bredere søgestreng, uden eksplicit at inkludere fejl- og verifikationsdimensionen. Ved screeningen af resultaterne blev publikationer om fejlhåndtering, verifikation og tillid i relation til Be My Eyes og lignende teknologier for blinde brugere prioriteret.

Litteraturen blev udvalgt baseret på systematiske kriterier for at sikre både relevans og kvalitet. Inklusionskriterier omfattede publikationer fra 2022-2025 på engelsk, der behandler generativ AI for blinde brugere med specifikt fokus på fejlhåndtering, brugerstrategier og tillidsaspekter.

Efter screening af abstracts på tværs af databaserne blev 45 artikler udvalgt til dybdeanalyse: 28 fra ACM Digital Library, 15 fra Google Scholar og 2 fra Scopus. Min vurdering prioriterede peer-reviewed artikler fra etablerede konferencer inden for HCI og tilgængelighed (CHI, ASSETS, TACCESS). Preprints fra Google Scholar blev kun inkluderet, hvis de bidrog med væsentlige nyere indsigter og blev vurderet ekstra kritisk med hensyn til metodisk stringens.

Eksklusion omfattede studier der udelukkende fokuserede på tekniske aspekter uden brugerperspektiv, studier af andre brugergrupper uden relevans for synshandicap, samt ældre studier af klassisk billedgenkendelse uden forbindelse til generative AI-modeller.

Det er slående, at jeg ikke fandt nogen relevante resultater om Be My Eyes eller generativ AI for blinde i ERIC og PsycINFO, uanset om jeg brugte snævre eller brede søgeord. Det viser tydeligt, at der mangler forskning i de pædagogiske og psykologiske databaser. Mens der er en fremvoksende mængde litteratur om de tekniske aspekter af generativ AI for blinde, synes de kognitive, psykologiske og læringsmæssige dimensioner af disse teknologier at være systematisk underbelyste. På baggrund af dette fund vurderer jeg, at dette projekt har en klar relevans for at belyse et overset område. Jeg vil bruge en kombineret sociokognitiv og socioteknisk tilgang til at undersøge samspillet mellem bruger, teknologi og den konkrete brugskontekst. Mit særlige fokus bliver på den situerede læring, der udspiller sig i brugernes daglige møde med Be My AI.

Gennem analysen af bibliografien fra Google Scholar fremtræder en tydelig kronologisk udvikling i forskningen om generativ AI for blinde brugere fra 2022 til 2025:

- 2022 (2 publikationer): Grundlæggende studier af udfordringer og behov relateret til digitale enheder for blinde (Busaeed m.fl., 2022) samt tidlige undersøgelser af visuelle beskrivelsestjenester (Alharbi m.fl., 2022).
- 2023 (2 publikationer): Studier af, hvordan brugere tilpasser sig gennem kompenserende teknologidesign (Herskovitz m.fl., 2023) og bredere kortlægning af forskningsfeltet inden for mobilitetsteknologier (Madake m.fl., 2023).
- 2024-2025 (11 publikationer): En tydelig stigning i antallet af studier med fokus på generativ AI, fejlhåndteringsstrategier og multimodale modeller. Blandt de vigtigste er studier af fejlverifikation (Alharbi m.fl., 2024), mentale modeller af AI (Adnin & Das, 2024) og håndtering af objektgenkendelsesfejl (Hong & Kacorri, 2024).

Der sker et skifte fra brede analyser af behov og udfordringer til mere konkrete studier af, hvordan brugere navigerer i fejl og vurderer teknologien. Samtidig ses et skift fra klassisk billedgenkendelse til generativ AI og multimodale modeller. Desuden skifter fokus fra teknologifokus til studier, der belyser samspillet mellem teknologi, bruger og kontekst i praksis.

Sammenlagt viser analysen, at forskningen har bevæget sig fra generelle studier af behov og teknologiske udfordringer til mere konkrete undersøgelser af fejl, brugerstrategier og tillid. For at give et klart overblik over, hvad der er mest relevant for mit projekt, har jeg udvalgt en række studier, som jeg nu vil gennemgå.

3.1 Generativ AI i assisterende teknologi for blinde

Generativ AI har revolutioneret tilgængelighedsteknologier for blinde brugere gennem systemer som Seeing AI og Be My AI. Det er teknologier balancerer mellem at tilbyde øget selvhjulpethed og at introducere nye udfordringer relateret til systemfejl og misforståelser (Gonzalez Penuela m.fl., 2024).

Anvendelseskonteksterne er mange og spænder fra navigation og objektidentifikation til social interaktion, informationssøgning og æstetisk oplevelse (Gonzalez Penuela m.fl., 2024). Denne

kategorisering danner grundlag for at forstå variationer i brugererfaringer og strategier på tværs af forskellige situationer.

Brugerne udvikler typisk en pragmatisk tilgang til AI-teknologi. Adnin og Das (2024) dokumenterer, hvordan de balancerer mellem teknologiens potentielle omformning af praksis for selvstændighed og en kritisk bevidsthed om dens begrænsninger. Denne aktive forhandling af teknologiens rolle i hverdagen resonerer med Deweys pragmatiske filosofi, der betoner erfaring, handling og refleksion som centrale for menneskers interaktion med deres omgivelser (Dewey, 1938 [1988], s. 29-30).

Mens forskningsfeltet er voksende, forbliver kulturelle og geografiske aspekter tilsyneladende underbelyst. Studier fra Kina (Zhao m.fl., 2024), Bulgarien (Tzvetkova-Arsova & Tomova, 2024) og Indien (Kameswaran m.fl., 2024) indikerer, at både design og anvendelse er påvirket af lokale kontekster og præferencer. Der mangler forskning om danske brugeres erfaringer, som jeg derfor fokuserer på.

3.2 Fejlmekanismer og brugerstrategier

Fejl i AI-systemer er centrale for at forstå brugeres interaktionsstrategier. Litteraturen kategoriserer disse fejl i flere forskellige kategoriseringer, der supplerer hinanden. Bisante m.fl. (2023) identificerer hallucinationer (ubegrundet information), upræcise beskrivelser (mangelfuld detaljegrad) og systematiske bias (skævheder fra træningsdata). Alharbi m.fl. (2024) skelner mellem faktuelle fejl (forkert objektidentifikation), kontekstuelle fejl (misforståelse af relationer) og relevansbaserede fejl (fokus på irrelevante detaljer). Hong og Kacorri (2024) beskriver desuden fejl relateret til timing og systemets selvtillid.

Brugeres tillid til AI-baserede billedbeskrivelser er relativt lav (2,43 ud af 4) (Gonzalez Penuela m.fl., 2024), hvilket understreger betydningen af verifikationsstrategier. Bendel (2024) argumenterer for, at tillid til AI forudsætter både teknisk pålidelighed og etisk ansvarlighed, hvilket er særligt relevant for blinde brugere. Alharbi m.fl. (2024) dokumenterer fire primære strategier: verifikation gennem alternativ sansning, systemskift, social validering og kontekstbaseret vurdering. Disse strategier udvikles og tilpasses over tid gennem erfaring med specifikke fejlmønstre, hvilket afspejler en form for situeret læring (Hong & Kacorri, 2024; Suchman, 1987).

Et etisk lag tilføjes af Stangl m.fl. (2023), der undersøger privatlivsbekymringer som en dimension af fejlrisikovurdering. Blinde brugere balancerer mellem behovet for detaljerede beskrivelser og risikoen for privatlivskrænkelser, hvilket viser, hvordan vurderinger af, hvad der er socialt eller etisk acceptabelt, påvirker deres tillid til teknologien og deres strategivalg – særligt når teknologien fejler på måder, der kan opleves som krænkende eller stødende.

At brugere kan ses som aktive “hackere” og meddesignere fremhæves af Herskovitz m.fl. (2023). De viser, hvordan brugere modificerer teknologi(er) for at overkomme begrænsninger. Dette understøtter projektets pragmatiske teknologiforståelse, hvor teknologi betragtes som redskab for menneskelig praksis, der kontinuerligt tilpasses og omformes i anvendelse.

3.3 Sociotekniske og pragmatiske perspektiver på AI

Nyere litteratur anlægger en socioteknisk forståelse af AI-teknologi, hvor den betragtes som indlejret i sociale praksisser frem for som isolerede værktøjer. Alharbi m.fl. (2024) placerer begrebet *misfitting* centralt i denne forståelse, idet teknologiens begrænsninger anskues som udgangspunkt for brugerens aktive engagement og tilpasning. Dette korresponderer med Orlikowskis dualitetsbegreb, hvor teknologi både former og formes af menneskelig praksis (Orlikowski, 1992). I blinde brugeres interaktion med generativ AI manifesterer denne dualitet sig ved, at brugere både tilpasser sig systemets begrænsninger og aktivt omformer teknologiens anvendelse gennem alternative strategier og "hacks" (Herskovitz m.fl., 2023). Et eksempel på dette er, når blinde brugere kombinerer Be My AI med andre lydbaserede genkendelsesværktøjer for at krydstjekke AI'ens påstande om visuelle elementer eller når de bevidst omformulerer samme forespørgsel flere gange med forskellige vinklinger for at teste konsistensen i AI'ens beskrivelser.

Suchmans situerede handlingsteori underbygger dette perspektiv ved at betone, hvordan teknologiinteraktion altid er forankret i konkrete kontekster (Suchman, 1987). Dette understreges også af Hong og Kacorri (2024), hvor brugernes fejlhåndteringsstrategier er uløseligt forbundet med specifikke brugssituationer og tilegnet erfaring. J. G. Lee m.fl. (2024) beskriver en tre-faset tilpasningsproces (identifikation af begrænsninger, strategiudvikling og kontinuerlig justering), som har klare paralleller til Ejsing-Duun og Skovbjergs forståelse af, hvordan brugere via praktisk afprøvning og erfaringsopbygning lærer at integrere teknologien i deres hverdag på trods af designmæssige begrænsninger (Ejsing-Duun & Skovbjerg, 2019). Denne proces illustrerer, hvordan brugerens kontinuerlige vekselvirkning mellem handling, refleksion og tilpasning danner grundlag for udvikling af praktisk anvendelig viden om teknologien.

3.4 Mangler i den eksisterende forskning

Min gennemgang af forskningslitteraturen viser fem væsentlige områder, hvor der mangler viden:

1. Ingen langvarige studier af tilpasning til fejl, der dokumenterer udviklingen af strategier over tid
2. Manglende forståelse af tillid ved synshandicap, når brugere ikke kan verificere information visuelt
3. Underbelyst kreativitet i fejlhåndtering i relation til de forskellige fejltyper, AI-systemer producerer
4. Begrænset indsigt i Be My AI trods platformens udbredelse, særligt fraværende er forskning om kognitive og læringsmæssige aspekter
5. Mangel på danske/skandinaviske perspektiver i et felt domineret af nordamerikanske og asiatiske studier.

Mit projekt adresserer punkterne 2, 3, 4 og 5 gennem en undersøgelse af danske blinde brugeres strategier med Be My AI. Med udgangspunkt i et socioteknisk perspektiv analyserer jeg,

hvordan fejlhåndtering og kreative tilpasninger udvikles i samspillet mellem bruger, teknologi og kontekst. Projektets centrale antagelser er: 1) at brugere udvikler deres strategier gennem daglig praktisk erfaring, 2) at tillid til AI både handler om teknologisk præcision og brugerens personlige/etiske vurdering af fejlenes karakter, og 3) at brugere ikke er passive modtagere, men aktive medskabere af teknologiens værdi gennem kreative anvendelsesmåder. Da punkt 1 ville kræve anderledes ressourcer, adressere jeg det ikke i denne opgave.

Eksisterende studier beskriver brugernes fejlhåndteringsstrategier, men mangler indsigt i hvordan disse strategier faktisk udvikles i praksis. Her kan perspektiver fra situeret læring og distribueret kognition bidrage til at forstå, hvordan brugere gennem gentagen erfaring lærer at samarbejde med AI'en - fx ved at eksperimentere med kameravinkler, udvikle fornemmelse for systemets styrker og svagheder, og tilpasse deres interaktion til konkrete brugssituationer.

Videnskabsteoretisk afsæt

Mit videnskabsteoretisk syn har afsæt i en pragmatisk forståelse, der vægter handlingens og erfaringens centrale betydning for erkendelse. Jeg forstår viden som situeret, kontekstafhængig og handlingsorienteret, hvilket har direkte indflydelse på min metodiske tilgang og analytiske perspektiv.

Dette videnskabsteoretiske ståsted muliggør en undersøgelse af teknologien, som den opleves og anvendes i praksis, frem for at evaluere den ud fra abstrakte, forudbestemte kriterier. Min interesse ligger i at forstå, hvordan teknologien fungerer som et medierende redskab i konkrete situationer, og hvordan brugerne gennem deres praksis former teknologiens betydning og funktion.

4.1 Deweys pragmatisme og inquiry-baseret læring

I min teoretiske ramme trækker jeg særligt på John Deweys pragmatiske perspektiv, hvor teknologi forstås som redskaber, der medierer menneskets erfaring og handlemuligheder. Deweys begreb om "experience" - forstået som den kontinuerlige, transaktionelle proces mellem organisme og omgivelser - giver mig mulighed for at beskrive, hvordan blinde brugere gennem deres interaktion med AI-teknologi både påvirker og påvirkes af denne.

Jeg finder Deweys fokus på problemløsning som drivkraft for læring særligt relevant i min undersøgelse. Når brugerne møder fejl eller begrænsninger i AI-systemet, opstår der problemsituationer, der kræver udforskning og tilpasning. Gennem denne proces udvikler brugerne nye strategier og forståelser, der igen påvirker deres fremtidige interaktion med teknologien. Elkjær (2025, s. 9) formidler med henvisning til Dewey (1938 [1988], s. 29–30): "Vi lever altid i den tid, vi gør, og ikke i nogen anden. Kun ved til enhver tid at forstå betydningen af hver enkelt erfaring, vi gør os, kan vi forberede os på at gøre det samme i fremtiden. Det er den eneste forberedelse, der i længden dur til noget".

Dette perspektiv hjælper mig til at se brugernes fejlhåndteringsstrategier som udtryk for en aktiv, kreativ læringsproces frem for blot reaktioner på teknologiske mangler.

Dewey som grundlag for situeret læring og distribueret kognition

Når Dewey definerer erfaring som en kontinuerlig proces mellem organisme og omgivelser, fremhæver han et perspektiv, hvor læring og kognition ikke kan adskilles fra de konkrete kontekster og redskaber, de udfoldes gennem (Dewey, 1988, s. 47). Inspireret af Dewey ser jeg

på fejlsituationer som læringsmuligheder, hvor brugerne gennem praktisk erfaring udvikler nye strategier for at navigere i AI'ens begrænsninger.

I konteksten af blinde brugeres interaktion med Be My AI betyder dette, at AI-teknologien kan forstås som et redskab, der indgår i en distribueret kognitiv proces. Den "visuelle" erkendelse, som traditionelt ville være lokaliseret i den individuelle krop, fordeles nu på tværs af brugerens taktile og auditive sanser, telefonens kamera, AI'ens billedprocessering og de sproglige beskrivelser, der genereres. Som Hutchins (1996, s. 287–289) demonstrerer, kan kognitive processer spredes på tværs af individer, artefakter og kontekster, idet "learning is adaptive reorganization in a complex system" hvor "the proper unit of analysis is, thus, not bounded by the skin or the skull. It includes the socio-material environment of the person".

Hutchins' begreb om adaptiv reorganisering er centralt for min forståelse af brugernes fejlhåndteringsstrategier. Når en blind bruger oplever hallucinationer eller fejlbeskrivelser fra Be My AI, sker der ikke blot individuel problemløsning, men en systemisk reorganisering af det distribuerede kognitive system. Brugeren reorganiserer billedtagningsstrategier, justerer promptingstrategier og kalibrerer tillidsforhold - ikke som isolerede handlinger, men som koordinerede tilpasninger af hele systemet bestående af bruger, teknologi og kontekst.

Samtidig understøtter Deweys problemløsningsforståelse teorier om situeret læring. Suchman (1987, s. 49-50, 179) viser, hvordan handlinger altid er situerede i konkrete kontekster og opstår i mødet mellem brugerens intentioner og omgivelsernes affordances. Når en blind bruger eksperimenterer med forskellige billedvinkler efter at have oplevet fejlbeskrivelser, eller udvikler nye promptingstrategier baseret på AI'ens responsemønstre, sker der en form for læring, der er dybt forankret i den konkrete brugssituation. Denne læring opstår ikke gennem abstrakt instruktion, men gennem praktisk engagement med teknologien i specifikke kontekster - præcis som både Deweys pragmatisme og Suchmans situated actions foreskriver.

Konkretisering gennem modes of inquiry

For at kunne anvende Deweys pragmatisk perspektiv i min analyse på det konkrete plan, trækker jeg på Ejsing-Duun & Skovbjergs (2019, s. 445-460) begreb om "design som inquiry", specifikt deres tre modes of inquiry: practice, research og politics. Disse tre niveauer af undersøgelse resonerer med Deweys forståelse af inquiry som en multidimensionel proces, der både adresserer praktiske problemer, genererer viden og engagerer sig med værdimæssige spørgsmål.

- Practice inquiry svarer til det niveau, hvor brugerne gennem Schöns (2017) begreb *reflection-in-action* udvikler konkrete strategier for at forbedre deres interaktion med Be My AI. Dette inkluderer billedtagningsstrategier, promptingstrategier og fejlhåndteringsteknikker, der opstår gennem direkte, eksperimentelt engagement med teknologien.
- Research inquiry omfatter de processer, hvor brugernes interaktion med teknologien genererer viden, der rækker ud over den umiddelbare praksis. Dette inkluderer brugernes kortlægning af AI'ens kapabiliteter og begrænsninger, deres udvikling af mentale modeller for teknologiens funktionalitet, og deres forståelse af systematiske fejlmønstre.
- Political inquiry refererer til de kritiske refleksioner, som brugernes praksis rejser om teknologiens rolle i forhold til autonomi, tillid og social deltagelse. Dette niveau af inquiry

adresserer spørgsmål om, hvordan teknologien former brugernes livsvilkår og samfundsmæssige position.

Denne tredelte operationalisering af Deweys inquiry-begreb giver mig et analytisk framework til at forstå blinde brugeres strategier som multi-dimensionelle former for undersøgelse, der opererer på tværs af praktiske, vidensskabende og etiske dimensioner.

4.2 Sociotekniske perspektiver og situeret handling

Mit teoretiske perspektiv er også informeret af sociotekniske tilgange, der betoner det gensidige forhold mellem teknologi og social praksis. Jeg trækker her på Orlikowskis begreb om *duality of technology*, der beskriver, hvordan teknologien både former social praksis og samtidig formes gennem brugernes konkrete anvendelse, frem for at være bestemt af teknologiens design alene (Orlikowski, 1992, s. 405–406).

Dette perspektiv er centralt for min undersøgelse, da det giver mig mulighed for at forstå, hvordan blinde brugere gennem deres specifikke praksisser kan omfortolke og omforme AI-teknologiens rolle og betydning. Jeg er særligt interesseret i, hvordan brugerne udvikler det, Suchman (1987, s. 49–50, 179) kalder ”situated actions” - kontekstspecifikke handlinger, der opstår i mødet mellem brugerens intentioner, teknologiens affordances og den konkrete situation.

Suchmans indsigt i handlingers situerede natur danner grundlaget for min forståelse af situeret læring, idet jeg ser brugernes læring og udvikling af strategier med Be My AI som noget, der sker direkte i de konkrete brugssituationer og formes gennem praktisk erfaring. Det er viden og færdigheder, der er tæt knyttet til handling og den specifikke kontekst, teknologien anvendes i.

4.3 Disability studies og relationel handicapforståelse

Som teoretisk perspektiv inddrager jeg disability studies, der undersøger, hvordan handicap konstrueres i samspillet mellem individuelle kropsfunktioner og sociale strukturer. Særligt anvender jeg den relationelle model for handicap, som Shakespeare (2013) har udviklet.

Den relationelle handicapmodel betoner, at handicap opstår i dynamikken mellem individet og omgivelserne. Som Shakespeare (2013, s. 217) fremhæver: ”In practice, it is the interaction of individual bodies and social environments which produces disability”. Dette perspektiv anerkender både kroppens materialitet og de sociale barrierers betydning.

Hutchins’ begreb om adaptiv reorganisering harmoniserer med denne relationelle forståelse af handicap. Når en blind bruger reorganiserer det distribuerede kognitive system (bruger + AI + smartphone + kontekst), reduceres ikke et individuelt synstab - i stedet skabes der nye systemiske kapaciteter, der kan være funktionelt overlegne i visse kontekster. Der kan således skabes nye muligheder for at opfatte og interagere med verden, snarere end blot at forsøge at efterligne seendes måde at opfatte på.

Det relationelle handicapperspektiv hjælper mig til at forstå AI-teknologiens dobbeltsidede rolle: På den ene side kan den reducere handicappende situationer ved at give adgang til visuel information, men samtidig kan den potentielt skabe nye barrierer gennem fejl eller manglende

tilgængelighed. Som Shakespeare (2013, s. 220) påpeger: "While environments and services can and should be adapted wherever possible, there remains disadvantage associated with having many impairments which no amount of environmental change could entirely eliminate".

Jeg finder dette perspektiv særligt værdifuldt, da det flytter fokus fra individuel kompensation til en bredere forståelse af, hvordan teknologisk design kan bidrage til mere eller mindre handicappende omgivelser. Dette informerer min analyse af, hvordan Be My AI's design påvirker brugernes handlemuligheder og oplevelse af agens, og hvordan brugerne aktivt tilpasser teknologien gennem deres praksis.

4.4 Teoretisk integration

De tre teoretiske perspektiver - Deweys pragmatisme, sociotekniske tilgange med situeret læring og distribueret kognition, og disability studies - supplerer hinanden i min forståelse af blinde brugeres interaktion med generativ AI.

Med afsæt i Deweys pragmatisme undersøger jeg, hvordan brugerne gennem konkret problemløsning udvikler praktiske strategier til at håndtere teknologiens begrænsninger. Dette perspektiv hjælper mig til at se, hvordan bruger og teknologi gensidigt former hinanden gennem praktisk erfaring.

De sociotekniske tilgange med situeret læring og distribueret kognition operationaliserer dette pragmatisk perspektiv ved at vise, hvordan læring og kognition konkret foregår i samspillet mellem bruger, teknologi og kontekst. Hutchins' begreb om adaptiv reorganisering forklarer, hvordan strategier udvikles gennem praksis som systemiske tilpasninger snarere end individuelle kompensationer.

Disability studies-perspektivet placerer denne interaktion i en bredere samfundsmæssig kontekst ved at undersøge, hvordan teknologidesign kan reducere eller skabe handicappende barrierer. Dette perspektiv sikrer, at analysen ikke reducerer brugernes strategier til individuel tilpasning, men forstår dem som navigering i teknologisk medierede sociale strukturer.

Samlet giver denne teoretiske ramme mig mulighed for at analysere blinde brugeres strategier med Be My AI som komplekse fænomener, der opererer på tværs af individuelle, teknologiske og sociale dimensioner - forankret i konkret praksis, men med bredere implikationer for forståelsen af teknologi, læring og tilgængelighed.

Metode

Dette kapitel redegør for mine metodiske valg i undersøgelsen af blindes oplevelser og strategier med Be My AI. Undersøgelsen er forankret i det pragmatiske videnskabsteoretiske afsæt beskrevet i kapitel 4, hvor erfaring, handling og problemløsning ses som centrale for erkendelsen. Dette kommer til udtryk i valget af en tilpasset contextual inquiry-undersøgelse som metodisk ramme, hvilket muliggør observation af, hvordan blinde brugere gennem praktisk erfaring arbejder med og udvikler strategier i mødet med teknologien. Metoden giver samtidig indblik i, hvordan kognitive processer distribueres mellem bruger, teknologi og kontekst, i tråd med Hutchins (1996) teori om distribueret kognition. I det følgende redegør jeg for det metodiske design, udvælgelsen af informanter, dataindsamlingsmetoder og analysestrategi, der alle er informeret af undersøgelsens pragmatiske afsæt.

Metodens design er specifikt udviklet til at adressere undersøgelsens forskningsspørgsmål. Spørgsmålet om hvordan blinde brugere anvender generativ AI i deres dagligdag (forskningsspørgsmål a) belyses gennem det indledende interview og observationen af, hvordan informanterne navigerer mellem forskellige strategier i praksis. Spørgsmålene om fejltyper og deres påvirkning af brugeroplevelsen (forskningsspørgsmål b) samt brugernes håndtering af fejl og begrænsninger (forskningsspørgsmål c) adresseres direkte gennem de praktiske opgaver, der er designet med varierende kompleksitet og potentiale for fejl. Tillidsaspektet (forskningsspørgsmål d) undersøges både gennem observation af informanternes beslutninger i øjeblikket om at stole på eller verificere AI'ens output og gennem det afsluttende reflekterende interview. Endelig muliggør metodens fokus på danske brugeres konkrete praksis en sammenligning med internationale studier (forskningsspørgsmål e) ved at forsøge at dokumentere de specifikke strategier og oplevelser, der kendetegner den danske kontekst.

5.1 Metodisk design

Metodisk er undersøgelsen inspireret af Holtzblatt & Beyers (2017) contextual inquiry, en kontekstbaseret interviewform, der betragter informanten som en partner, der udfører sine vanlige rutiner, mens interviewerens veksler mellem at observere og drøfte, hvad der sker. Som Holtzblatt & Beyer fremhæver: "Because they talk about their activity while doing it, it's easy for users to see details that would otherwise be hard to discover" (2017, s. 47). Dette princip har jeg bevaret ved at lade informanterne interagere med Be My AI i konkrete interaktionsopgaver.

I dette projekt har jeg tilpasset contextual inquiry ved at bruge standardiserede opgaver frem for at lade informanterne gå spontant til værks i deres eget miljø. Tilpasningen har gjort det muligt at lave en mere systematisk sammenligning af strategier på tværs af informanter med deres forskellige erfaringsniveauer, samtidig med at den bevarer metodens fokus på observation af konkret praksis. Med denne tilpasning har min metode stræbt efter at bevare kerneprincipperne: 1) Jeg behandler informanten som ekspert på sin egen praksis, 2) jeg veksler mellem observationer og dialog med informanten om det der sker, og 3) jeg inviterer informanten til at deltage i fortolkning, så vedkommende aktivt bidrager til analysen af egne handlinger.

Dette er særligt relevant for blinde brugeres interaktion med AI-teknologi, hvor mange af strategierne er indlejret i selve interaktionen og derfor svære at artikulere uden for konteksten, som Holtzblatt & Beyer også påpeger: "Sometimes users cannot describe an activity at all when they are not doing it. Their memory of how to do the task is partially incorporated into the objects they use, so without those objects and the context of the activity, they really don't know what they do" (2017, s. 47). Elementer som kameraposition eller måden interaktionen med AI'en startes på kan være vanskelige at beskrive abstrakt. Som en del af den indledende samtale beskrev jeg derfor også for informanterne, at jeg undervejs, udover vores samtale i form af lyd, også ville indsamle billeder i situationer, hvor jeg fandt det relevant.

Teoretisk ramme

Denne tilpassede form af contextual inquiry harmonerer med det pragmatiske fokus på at observere handling i kontekst. Når informanterne engagerer sig i problemløsning og fx møder en fejlagtig AI-beskrivelse, kan jeg se, hvordan de gennem eksperimenteren udvikler eller afsøger nye strategier - helt i tråd med Deweys begreb om læring igennem problemløsning.

Samtidig gav metodens fokus på observation af brugere i praktisk interaktion med Be My AI mig mulighed for at undersøge distribueret kognition i praksis. Gennem dokumentation af brugernes fysiske håndtering af telefonen, deres sproglige interaktion med AI'en og systemets respons kunne jeg observere, hvordan den "visuelle erkendelse" blev fordelt mellem bruger, teknologi og miljø. Dette perspektiv på samspillet mellem menneske og teknologi bygger på Hutchins' teori om, at kognition strækker sig ud over den individuelle hjerne og manifesterer sig som adaptiv reorganisering af komplekse systemer (Hutchins, 1996, s. 287–289).

5.2 Udvælgelse af informanter

Undersøgelsen bygger på deltagelse af fem personer med blindhed, som alle har erfaring med brug af smartphone og assisterende teknologier. Den sjette informant er social blind¹. Informanterne blev strategisk udvalgt for at sikre en variation i alder, køn, teknologierfaring og grad af synshandicap, da disse faktorer potentielt kan påvirke interaktionen med generativ AI.

Rekrutteringen foregik gennem mit professionelle netværk ved Institut for Blinde og Svagsynede (IBOS), hvor jeg arbejder som konsulent. Dette gav mig adgang til informanter med

¹Social blindhed betyder at personen har synsfunktion, men i praksis fungerer som blind i sociale og hverdagssituationer på grund af synets begrænsede anvendelighed jf. den danske sociale klassifikation af synshandicap (ibos.dk, 2025).

forskellige baggrunde, men som alle havde det til fælles, at de anvendte digitale teknologier i deres hverdag. Kriteriet for inklusion i undersøgelsen var, at informanterne skulle have et synshandicap, der gjorde dem afhængige af ikke-visuelle teknologiske løsninger, samt at de havde grundlæggende erfaring med smartphone-brug - særligt skærmlæsere da dette er noget man typisk uddannes i at anvende. Se Bilag B: Skærmlæsere på s. 59 for en kort beskrivelse.

Det var ikke et intentionelt valg, at alle informanter var iPhone-brugere, men det viste sig, at samtlige deltagere anvendte iOS som platform. Mit undersøgelsesdesign var bygget op om, at informanterne anvendte deres eget hardware, da smartphones og særligt deres tilgængelighedsindstillinger udgør et meget personligt hjælpemiddel. De individuelle indstillinger, opsætninger og vaner, som brugerne har etableret løbende i deres hverdag, blev således bevaret under dataindsamlingen for at sikre en så naturlig interaktion som muligt med Be My AI.

De seks informanter repræsenterer forskellige grader af erfaring med Be My AI. To informanter (IP2 og IP3) havde omfattende erfaring med appen og anvendte den regelmæssigt. En informant (IP1) havde begrænset erfaring rent teknisk, men var fortrolig med Be My AI. En informant (IP4) var teknisk erfarer og havde en positiv teknologiindstilling. De sidste to informanter (IP5 og IP6) havde minimal eller ingen tidligere erfaring med Be My AI, men var fortrolige med andre billedbeskrivende teknologier som fx Seeing AI.

Denne spredning i erfaring var tilsigtet, da den gav mulighed for at observere både etablerede strategier hos erfarne brugere og den første tilegnelsesproces hos mindre erfarne brugere. Det er særligt i mødet med nye og udfordrende situationer, at situeret problemløsning og den adaptive reorganisering af praksis bliver tydelig og observerbar.

For at give et hurtigt overblik over de informanter, der indgår i undersøgelsen, præsenterer tabel 5.1 deres køn, alder, erfaring med Be My AI og generelle teknologibrug. Fælles for alle informanter er, at de alle er brugere af skærmlæseren VoiceOver, da de alle er iOS/iPhone-brugere.

Det relativt begrænsede antal informanter er i overensstemmelse med contextual inquiry's fokus på dybde frem for bredde. Denne tilgang er forankret i Holtzblatt & Beyers (Holtzblatt & Beyer, 2017, s. 47) principper om at observere brugere i deres autentiske kontekst for at opnå dybere indsigt i deres praksis. Dette dybdeperspektiv kræver en intensiv dataindsamlingsproces, hvilket naturligt begrænser antallet af informanter. Jeg lægger mig desuden op ad Brinkmann & Tanggaard (2022, s. 36) der peger på at forskningsinterviewet i studenterprojekter typisk omfatter tre til fem informanter.

5.3 Dataindsamling

Dataindsamlingen til denne undersøgelse blev gennemført i perioden marts-april 2025 gennem seks individuelle sessioner med informanterne. Interviewsessionerne varede mellem 42-72 minutter og bestod af tre hovedkomponenter. 1) Et indledende semistruktureret interview, 2) observation af praktisk interaktion med Be My AI, og 3) et afsluttende reflekterende interview. I praksis var interviewet ikke stramt styret, hvorved de enkelte interviews adskilte sig en smule i form og længde.

Tabel 5.1: Oversigt over informanter

Informant	Køn	Alder	Erfaring	Teknologibrug
IP1	Mand	45-55	Moderat	Praktisk orienteret teknologibruger med fokus på effektivitet. Forholder sig kritisk til nøjagtighed. Foretrækker menneskelig hjælp til komplekse opgaver.
IP2	Kvinde	25-35	Omfattende	Regelmæssig bruger af forskellige teknologier. Har udviklet strategier for kameraplacering og spørgsmålsformulering. Kombinerer punktskrift med AI og forstår hvornår menneskelig assistance er nødvendig.
IP3	Mand	60-70	Omfattende	Daglig bruger af Be My AI til praktiske formål. Teknologipositivt indstillet. Formuleringer præcise spørgsmål og eksperimenterer med avancerede anvendelser.
IP4	Mand	25-35	Moderat til omfattende	Teknisk avanceret bruger. Pragmatisk tilgang til fejl og foretrækker ofte AI frem for menneskelig assistance.
IP5	Mand	50-60	Minimal	Kritisk over for teknologiens pålidelighed. Reagerer på fejlfortolkninger. Foretrækker lavrisiko-anvendelser.
IP6	Kvinde	25-35	Minimal	Relativt nyblind og foretrækker at klare sig uden hjælpemidler når muligt. Har erfaring med Seeing AI. Finder detaljemængden i beskrivelser overvældende. Foretrækker menneskelig hjælp til opgaver.

Arrangement og de fysiske rammer

Interviewene blev afholdt på mit kontor på IBOS i Hellerup, da alle mine informanter har tilknytning til stedet. På bordet lå de objekter, som informanterne skulle interagere med via Be My AI-appen. For at skabe realistiske brugssituationer, arrangerede jeg nogle ”stationer” med forskellige typer genstande, der repræsenterede typiske anvendelsessituationer baseret på litteraturgennemgangen. Dette inkluderede:

1. Forskellige indpakkede fødevarer med varierende kompleksitet i emballage (Anthon Berg-chokolade formet som kanin, dato på mælkekarton i køleskabet)
2. Forskelligt trykt tekst på varierende materiale med varierende layout og typografi.

3. To forskellige typer panodil som let forveksles. Valgt pga. muligheden for at bruge punkt som reference.
4. Forskellige plakater og billeder ophængt på væggene

Dette arrangement var inspireret af Holtzblatt & Beyers Holtzblatt og Beyer (2017) princip om at observere brugere i kontekster, der udløser realistisk interaktion, samtidig med at det gav mulighed for at studere strategier på tværs af forskellige typer visuelle udfordringer.

Den trækker også på Sahib m.fl. (2013) og arbejdet med partecipatorisk design med blinde brugere, hvor de fremhæver værdien af at kombinere konkrete referencepunkter med dialog. I modsætning til deres primært tekstuelle scenarier, valgte jeg en mere taktil tilgang med fysiske stationer, der gav informanterne mulighed for at interagere direkte med objekterne, samtidig med at vi gennem dialog udforskede deres strategier og oplevelser med Be My AI. Denne kombination af det taktile og den mundtlige refleksion muliggjorde en dybere forståelse af, hvordan informanterne navigerede mellem fysisk objektforståelse og AI-genererede beskrivelser. Arrangementet var således et kompromis mellem naturalistisk observation og kontrolleret eksperiment.

De fire ”stationer” var forsøgt designet til at belyse centrale aspekter af mine forskningsspørgsmål, særligt spørgsmålene om fejltyper (b) og fejlhåndteringsstrategier (c). Fødevarestationen med emballage fremprovokerede forskellige fejltyper, mens tekststationen udfordrede AI’ens evne til at prioritere relevant information. Stationen med let forvekslelige medicintyper skabte situationer, hvor tillidsaspekter (forskningsspørgsmål d) blev særligt fremtrædende gennem potentielle konsekvenser af fejlidentifikation.

Arrangementet muliggjorde samtidig observation af distribueret kognition i praksis. Fx blev datoaflysningen på mælkekartonen distribueret mellem informantens taktile udforskning, kamerapositionering, AI’ens billedprocesering og den sproglige beskrivelse. I denne proces fungerede teknologien ikke blot som et passivt værktøj, men som en aktiv deltager i en kognitiv proces.

Valget af et kontrolleret arrangement frem for helt naturalistisk observation var en metodisk afvejning, der tjente flere formål. For det første muliggjorde den standardiserede opgaver en systematisk sammenligning af strategier på tværs af informanter med forskellige erfaringsniveauer. For det andet tillod den kontrollerede introduktion af potentielle fejlsituationer en fokuseret observation af fejlhåndteringsstrategier, som måske ikke ville opstå spontant under en kortere observation i naturlige omgivelser.

Samtidig er det vigtigt at anerkende begrænsningerne ved denne tilgang. Dette delvist kontrollerede arrangement kan potentielt påvirke brugernes naturlige interaktionsmønstre, og opgaverne repræsenterer ikke nødvendigvis det fulde spektrum af anvendelseskontekster. Disse begrænsninger adresseres delvist gennem den selvvalgte opgave, hvor informanterne selv kunne definere en relevant anvendelse, samt gennem det indledende interview, hvor tidligere erfaringer med teknologien i naturlige kontekster blev udforsket.



Figur 5.1: En informant anvender Be My AI til at aflæse information på en Thise mælkekarton. Billedet illustrerer den distribuerede kognitive proces, hvor bruger, objekt og teknologi indgår i et samspil. Kilde: Egen optagelse.

Dataindsamlingsmetoder

For at indfange de komplekse interaktioner mellem brugere og teknologi, anvendte jeg triangulering af flere datakilder:

1. Lydoptagelser: Alle sessioner blev optaget med informanternes informerede samtykke (se Bilag E: Samtykkeerklæring). Disse optagelser indfangede tre vigtige lydspor: informanternes verbalisering af tanker, min dialog med informanterne, og – særligt vigtigt – VoiceOver-output (skærmlæseren) og Be My AI's beskrivelser, der fungerede som "den tredje stemme" i interaktionen. Denne tredobbelte lydoptagelse gav mig mulighed for at dokumentere hele interaktionsforløbet, inklusive hvordan informanterne reagerede på AI'ens beskrivelser. Af hensyn til informanternes anonymitet og opgavens omfang er de fulde transskriptioner ikke vedlagt.
2. Feltnoter og verbal kontekstbeskrivelse: I tråd med god praksis i samvær med blinde personer verbaliserede jeg kontinuerligt mine handlinger og ændringer i det fysiske rum ("Jeg flytter lige denne her ind mod væggen" eller "Jeg lægger to æsker panodil foran dig"). Denne verbale kontekstbeskrivelse blev dermed dokumenteret i lydoptagelserne og suppleret med begrænsede feltnoter om non-verbale reaktioner, som jeg senere integrerede i transskriptionerne.
3. Fotografier: I udvalgte situationer, og med informanternes samtykke, tog jeg fotografier af interaktionssituationen for at dokumentere fysiske placeringsstrategier og kontekstuelle faktorer.

Kombinationen af disse datakilder gav et nuanceret billede af, hvordan brugerne interagerede med teknologien. Lydoptagelserne med den "tredje stemme" fra VoiceOver og Be My AI var særligt værdifulde, da de gav direkte indblik i den information, informanterne modtog, og hvordan de responderede på den. For eksempel kunne jeg høre, hvordan en informant justerede sin promptingstrategi efter at have modtaget en uklar beskrivelse fra AI'en.

Ved at dokumentere både informanternes sproglige udtryk, AI'ens beskrivelser og de fysiske aspekter af interaktionen, skabte jeg et datamateriale, der muliggjorde en analyse af, hvordan blinde brugere udvikler strategier i samspil med AI-teknologien. Dette afspejler en forståelse af, at interaktionen ikke kun handler om mentale processer, men udspiller sig i det konkrete samspil mellem bruger, teknologi og kontekst.

Praktisk gennemførelse

Hver interviewsession fulgte denne overordnede struktur:

1. Introduktion og indledende interview (20-30 min): Efter en gennemgang af forskningsetik og informeret samtykke, indledte jeg med et semistruktureret interview om informanternes baggrund, tidligere erfaringer med Be My AI og andre assisterende teknologier, samt generelle oplevelser med AI-beskrivelser.

2. Praktisk interaktion (40-60 min): Informanten fik en række opgaver, der involverede brug af Be My AI til at identificere og beskrive forskellige objekter. Opgaverne var designet til at variere i kompleksitet og potentiale for fejl. Under denne fase opfordrede jeg informanten til løbende at udtrykke sine tanker og beslutninger højt for at skabe indblik i deres refleksioner og interaktionsstrategier.

De praktiske opgaver fulgte en grundstruktur, der startede med enkle objektbeskrivelser og gradvist blev mere komplekse. Jeg tilpassede dog rækkefølge og fokus undervejs, baseret på den enkelte informants reaktioner og erfaringsniveau. Opgavetyper i forløbet omfattede:

- Beskrivelse af et velkendt objekt (fx en æske chokolade)
- Identifikation af information på emballage (fx ingredienser eller udløbsdato)
- Aflæsning af tekst på dokumenter eller etiketter
- Beskrivelse af visuelle elementer (plakater eller billeder)
- En selvvalgt opgave, hvor informanterne selv bestemte, hvad de ville undersøge
- En beskrivelse af et fotografi der kræver kontekstuel viden om stedet man befinder sig, for at afkode det
- Finde en udløbsdato på en mælkekarton

3. Afsluttende reflekterende interview (20-30 min): Sessionen afsluttedes med et reflekterende interview, hvor informanterne blev bedt om at vurdere oplevelsen, diskutere strategier og reflektere over tillidsaspekter.

Jeg anvendte en adaptiv tilgang til dataindsamlingen, hvor spørgsmål og opgaver blev justeret baseret på informanternes svar og interaktionsmønstre. Dette er i overensstemmelse med Holtzblatt & Beyers anbefalinger om at følge informantens naturlige arbejdsgang og fokusere på de aspekter, der fremstår som mest relevante i situationen.

Forskerposition og etiske overvejelser

Min dobbeltrolle som både forsker og IKT-konsulent kræver særlige metodiske og etiske refleksioner. På den ene side giver min professionelle baggrund mig en dybere forståelse af feltet og adgang til relevante informanter. På den anden side indebærer denne nærhed en risiko for forudindtagethed og påvirkning af informanternes svar. For at adressere dette potentielle bias forsøgte jeg at implementere flere metodiske greb:

- 1) Jeg ekspliciterede min forskerrolle ved sessionens begyndelse og understregede, at undersøgelsen ikke var en evaluering af informanternes kompetencer, men en udforskning af deres strategier og oplevelser.

- 2) I interviewguiden inkluderede jeg bevidst åbne spørgsmål, der inviterede til kritiske perspektiver på teknologien, for at modvirke en potentiel positiv bias.

- 3) Under dataanalysen har jeg været særligt opmærksom på at identificere både positive og negative aspekter af brugernes interaktion med Be My AI.

I praksis betød min tilgang, at jeg under interviews var eksplicit om min dobbeltrolle, fx ved at understrege, at min interesse lå i oplevelsen som bruger. Denne åbenhed om forskerposition

er i tråd med Deweys pragmatisk tilgang, hvor viden skabes gennem åben udforskning snarere end gennem fastlåste kategorier (Elkjær, 2025).

Metodiske justeringer i processen

I tråd med det kvalitative forskningsinterview gennemgik undersøgelsens metodiske design en naturlig udvikling fra planlægning til udførelse (Brinkmann & Lene Tanggaard, 2022). Undervejs i gennemførslen af mine interviews, foretog jeg nogle ændringer. Den planlagte frie observation gled ret hurtigt over i de formaliserede strukturerede "stationer". Dette skete som en naturlig konsekvens af at mangle noget "kontekst at tale ud fra". Det var derfor nærliggende efter lidt indledende indstillinger på informanternes mobiltelefoner (langsommere talehastighed, højere lyd mv.) at gribe ud efter de objekter der var inden for rækkevidde. Hermed var første "station" i gang. En konsekvens af det, muliggjorde dermed i endnu højere grad den systematiske sammenligning på tværs af informanter. Interviewsessionernes varighed viste sig at blive for lange, da VoiceOver-feedback også tager lang tid, da man ofte gerne vil lytte flere gange. Interviewsessionerne blev optimeret fra de oprindeligt planlagte 100-125 minutter, hvilket betød, at det korteste interview varede 42 minutter og det længste 72 minutter. Denne justering reducerede så omvendt informanternes kognitive belastning og resulterede i mere fokuserede data, i tråd med Suchmans 1987, s. 179 teori om situeret handling og dynamisk tilpasning i konteksten. Endelig førte refleksioner fra de første interviews til større opmærksomhed på billedtagningsstrategier i de senere interviews, hvor jeg derfor tilføjede specifikke spørgsmål om kameraposition og objektmanipulation. Samtidig blev balancen justeret mod flere strukturerede opgaver og færre elementer af fri udforskning, da dette gav øget indsigt i informanternes strategier. Erfaringerne fra de første interviews påvirkede min tilgang til de efterfølgende og skærpede mit blik for de konkrete strategier, som informanterne anvendte - strategier der kun blev synlige, når jeg observerede den faktiske interaktion med Be My AI. På denne måde formede den løbende tilpasning af metoden også min efterfølgende analyse af datamaterialet.



Figur 5.2: Reflection-in-action. Eksempel på objektmanipulation, miljøkontrol og komposition som implicit intentionskommunikation. Billedet illustrerer den strategiske positionering af en mælkekarton med henblik på at "indkode" intentionen om at finde holdbarhedsdato i selve billedet. I den konkrete situation kæmper informanten med at få oplysninger om hvorvidt det overhovedet er mælk. Kilde: Egen optagelse.

5.4 Analysestrategi

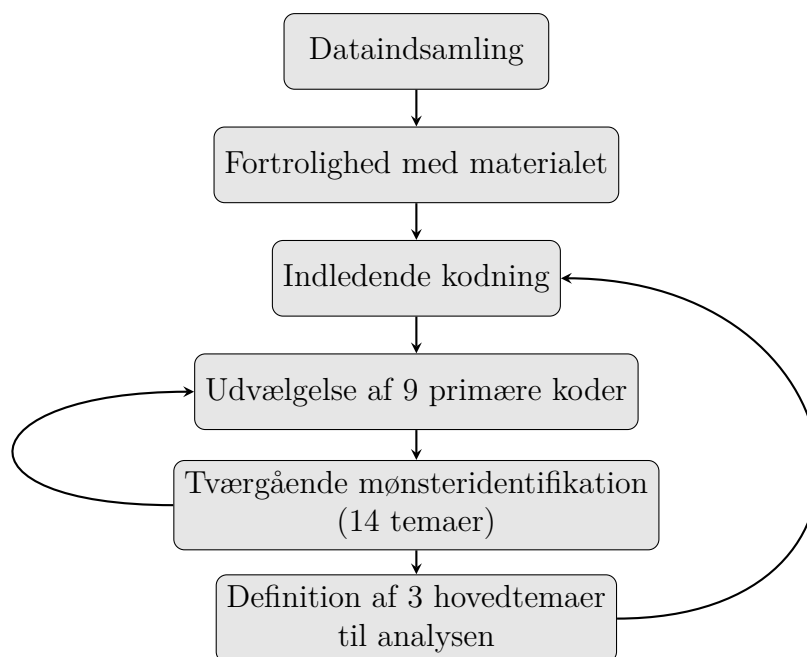
I min kodningsproces anvendte jeg Ejlsing-Duun & Skovbjergs 2019 tre modes of inquiry - practice, research og politics - som analytisk ramme til at identificere forskellige dimensioner af brugernes interaktion med teknologien. Dette gjorde jeg ved at stille specifikke spørgsmål til materialet, der fokuserede på henholdsvis praktiske strategier, systematisk vidensopbygning og værdimæssige overvejelser. Gennem denne proces identificerede jeg tre overordnede temaer, der danner grundlag for analysen.

Analysen fulgte en iterativ, tematisk tilgang inspireret af både Holtzblatt & Beyers (2017) contextual design-principper og Deweys (1988) begrebsapparat omkring inquiry. Processen bestod af tre hovedfaser: (1) udarbejdelse af individuelle brugerprofiler baseret på det indsamlede datamateriale, (2) tværgående mønsteridentifikation gennem en adapteret version af affinitetsdiagrammer, og (3) teoretisk fortolkning med udgangspunkt i de tre modes of inquiry fra Ejlsing-Duun & Skovbjerg (2019).

Min kodningsproces fulgte systematisk Braun & Clarkes (2006) seks-trins model:

1. Jeg gjorde mig fortrolig med materialet gennem gentagne gennemlæsninger og lytninger til lydoptagelser.
2. I den indledende kodning udviklede jeg et kodningsskema med seks kolonner: citat, primær kode, begrundelse for koden, ekstra observationer, sekundær kode, og tværgående tema. Dette muliggjorde en systematisk dokumentation af både manifesterede handlinger og latente betydninger i materialet.
3. Gennem kodning identificerede jeg 23 primære koder. Af disse udvalgte jeg de mest fremtrædende til videre analyse baseret på deres hyppighed og relevans for forskningsspørgsmålene. De centrale koder omfatter "promptingstrategier", "billedtagningsstrategier", "fejlbeskrivelser/hallucinationer", "tillid og verifikationsstrategier", "risikohåndtering", "begrænsninger i appens design" og "hensigt". Særligt fremtrædende var også de emotionelle reaktioner, der spændte fra "positivt overrasket" til "negativ oplevelse". Se tabel 9.2 på Bilag D: Kodningsoversigt.
4. Ved gennemgangen af koderne undersøgte jeg mønstre på tværs af de seks informanter, idet jeg fokuserede på både fællestræk og individuelle variationer. Dette førte til identifikation af 14 tværgående temaer som "visuel tilgang og billedforståelse", "promptstrategier og interaktion", "sikkerhed og kritisk anvendelse" og "teknologisk progression".
5. Jeg definerede og navngav temaerne ved at identificere essensen af hvert tema og bestemme, hvilke aspekter af dataene hvert tema indfangede.
6. Jeg grupperede de tværgående temaer til tre overordnede hovedtemaer, der strukturerer analysens opbygning og præsentation. Disse hovedtemaer præsenteres i begyndelsen af analysekapitlet.

For at give et visuelt overblik over kodningsprocessens faser præsenteres figur 5.3.



Figur 5.3: Systematisk kodningsproces fra rådata til hovedtemaer

I praksis bevægede jeg mig mellem kodningsniveauerne, hvilket illustreres af de tilbagevendende pile i figur 5.3. Nye indsigter førte ofte til revurdering af tidligere koder, ligesom formuleringen af hovedtemaerne påvirkede min forståelse af de mellemliggende mønstre. Denne iterative proces sikrede en vedvarende forbindelse mellem det datamateriale og de analytiske kategorier.

I min bevægelse fra koder til temaer fulgte jeg en systematisk proces, hvor jeg grupperede beslægtede koder og undersøgte deres indbyrdes relationer. Fx bemærkede jeg, hvordan visse koder relateret til fysisk interaktion med teknologien naturligt grupperede sig, mens andre koder vedrørende vurdering af information dannede et separat mønster. Gennem iterativ analyse af disse grupperinger identificerede jeg de tre hovedtemaer, der præsenteres i analysekapitlet.

For at operationalisere Ejlsing-Duun & Skovbjergs (2019) tre modes of inquiry i min analyse, stillede jeg følgende spørgsmål til materialet:

- **Practice inquiry:** "Hvilke konkrete strategier udvikler brugerne i mødet med teknologien?". Dette afdækkede både fysiske billedtagningsstrategier (f.eks. kamerapositioner og objektmanipulation), interaktionsstrategier (prompt-formulering og feedback) og verifikationsstrategier (taktil verificering og kontekstvalidering).
- **Research inquiry:** "Hvordan opbygger brugerne systematisk viden om AI'ens muligheder og begrænsninger?". Dette belyste, hvordan brugerne gennem praktisk erfaring kortlægger teknologiens præcision ved forskellige opgaver og udvikler mentale modeller for, hvornår og hvordan teknologien kan anvendes.
- **Political inquiry:** "Hvilke værdimæssige overvejelser gør brugerne sig om teknologiens rolle i forhold til selvhjulpethed, tillid og social deltagelse?". Dette fokus afdækkede brugernes afvejninger mellem teknologisk og menneskelig assistance, samt deres overvejelser omkring risici, privatliv og teknologiafhængighed.

I min analyse vekslede jeg mellem to tilgange: en induktiv, hvor informanternes egne ord og handlinger styrede, hvad jeg så efter i materialet, og en deduktiv, hvor jeg anvendte mine teoretiske begreber om fejltyper, brugerstrategier og tillidsforhold til at strukturere deres erfaringer.

Denne vekslen mellem induktion og deduktion afspejler mit pragmatiske videnskabsteoretiske afsæt, hvor teori og praksis gensidigt informerer hinanden.

På denne måde kunne jeg både være tro mod informanternes konkrete, situerede oplevelser og samtidig identificere, hvordan disse erfaringer manifesterer mønstre af distribueret kognition og adaptive strategier i interaktionen med teknologien.

Gennem kodningsprocessen identificerede jeg tre overordnede hovedtemaer, der danner grundlag for den efterfølgende analyse:

1. **Tillidsforhold og risikohåndtering i fravær af visuel verifikation**, der belyser, hvordan brugerne navigerer i tillidsrelationen til teknologien, når de ikke selv kan verificere information visuelt.
2. **Brugerstrategier for intentionskommunikation**, der undersøger, hvordan brugerne gennem praksis udvikler måder at kommunikere deres formål til systemet.
3. **Systemets muligheder og begrænsninger i hverdagspraksis**, der fokuserer på, hvordan Be My AI's funktionalitet og fejltyper manifesterer sig i brugernes konkrete anvendelsessituationer.

5.5 Metodiske begrænsninger

Trods grundig metodisk planlægning har undersøgelsen flere begrænsninger, der må tages i betragtning ved vurdering af resultaterne.

Informantudvælgelsens begrænsninger

Med seks informanter ($n=6$) er undersøgelsen naturligt begrænset i sin demografiske bredde. Selvom informanterne varierer i alder, køn og teknologierfaring, repræsenterer de ikke den fulde diversitet af blinde brugere. Fx er ældre brugere (70+) og personer med senere erhvervet blindhed underrepræsenterede. IBOS' målgruppe er defineret som unge og voksne (over 14 år) med et synshandicap (blindhed og alvorlig synsnedsættelse). I det daglige er aldersspredningen derfor i praksis 16+ til og med pensionsalderen. Mit fokus har derfor ikke været på brugere over 70 år, da disse ikke kunne rekrutteres via mit eget netværk.

Netop rekrutteringen gennem mit professionelle netværk ved IBOS kan have resulteret i en gruppe af informanter, der er mere fortrolige med teknologi og mere åbne over for nye løsninger end blinde generelt. Jeg har søgt efter frivillige i mit eget netværk ved blot at spørge - og har har på den måde med stor sandsynlighed tiltrukket personer med særlig interesse for eller meninger om teknologien, hvilket potentielt kan skabe bias i dataene.

Rammebaserede begrænsninger

Selve dataindsamlingen er inspireret af contextual inquiry, men den foregik i et konstrueret arrangement, der ikke fuldt ud afspejler den naturlige anvendelseskontekst. Dette kan have påvirket informanternes interaktionsmønstre og strategier. Hertil kan min tilstedeværelse som forsker have påvirket informanternes adfærd og udsagn. Særligt kan informanterne have følt sig tilskyndet til at demonstrere kompetence eller udtale sig mere positivt om teknologien end de ellers ville have gjort - på trods af opfordringer om både at nævne det negative og det positive. Interviewsessionerne var tidsmæssigt begrænsede, hvilket ikke gav mulighed for at observere, hvordan strategier udvikles over længere tid.

Tekniske begrænsninger

Undersøgelsen fokuserer på en specifik version af Be My AI (foråret 2025), og resultaterne afspejler derfor denne specifikke versions styrker og svagheder. Hurtig teknologisk udvikling kan betyde, at nogle af de identificerede udfordringer allerede er adresseret i nyere versioner. Dog skal man her skelne imellem selve appens design og den sprogmodel der driver den. Disse er adskilt, da AI-modellen frit kan ændres uden appen behøver at undergå udvikling. Dermed kan der potentielt ske ændringer i Be My AIs evne til at genererer beskrivelser uden brugeres kendskab.

Som forsker har jeg ikke indsigt i de præcise systemprompts der driver henholdsvis GPT-4v (OpenAI) samt den systemprompt Be My AI (Be My Eyes) styres efter. Dette begrænser min mulighed for at forklare visse typer fejl eller inkonsistenser. Under interviewsessionerne oplevedes varierende kvalitet af internetforbindelse, hvilket kan have påvirket AI'ens svartid. Dette kan dermed potentielt have haft en indflydelse på informanternes oplevelse og dermed deres vurdering af teknologien.

Analytiske begrænsninger

Min position som seende forsker, der undersøger blindes teknologibrug, indebærer en fundamental erkendelsesmæssig asymmetri. Selvom jeg gennem min faglige baggrund har kendskab til målgruppen, mangler jeg den kropslig erfaring ift. blindhed, hvilket begrænser min forståelse af visse aspekter af informanternes oplevelser. Som kvalitativ undersøgelse sigter projektet ikke mod statistisk generalisering, men mod teoretisk forståelse og overførbare. Resultaterne bør derfor ikke ukritisk generaliseres til alle blinde brugere eller alle kontekster.

Disse begrænsninger understreger behovet for at se denne undersøgelse som et bidrag til en bredere samtale. Resultaterne bør suppleres med andre metodiske tilgange, inklusive kvantitative undersøgelser, længerevarende studier og partcipatoriske designprocesser, for at opnå en mere komplet forståelse af blinde brugeres interaktion med generativ AI.

5.6 Opsamling

Metodisk er undersøgelsen designet til at indfange interaktioner mellem blinde brugere og generativ AI gennem en tilpasset contextual inquiry-tilgang. Ved at kombinere semistrukturerede interviews, praktisk observation og reflekterende dialog skaber metoden et empirisk grundlag for at besvare undersøgelsens forskningsspørgsmål.

Den pragmatiske og sociotekniske forankring afspejles i mit fokus på konkret praksis og distribueret kognition, hvor brugernes strategier forstås som situeret problemløsning i et samspil mellem bruger, teknologi og kontekst. I det følgende kapitel præsenteres analysens resultater, struktureret omkring de tre identificerede hovedtemaer.

Analyse

Gennem kodningsprocessen, beskrevet i kapitel 5, har jeg identificeret tre overordnede temaer: ”Tillidsforhold og risikohåndtering i fravær af visuel verifikation”, ”Brugerstrategier for intentionskommunikation” og ”Systemets muligheder og begrænsninger i hverdagspraksis”.

Disse temaer belyser forskellige aspekter af min problemformulering om blinde brugeres oplevelser og strategier med Be My AI og adresserer forskningsspørgsmålene om brugsmønstre, fejltyper, håndteringsstrategier og tillidsaspekter.

Analysen er understøttet af den teoretiske ramme præsenteret i kapitel 4, hvor jeg ser på brugernes interaktion med Be My AI som en form for distribueret kognitivt system, hvor den ”visuelle” erkendelse fordeles mellem bruger, teknologi og kontekst.

6.1 Tillidsforhold og risikohåndtering

I dette afsnit analyserer jeg, hvordan blinde brugere navigerer i tillidsrelationen til teknologien, når de ikke kan verificere information visuelt. Dette skaber en grundlæggende informationsasymmetri, hvor brugeren er afhængig af systemets beskrivelser uden mulighed for selv at bekræfte eller afkræfte disse gennem direkte visuel observation. Med udgangspunkt i begrebsafklaringen af tillidsformer i fravær af visuel verifikation i afsnit 2.4 undersøger jeg, hvordan de tre former for tillid - funktionel, kontekstuel og konsekvensbaseret tillid - giver sig til kende i brugernes praksis. Analysen viser, hvordan brugerne gennem konkret erfaring udvikler strategier til at håndtere den grundlæggende informationsasymmetri.

Funktionel tillid

Funktionel tillid handler om brugerens vurdering af systemets grundlæggende tekniske kvalitet. IP1 udviser en grundlæggende skepsis: ”Det er jo altid i forhold til, jeg er altid en lille smule skeptisk i forhold til kunstig intelligens”. Denne skepsis uddybes: ”Det irriterer mig, at netop, hvis den siger noget, som lyder så sikkert... og så erfaringsmæssigt, at det kan være noget sludder”.

Dette udgangspunkt afspejler, hvad Hutchins (1996, s. 287–289) beskriver som brugerens aktive rolle i det distribuerede kognitive system, hvor brugeren må vurdere outputtets pålidelighed i fraværet af visuel verifikation (se 4.1). Den distribuerede kognition kompliceres, når AI fremstår skråsikker, men gennem erfaring har vist sig upålidelig.

IP4 indtager en anden position: "Det gør jeg jo i de fleste tilfælde. Jeg stoler på den". Denne tillid er dog betinget: "Det er da meget dejligt, at den sagde, at der var tal, den ikke kunne læse" Dette eksempel illustrerer, hvordan systemets transparens om egne begrænsninger styrker tilliden, hvilket fra et distribueret kognitivt perspektiv gør det samlede system mere robust. Dette reflekterer hvad Ejsing-Duun & Skovbjerg (2019) kalder *research inquiry*. Her opbygger brugeren systematisk viden om systemets kapabiliteter/funktionalitet.

Dette peger på, at brugerne indtager en afvejet position i forhold til funktionel tillid, hvor vurderingen af systemets tekniske pålidelighed udvikles gennem konkrete erfaringer med teknologien. IP4's betingede tillid og IP1's grundlæggende skepsis illustrerer, hvordan denne funktionelle tillid kalibreres gennem praktisk erfaring, helt i tråd med Deweys forståelse af erfaring som grundlag for viden og brugerens rolle i det distribuerede kognitive system.

Kontekstuel tillid

Kontekstuel tillid vedrører brugerens vurdering af, hvorvidt AI'en forstår deres intention og leverer relevant information. Dette tillidsaspekt bliver særligt udfordret i *Be My AI's* "billede først - intention bagefter"-design.

IP1 nævner denne udfordring: "Jeg har været i nogle situationer, hvor jeg ved præcis, hvad det er for en information, jeg søger. Og så er det irriterende, at det er det, den ikke vil give mig". Set gennem Suchmans (1987) begreb om situeret handling kan vi forstå dette som en fundamental udfordring i interaktionsdesignet. Suchmans teori fremhæver, hvordan handlinger altid er situerede i konkrete kontekster og opstår i mødet mellem brugerens intentioner og omgivelsernes affordances. Frustrationen opstår, fordi interaktionsmodellen skaber en diskrepans mellem den situerede handling og systemets designede flow.

IP2 beskriver en kompenserende strategi: "Fordi jo mere konkret jeg er, jo bedre er den i stand til at forstå, hvad det er, jeg vil have den til" Dette illustrerer Orlikowskis (1992) pointe om teknologiens dualitet, hvor brugeren aktivt tilpasser sin praksis til teknologiens begrænsninger, samtidig med at denne praksis omformer teknologiens anvendelse. Gennem situeret læring har brugeren udviklet specifikke, meget tydelige formuleringsstrategier.

IP5 udtaler frustrationen over irrelevant information: "Det er fint nok med en flot beskrivelse, men det er ikke det, jeg har brug for". Og senere: "Den har jo ikke nogen idé om, hvad jeg egentlig leder efter" Dette indrammer den grundlæggende udfordring ved kontekstuel tillid og illustrerer hvad Ejsing-Duun & Skovbjerg ville betegne som *practice inquiry* - brugerens praktiske udfordring med at få systemet til at levere situationsrelevant information.

Analysen viser, hvordan brugerne aktivt må arbejde med at kommunikere deres intention. Dette illustrerer både Suchmans pointe om interaktionens situerede natur og Orlikowskis forståelse af, hvordan brugere gennem praksis både tilpasser sig til og omformer teknologien.

Konsekvenstillid

Konsekvenstillid involverer brugerens vurdering af potentielle risici ved at stole på AI'ens output i forskellige situationer. Denne tillidsform er tæt knyttet til, hvilke konsekvenser fejlinformation kan have i specifikke kontekster.

IP2's risikovurdering illustrerer dette, da AI'en hævder at en almindelig minimælk er laktosefri uden af være spurgt konkret til dette: "Hvis jeg nu havde været laktoseintolerant, så havde jeg måske troet, at så var det sikkert for mig at drikke den her mælk". Senere understreger samme informant en klar grænse: "Så skal du ikke bruge den til at tage billeder af lyskrydset". Denne kategoriske afgrænsning demonstrerer, hvad Ejsing-Duun & Skovbjerg betegner som political inquiry - en værdibaseret vurdering af teknologiens rolle i forhold til personlig sikkerhed, sundhed og autonomi.

IP6 udtrykker lignende bekymring ved fejlinformation om panodil: "Det går jo faktisk hen og bliver en lille smule halvfarligt, vil jeg jo mene". Og om navigation: "Sådan noget med vejvisning og sådan noget... det vil jeg sgu heller ikke rigtigt turde". Dette viser, hvordan brugeren aktivt afgrænser teknologiens anvendelse baseret på risikovurdering¹.

Set gennem Shakespeares (2013, s. 323–334) relationelle handicapmodel ser vi, hvordan brugerne navigerer i spændingsfeltet mellem teknologiens muligheder og begrænsninger. Som Shakespeare påpeger, opstår handicap i interaktionen mellem individet og omgivelserne, og gennem strategisk risikovurdering minimerer brugerne situationer, hvor denne interaktion bliver problematisk. Dette skaber ligeledes en bro til Orlikowskis forståelse af, hvordan brugere aktivt former teknologiens rolle gennem deres praksis.

Tillid som adaptiv proces i fravær af visuel verifikation

På tværs af de tre tillidsformer viser analysen, at tillid ikke er en statisk tilstand, men en dynamisk, adaptiv proces, der udvikles gennem praktisk erfaring. Brugere befinder sig til stadighed i en situation hvor de må stole på information, de ikke selv kan verificere visuelt.

IP4 formulerer en pragmatisk position: "Man skal jo ikke stole blindt på noget eller nogen". Denne metaforiske formulering illustrerer Deweys (1988) forståelse af inquiry som en aktiv, undersøgende proces. Set gennem Hutchins distribuerede kognitionsperspektiv positionerer brugeren sig som den kritiske komponent i det kognitive system, der må evaluere og integrere information fra forskellige kilder.

Tillidsrelationen til Be My AI viser, hvad Ejsing-Duun & Skovbjerg kalder adaptive reorganisering - en kontinuerlig justering af praksis baseret på erfaring. Brugere udvikler gennem denne proces en nuanceret forståelse af, hvornår og hvordan teknologien kan anvendes. Dette afspejler både Orlikowskis dualitetsbegreb, hvor teknologi og praksis gensidigt former hinanden, og Suchmans forståelse af, hvordan handlinger altid er situerede i konkrete kontekster.

Samlet set bekræfter analysen af tillidsrelationerne, at blinde brugere tilgår teknologien pragmatisk - som et redskab der kontinuerligt vurderes, justeres og tilpasses i praksis.

Verifikationsstrategier og kalibreret tillid

I mødet med informationsasymmetrien udvikler brugerne konkrete verifikationsstrategier, der kan ses som eksempler på det Ejsing-Duun & Skovbjerg kalder *practice inquiry* – praktiske

¹Brug af Be My Eyes (inkl. Be My AI) forudsætter, at brugeren er mindst 18 år samt at brug sker på eget ansvar. Tjenesten må ikke erstatte mobility-hjælpe midler, medicinsk rådgivning eller nødhjælp. Brugeren licenserer alt indhold til Be My Eyes, og deling af personfølsomme oplysninger er ikke tilladt jf. Terms of Service (Be My Eyes, 2025)

løsninger udviklet gennem erfaring.

En almindelig strategi er at stille opfølgende spørgsmål til systemet. IP4 gør dette direkte, da AI'en fejltolker en flæsketeg som et franskbrød: "Så spurgte jeg den, er du sikker? Nej, undskyld, det er en flæsketeg". Her udfordrer brugeren aktivt systemet og dermed tester dets troværdighed.

Brugerne trækker også på deres almene viden til at vurdere rimeligheden i AI'ens påstande. IP1 illustrerer dette med et konkret eksempel: "Kan det virkelig være rigtigt, at jeg skal betale 135.000 kr. i husleje denne måned". Denne kobling mellem systemets output og brugerens erfaringsverden er helt i tråd med Deweys forståelse af, hvordan vi skaber mening gennem erfaring.

Flere informanter nævner også, at de skifter mellem forskellige teknologier ved usikkerhed. IP4 beskriver det da samtalen handler om tekstlæsning og nogle tydelige fejl: "Jeg ville bare gå ud af Be My AI og gå ind i et OCR-program"². Dette viser, hvordan brugerne aktivt sammensætter forskellige værktøjer for at opnå den information, de har brug for.

Gennem interviewene tegner der sig et mønster, hvor brugerne kategoriserer situationer efter risiko:

- Ved **lav risiko** (f.eks. at identificere genstande eller få æstetiske beskrivelser) accepterer brugerne ofte AI'ens svar uden videre. IP4 siger det ligeud: "Det gør jeg jo i de fleste tilfælde. Jeg stoler på den".
- Ved **moderat risiko** (f.eks. at identificere madvarer eller læse dokumenter) bruger de typisk ekstra kontrolstrategier. IP1 forklarer: "[...] Man skal selv vurdere lidt... giver det her mening, kan det virkelig være rigtigt".
- Ved **høj risiko** (f.eks. medicin eller trafiksituationer) vælger de ofte helt andre løsninger. IP6 er ikke i tvivl: "Sådan noget med vejvisning og sådan noget... det vil jeg sgu heller ikke rigtig turde".

Denne opdeling viser, hvordan brugerne gennem erfaring opbygger systematisk viden om, hvor og hvordan teknologien kan bruges - hvad Ejsing-Duun & Skovbjerg ville kalde *research inquiry*. Det knytter an til Suchmans forståelse af, at teknologiens brugbarhed altid vurderes i specifikke kontekster.

Gennem disse strategier udvikler brugerne det, jeg kalder *kalibreret tillid* – en nuanceret forståelse af, hvornår systemet kan stoles på. IP3 formulerer det præcist: "Jamen tilliden er fin, men nu er jeg jo også sådan lidt skeptisk-agtig. Så jeg tager ikke alting for gode varer, vel". Dette viser en pragmatisk tilgang, hvor tillid ikke er et enten-eller, men en løbende justering baseret på erfaring. I formuleringen af begrebet "kalibreret tillid" låner jeg kraftigt fra Lee & See (2004).

Set fra Orlikowskis perspektiv om teknologiens dualitet viser disse strategier, hvordan brugerne aktivt omformer teknologiens anvendelse gennem deres daglige praksis. Samtidig illustrerer det, i tråd med Shakespeares relationelle handicapmodel, hvordan brugerne gennem aktiv handling reducerer de barrierer, der opstår i mødet mellem menneske og teknologi.

²OCR står for Optical Character Recognition (optisk tegngenkendelse), eller tekstgenkendelse på dansk.

6.2 Brugerstrategier for intentionskommunikation

I dette afsnit analyserer jeg, hvordan blinde brugere udvikler strategier til at kommunikere deres intentioner til et system, der fundamentalt er designet med en ”billede først - intention bagefter”-struktur. Begrebet opstår i forbindelse med kodningen af datamaterialet og favner det, at en bruger kommunikerer med AI'en med en bestemt intention (se afsnit 5.4).

Med afsæt i begreberne fra distribueret kognition og situeret læring undersøger jeg, hvordan brugerne gennem praktisk erfaring udvikler både proaktive billedtagingsstrategier og reaktive promptingstrategier. Analysen viser, hvordan disse strategier fungerer som kreative tilpasninger i et komplekst socioteknisk system, hvor brugeren må kompensere for designmæssige begrænsninger.

At navigere i ”billede først - intention bagefter”

Be My AI's grundlæggende interaktionsdesign, hvor brugeren må indlede med at tage et billede før intentionen kan udtrykkes, skaber en fundamental udfordring for blinde - hvilket jo er paradoksalt, da appens intenderede målgruppe netop er blinde.

Denne designbeslutning kan forstås gennem både Orlikowskis (1992) dualitetsbegreb og Suchmans (1987) begreb om situeret handling, hvor designet skaber en indbygget spænding mellem systemets prioriteter og brugerens faktiske intentioner.

IP2 bekræfter denne præmis, når hun fremhæver at billedet uomgængeligt er startpunktet for enhver interaktion. Konsekvensen er en grundlæggende kamp om fortolkningsretten, som flere informanter udtrykker frustration over. IP5: ”Det er fint nok med en flot beskrivelse, men det er ikke det, jeg har brug for” og ”Den har jo ikke nogen idé om, hvad jeg egentlig leder efter”. IP3 er ligeså direkte: ”Jeg er sgu ligeglad med, om bordet er rundt, eller lampen er moderne”.

IP2 identificerer kernen i problemet: ”Det er jo den, der til at starte med udvælger, hvad den synes er vigtigt”. Denne asymmetriske magtrelation forstærkes yderligere af informationsmængden, som informanterne finder overvældende.

IP6: ”Ved det første billede... Du får jo virkelig mange inputs” og IP3 fortsætter: ”Verbositeten er overvældende ofte”. Under interviewene observerede jeg gentagne gange, hvordan informanterne afbrød oplæsningen – et forsøg på at genvinde kontrol over informationsstrømmen. Den kognitive belastning bliver tydelig i IP4's beskrivelse af processen som en ”kortslutning” da han forsøger at dikterer et opfølgende spørgsmål og bytter rundt på ordene i sætningen: ”Det var derfor min hjerne kortsluttede lidt, da jeg skulle spørge den ind til [det] derude på billedet”.

Flere informanter efterspørger derfor muligheden for at kommunikere intention før billedtagning. IP4 adspurgt om muligheden: ”Jeg synes, det skal være en mulighed. Helt klart”. IP1: ”At sige, at nu tager jeg et billede, og det skal jeg bruge til... Altså, det du skal svare på, det er det her”.

Disse ønsker kan, set gennem Hutchins (1996) teori om distribueret kognition, forstås som forsøg på at omfordele den kognitive arbejdsbyrde mere hensigtsmæssigt mellem bruger og system. I fraværet af denne mulighed udvikler brugerne alternative strategier for at genvinde kontrollen, særligt gennem billedtagingsstrategier der forsøger at ”indkode” intentionen direkte i selve billedet. Det handler altså grundlæggende om kampen om hvem der besidder fortolkningsretten.

Billedtagningsstrategier som implicit intentionskommunikation

I fraværet af muligheden for et direkte valg af intention udvikler brugerne billedtagningsstrategier. Med Ejsing-Duun & Skovbjergs (2019) teoretiske ramme udgør disse strategier en form for practice inquiry. En praktisk, eksperimentel tilgang hvor brugeren gennem konkret handling udforsker og forbedrer interaktionen med teknologien. Samtidig illustrerer strategierne og Orlikowskis dualitetsbegreb og Hutchins' distribuerede kognition, hvor brugerne aktivt omformer teknologiens anvendelse gennem kreative tilpasninger.

Informanterne beskriver tre typer af strategier, der alle tjener til at "indkode" deres intention i selve billedet:

- Den første strategi involverer *bevidst kamerapositionering*, som IP1 beskriver: "Jeg tænkte egentlig, at jeg skulle sidde helt nonchalant og pege derhen, hvor cirka det er, eller at jeg skulle rejse mig og prøve at tage lidt oppefra, for at det blev lidt bedre". Her tilpasser brugeren sin kropslige position for at kompensere for teknologiens begrænsninger. Se figur 5.2.
- Den anden strategi fokuserer på *miljøkontrol og komposition*, som IP3 demonstrerer: "Jeg sørger for, at der er nogenlunde frit omkring det væsentlige". Denne bevidste manipulation af den fysiske kontekst er et tydeligt eksempel på, hvordan den kognitive proces fordeles mellem bruger, teknologi og omgivelser/miljøet.
- Særligt interessant er den tredje strategi – *objektmanipulation* – som IP2 beskriver: "Jeg roterer den lige en gang mere, fordi det var det der med, at holdbarhedsdatoen vi var interesseret i". Her ser vi, hvordan brugeren aktivt manipulerer det fysiske objekt for at fremhæve den relevante information og dermed implicit kommunikere sin intention til systemet.

Disse strategier viser, hvordan brugerne gennem praktisk erfaring udvikler et repertoire af teknikker, der kompenserer for designets iboende begrænsninger og skaber en mere effektiv interaktion på deres egne præmisser.

Promptingstrategier for intentionsafklaring

Efter den indledende billedbeskrivelse udvikler brugerne forskellige sproglige promptingteknikker gennem eksperimentering med formuleringer.

IP2: "Jo mere konkret jeg er, jo bedre er den i stand til at forstå, hvad det er, jeg vil have den til". Dette illustrerer sammenhængen mellem promptens præcision og kvaliteten af AI'ens respons.

En anden strategi involverer validerende prompts, som IP4 viser: "Så spurgte jeg den bare: er du sikker? Og så rettede den sig faktisk". Denne tilgang fungerer som verifikationsmekanisme, men viser sig problematisk, idet IP2 oplevede, at systemet, selv efter en rettelse, stadig fastholdt den oprindelige fejl: "Hvis jeg nu havde skrevet, at den er formet som en fodbold [...] så ville den stadigvæk have skrevet, tak for at rette mig".

IP5 understreger frustrationen: "Det ærgrer mig da. Især når man prøver at belære den. Så ærgrer det mig, at den så ikke lærer noget." Dette peger på manglende forståelse for teknologiens

begrænsninger. IP5 er selvreflekterende i citatet: ”Man skal jo selv lige træne sig selv i at stille det rigtige spørgsmål” – et centralt element i Deweys forståelse af inquiry-baseret læring.

En fjerde strategi handler om koncise kommunikation, som IP3 illustrerer: ”Jeg er relativt kort for hovedet. Også fordi ellers kan man skrive sig synder sammen”.

Disse promptingstrategier udgør en dialogisk kompetence, der udvikles gennem praktisk erfaring. De demonstrerer, hvordan brugerne er aktive medskabere af teknologiens anvendelse gennem kreativ tilpasning og kontinuerlig læring jf. Ejning-Duun & Skovbjergs begreb om research inquiry.

Hensigtsstyringens kompleksitet

Gennemgående er der i mit datamateriale et ønske fra informanternes side om mere direkte hensigtsstyring – altså at kunne styre systemets opmærksomhed, før billedet tages. Dette ønske afspejler en dybere værdibaseret refleksion over teknologiens rolle i forhold til brugerens autonomi. Dette kan forstås som det Ejning-Duun & Skovbjerg betegner som political inquiry.

IP1 artikulerer dette tydeligt: ”Så tænker jeg, at i en del tilfælde kunne det være rigtig rart. [...] At sige, at nu tager jeg et billede, og det skal jeg bruge til... Altså, det du skal svare på, det er det her”.

Samtidig fremkommer der i materialet en nuanceret forståelse af designets intentioner. IP3's overvejelse illustrerer denne kompleksitet: ”Ja, ja. Så må man spørge ind til det, ikke? Men det er jo også et... Altså meningen er, at det skal være enkelt at bruge, ikke”. Her balanceres mellem anerkendelsen af appdesignets mål om enkelhed og de praktiske begrænsninger, dette medfører.

Denne afvejning mellem brugervenlighed og kontrol kan forstås gennem Shakespeares (2013, s. 217) relationelle handicapmodel. De strategier brugerne udvikler repræsenterer ikke blot individuel tilpasning, men en aktiv forhandling af teknologiens rolle i en social kontekst. I dette perspektiv opstår ”handicappet” ikke alene fra den individuelle funktionsnedsættelse, men i mødet mellem brugerens behov og systemets design.

Hensigtsstyringsproblematikken illustrerer kernen i intentionskommunikation: når Be My AI forhindrer brugere i at udtrykke deres intention før billedtagning, skabes der ikke blot en teknisk begrænsning, men en handicappende barriere i Shakespeares forstand - en systemisk hindring for brugerens agens. Paradoksalt opstår der således en situation, hvor assisterende teknologi selv bliver kilde til nye barrierer. Be My AI's design, der prioriterer teknologisk enkelhed over brugerkontrol, skaber utilsigtet handicappende strukturer for netop de brugere, teknologien skulle støtte.

Samlet viser analysen af brugernes strategier for intentionskommunikation, hvordan de aktivt navigerer i et design, der ikke umiddelbart understøtter deres behov. Gennem praktisk erfaring udvikles et repertoire af teknikker – fra billedtagningsstrategier til promptformuleringer – der kompenserer for designets begrænsninger. Disse adaptive praksisser demonstrerer både Deweys forståelse af læring gennem problemløsning og Hutchins' begreb om reorganisering af komplekse kognitive systemer.

I det følgende afsnit vil jeg undersøge, hvordan disse strategier udspiller sig i mødet med systemets konkrete muligheder og begrænsninger i hverdagspraksis.

6.3 Muligheder og begrænsninger i praksis

I dette afsnit analyserer jeg, hvordan Be My AI's funktionalitet og begrænsninger giver sig til kende sig i praksis. Med udgangspunkt i de foregående afsnit, undersøger jeg, hvordan brugerne oplever mulighederne i Be My AI samt de begrænsninger de oplever i konkrete anvendelsessituationer.

Analysen viser, hvordan brugerne veksler mellem positiv indstilling og frustration, samt hvordan de kontinuerligt afvejer teknologiens værdi i forhold til deres behov. Dette kan forstås gennem Hutchins (1996) perspektiv på distribueret kognition.

Hallucinationer og fejlfortolkninger

Informanterne oplever forskellige fejltyper, hvor AI'en enten fejlidentificerer objekter eller beskriver ikke-eksisterende elementer. IP4 beskriver en markant fejlidentifikation: "Min største hallucination, den sagde, da den sagde min flæskesteg var et stykke franskbrød." Og fortsætter: "Så spurgte jeg den, er du sikker? Nej, undskyld, det er en flæskesteg." Dette viser systemets evne til selvkorrektion³, men afslører også den usikkerhed, som fejl skaber jf. afsnittet om promptingstrategier for intentionsafklaring (s. 38).

Særligt alvorligt bliver det i sikkerhedskritiske kontekster. IP6 oplevede, at medicin blev fejlbeskrevet: "Det var nogle Fishermans... Et eller andet... Jeg ved ikke, om det er noget lakrids..." Objektet var faktisk Panodil, hvilket IP6 reflekterer over: "Det går jo faktisk hen og bliver en lille smule halvfarligt, vil jeg jo mene." jf. afsnittet om konsekvenstillid (s. 34).

Disse oplevelser medfører ændringer i brugernes praksis. IP5 udtrykker dette direkte: "Men jeg vil dog være varsom. Jeg er mere varsom nu." Her ses Orlikowskis Orlikowski (1992) dualitetsbegreb ligeledes, da Be My AI aktivt omformer brugeradfærd.

Informationsmængde og relevans

Et gennemgående tema er paradokset, hvor systemet leverer for meget generel information samtidig med at mangle specifik, kontekstrelevant information. IP3 udtaler: "Verbositeten er overvældende ofte," og IP4 reagerer: "Hold da op, det var lidt meget information."

Samtidig oplever informanterne også det modsatte problem. IP5 bemærker ved gennemgang af en fødevarers ingredienser. IP5 siger efter jeg, som interviewer, har sammenholdt det Be My AI



Figur 6.1: Kamerapositionering over Panodil Zapp-æske. Billedet illustrerer den grundlæggende udfordring med præcis placering af kamera, når brugeren ikke har visuel feedback. Kilde: Egen optagelse.

³Systemet er selvkorrigerende inden for samme session. Det har ingenting lært. Se evt. afsnit 6.2.

beskriver med den faktiske ingrediensliste: ”Den nævner jo ikke alle ingredienser” – en potentiel risiko hvis man har fødevareallergi.

Teorien om situeret handling belyser her, hvordan disse problemer opstår når AI-systemets generelle formål ikke harmonerer med brugerens situerede informationsbehov. Kløften mellem generel beskrivelse og specifikke behov skaber kognitive udfordringer for brugeren (Suchman, 1987).

Teknologiens styrker

Trods udfordringerne beskriver informanterne også situationer, hvor Be My AI overgår forventningerne og skaber ny værdi i deres dagligdag. IP3 udtrykker entusiastisk: ”Kæmpe, kæmpe, kæmpe fordel,” og bekræfter at teknologien har gjort ham mere selvhjulp.

Flere informanter oplever, at AI'en opdager detaljer, som de selv havde overset. IP5 illustrerer dette: ”Hold da op. Det var da godt en lille cylindrisk [genstand]. Den havde jeg da fuldstændig overset.” Denne evne til at supplere brugerens egen sansning med systematisk billedanalyse repræsenterer en unik styrke ved teknologien, der rækker ud over traditionel beskrivelse.

IP3 fremhæver særligt værdien af privatliv og tidsmæssig fleksibilitet: ”Det er jo ikke alting, der rager andre mennesker, man spørger om... det kan jo være på alle mulige mystiske tidspunkter af døgnet.” Denne autonomi i informationssøgning udgør en væsentlig gevinst sammenlignet med at skulle bede andre om hjælp.

Teknologien muliggør også spontan adgang til visuel information i situationer, hvor alternative løsninger ikke er tilgængelige. Denne umiddelbare tilgængelighed skaber handlemuligheder, der ellers ikke ville eksistere for brugerne, og den understøtter brugernes ønske om at bevare kontrol over privat visuel information uden at involvere andre i deres personlige forhold (Stangl m.fl., 2023).

Set gennem Shakespeares 2013 relationelle handicapmodel kan disse positive anvendelser forstås som situationer, hvor teknologien effektivt reducerer handicappende barrierer ved at muliggøre selvstændig adgang til visuel information og dermed øge brugerens agens i hverdagssituationer.

Pragmatisk navigation i teknologiens muligheder og begrænsninger

Gennem analysen af brugernes praktiske interaktion med Be My AI fremtræder et mønster. Denne kontinuerete forhandling handler dermed ikke blot om at overkomme systemets begrænsninger. Brugerne udvikler aktive strategier for at øge teknologiens potentiale, samtidig med at de kompenserer for dens svagheder gennem kreativ kombination med andre værktøjer.

IP4 beskriver fx denne tilgang konkret: ”Jeg ville bare gå ud af Be My AI og gå ind i et OCR-program” og tilføjer: ”Der er jo en app, Zusanka. Den er jo skabt til at lede efter udløbsdatoer.” Dette illustrerer, hvordan practice inquiry hjælper til at kombinere forskellige værktøjer afhængigt af den konkrete opgave. Ud over denne adaptive værktøjsanvendelse engagerer brugerne sig aktivt i refleksion over teknologiens fremtidige udvikling.

IP3 foreslår konkrete designforbedringer: “Man kunne måske have nogle forskellige prompts, hvor man fokuserede på det, der var centralt i billedet,” og tilføjer: “Jeg synes måske, at det ville være rigtig dejligt, hvis vi fik noget hovedbårent kamera, så man fik en håndfri fidus [red., forstået som et praktisk hjælpemiddel, der muliggør håndfri brug af Be My AI, som ellers kræver, at brugeren holder telefonen i hånden]”. Dette skal ses i lyset af, at blinde ofte ”løber tør for hænder”, da der både skal jongleres med fx stok, mobiltelefon, dagligvare og måske en hund.

Disse forslag repræsenterer eksempler på political inquiry, hvor brugerne aktivt reflekterer over, hvordan teknologien bør udvikles for bedre at understøtte deres autonomi og handlefrihed.

Den overordnede praksis kan forstås som en kontinuerlig forhandling af teknologiens rolle i dagligdagen, hvor brugerne balancerer mellem nye muligheder og (frustrerende) begrænsninger.

IP4’s pragmatiske tilgang illustreres i udsagnet: ”Jeg er egentlig bare glad for, at den har beskrevet noget.” Dette afspejler en forståelse af teknologi som praktisk problemløsning – brugerne bedømmer Be My AI ud fra dens konkrete nytte i specifikke situationer.

Set gennem Hutchins’ perspektiv på distribueret kognition kan brugernes praksis forstås som adaptiv reorganisering af det kognitive system, hvor de udvikler strategier for at optimere fordelingen af kognitive opgaver mellem bruger, AI og supplerende teknologier. Denne reorganisering demonstrerer, hvordan læring i distribuerede systemer ikke blot handler om at tilpasse sig teknologiens begrænsninger, men om at omforme hele systemets funktionalitet gennem kreativ praksis.

6.4 Opsamling på analyse

Analysen af danske blinde brugeres erfaringer med Be My AI har afdækket tre centrale dimensioner af interaktionen, som tilsammen tegner et nuanceret billede af, hvordan generativ AI fungerer som assisterende teknologi i hverdagspraksis.

Tillidsanalysen viser, hvordan brugerne navigerer i en grundlæggende informationsasymmetri - de må stole på beskrivelser, de ikke selv kan verificere visuelt. Gennem praktisk erfaring udvikler de strategier for kalibreret tillid, der balancerer funktionel, kontekstuel og konsekvensbaseret vurdering af systemets output. Denne tillidsopbygning er en aktiv proces, hvor brugerne kontinuerligt justerer deres afhængighed af teknologien baseret på konkrete erfaringer med både fejl og succeser.

Intentionskommunikation og designparadokser Analysen af intentionskommunikation afdækker et fundamentalt designproblem: Be My AI’s struktur - ”billede først, intention bagefter” - skaber barrierer for netop den målgruppe, teknologien skulle assistere. Som respons udvikler brugerne kreative kompensationsstrategier gennem billedkomposition, objektmanipulation og promptingteknikker. Dette demonstrerer deres rolle som aktive medskabere af teknologien, men illustrerer også hvordan designvalg kan skabe nye former for ekskludering.

Undersøgelsen af teknologien i hverdagspraksis viser brugernes strategiske navigation mellem muligheder og begrænsninger. De fungerer som ”teknologihackere”, der adaptivt kombinerer forskellige værktøjer og reorganiserer deres distribuerede kognitive system. Dette demonstrerer både teknologiens transformative potentiale og de strukturelle udfordringer, når generelle AI-systemer møder specifikke brugerbehov.

Analysen peger på nogle paradokser:

Hvordan kan assisterende teknologi designes til at understøtte brugernes autonomi? Hvordan afstemmes generativ AI mellem generalitet og kontekstspecifik anvendelighed? Hvilke design sikrer integration af brugernes praktiske viden i teknologiudvikling?

Disse spørgsmål danner grundlag for diskussionen af projektets bredere implikationer for både teori og praksis inden for tilgængelighedsteknologi.

Diskussion

Denne undersøgelse af en lille gruppe danske blinde brugeres erfaringer med Be My AI viser nye perspektiver, der både supplerer og udfordrer den internationale forskning. Gennem analysen af de seks deltageres konkrete strategier kan jeg pege på tre hovedområder, hvor mine fund nuancere viden, samtidig med at jeg anerkender undersøgelsens begrænsede omfang.

I det følgende diskuterer jeg først, hvordan brugernes tillidsstrategier kan forstås som en form for kalibreret tillid, der udfordrer eksisterende forståelser af tillidsrelationer til AI-teknologi. Dernæst belyser jeg, hvordan Be My AI's designstruktur skaber en paradoksal ekskludering, der kræver kreative kompensationsstrategier fra brugernes side. Afslutningsvis analyserer jeg, hvordan brugernes teknologihacking-praksis peger på en adaptiv ekspertise, der har både teoretiske og praktiske implikationer. Gennem disse tre temaer adresserer jeg systematisk mine forskningsspørgsmål om brugsmønstre, fejltyper, håndteringsstrategier og tillidsaspekter.

7.1 Kalibreret tillid som udfordring til eksisterende tillidsforståelser

Mine resultater udfordrer antagelser i den internationale litteratur om tillid ved synshandicap i forbindelse med AI. Eksisterende studier beskriver primært tillid som binært, hvor brugere enten stoler på AI-systemet eller udviser generel skepsis. Dette afspejles i målinger som Gonzalez Penuela m.fl. (2024) fund af relativt lav tillid til AI-baserede billedbeskrivelser (2,43 ud af 4).

I modsætning hertil viser mine data en langt mere nuanceret praksis blandt mine informanter. Denne systematiske variation i tillid kan bedst karakteriseres med det J. D. Lee og See (2004) betegner som *kalibreret tillid* - et begreb der beskriver overensstemmelsen mellem brugerens tillid til systemet og systemets faktiske kapabiliteter. Kalibreret tillid indebærer, at brugeren hverken udviser overtillid eller mistillid, men aktivt afstemmer sin tillid efter teknologiens faktiske formåen i forskellige kontekster.

Begrebet kalibreret tillid blev først inddraget i opgavens senere faser, da jeg først her stødte på teorien. Dens relevans for at forklare deltagernes systematiske variation i tillid blev hurtigt tydelig, hvorfor det blev integreret i analysen på dette tidspunkt.

Mine observationer viser, at deltagerne systematisk justerer deres tillidsrelation baseret på situationens karakter. I lavrisikosituationer, som objektidentifikation eller æstetiske beskrivelser,

udviser de høj tillid til AI'en. IP4 formulerer det direkte: "Det gør jeg jo i de fleste tilfælde. Jeg stoler på den." Modsat udviser de samme brugere systematisk mistillid i højrisikokontekster som trafiksituationer eller medicinidentifikation. Dette illustreres tydeligt når IP2 kategorisk afgrænser anvendelsen: "Så skal du ikke bruge den til at tage billeder af lyskrydset."

Denne situationsbestemte graduering af tillid er tæt forbundet med brugernes behov for ikke-visuel verifikation. IP4's metaforiske udsagn: "Man skal jo ikke stole blindt på noget eller nogen", illustrerer denne pragmatiske tilgang. Set gennem Deweys begreb om inquiry repræsenterer dette ikke mangel på tillid, men aktiv, reflekteret tilgang til teknologiens muligheder og begrænsninger. Brugerne engagerer sig i det Dewey ville beskrive som erfaringsbaseret læring, hvor praktisk handling og refleksion gensidigt informerer hinanden.

Kalibreret tillid kunne alternativt fortolkes som udtryk for grundlæggende teknologiskepsis. Denne fortolkning ville se brugernes selektive anvendelse af Be My AI som en afspejling af generel mistillid til AI-teknologi.

Jeg finder dog denne fortolkning utilstrækkelig af to grunde. For det første viser informanterne systematisk variation i deres tillidsstrategier på tværs af forskellige brugssituationer - samme bruger udviser både høj tillid *og* systematisk skepsis afhængigt af konteksten. For det andet nævner informanterne eksplicit afvejning af risici som bevidst strategi snarere end generel mistro. Det viser IP4's metaforiske udsagn også, da jeg ikke tolker det som teknologiskepsis, men en pragmatisk forståelse af teknologiens situerede anvendelighed (Suchman, 1987).

En metodisk begrænsning, der påvirker fortolkningen, er undersøgelsens relativt homogene og teknologivante informantgruppe. En mere varieret sammensat gruppe, med forskellige grader af teknologierfaring, alder og køn, kunne have udvist mindre systematiske eller mere polariserede tillidsstrategier. Ældre brugere med begrænset teknologierfaring kunne eksempelvis have haft en mere generel skepsis over for teknologien, mens yngre brugere måske ville have vist en større tilbøjelighed til tillid. Omvendt kan det også tænkes, at ældre brugere – på trods af lav teknologierfaring – viser større tillid, fordi de ikke nødvendigvis forstår de bagvedliggende teknologiske strukturer. Ikke desto mindre vurderer jeg, at den teoretiske konsistens på tværs af de seks informanter styrker argumentet for kalibreret tillid som et relevant analytisk begreb i denne kontekst.

Studier af Bendel (2024) og Alharbi m.fl. (2024) fokuserer på verifikationsstrategier som respons på lav tillid. Mine fund viser derimod, at verifikation ikke nødvendigvis indikerer mistillid, men snarere repræsenterer behovet for at kunne vurdere risici. Dette nuancerer forståelsen af verifikation fra et udtryk for skepsis til en integreret del af forhandling om tillidsrelation.

Denne forståelse af tillid kan ses som en adaptiv kompetence, der udfordrer den teknologi-centrerede tilgang, som dominerer feltet. I stedet for udelukkende at fokusere på at forbedre AI'ens tekniske præcision, peger mine fund på behovet for at anerkende og understøtte brugernes tillidsstrategier som en ressource fremfor et problem, der udelukkende skal "løses" gennem teknologisk udvikling. En slags perfektionering, der i øvrigt aldrig vil kunne opnås, hverken for teknologier eller mennesker.

7.2 ”Billede først - intention bagefter” som designbarrierer

Mens første del af diskussionen fokuserede på, hvordan brugerne håndterer usikkerhed gennem kalibreret tillid, behandler dette afsnit, hvordan den praktiske oplevelse af teknologien skaber særlige udfordringer, der går ud over individuelle fejl. Jeg analyserer mig frem til begrebet *intensionskommunikation*, som illustrerer, hvordan specifikke designvalg ikke blot fører til fejl, men også kræver kreative håndteringsstrategier fra brugernes side.

Dette bidrager med et underbelyst perspektiv på designbarrierer i AI-teknologi til kompenserende brug. Eksisterende litteratur behandler typisk fejl som tekniske udfordringer (Bisante m.fl., 2023; Hong & Kacorri, 2024), hvor fokus ligger på præcision og nøjagtighed i billedgenkendelse. I modsætning hertil viser min analyse, at designets fundamentale interaktionslogik – ”billede først - intention bagefter” – skaber systemiske barrierer for brugerne.

Observationerne viser, hvordan Be My AI’s interaktionsflow skaber det, jeg vil kalde *paradoksals ekskludering* – en dynamik, hvor teknologien utilsigtet skaber barrierer for den målgruppe, den er designet til at støtte. Dette bliver særligt tydeligt, når brugerne har vanskeligt ved at positionere kameraet korrekt, eller når AI’en fortolker billedet på en måde, der ikke afspejler brugerens intention. Disse udfordringer er ikke blot tilfældige, men konsekvensen af systemiske designvalg som prioritering af automatisk billedfortolkning uden mulighed for at specificere formål eller kontekst på forhånd.

Dette perspektiv er i tråd med Shakespeares (2013) relationelle handicapmodel, der fremhæver, at handicap ikke blot opstår i individet, men i interaktionen mellem individ og omgivelser. Min analyse udvider denne forståelse ved at vise, hvordan ekskludering kan opstå som en konsekvens af teknologisk design, snarere end blot tilfældige tekniske fejl.

Et konkret eksempel på denne dynamik ses hos IP1, der under en af ”stationer” fik til opgave at få beskrevet et fotografi af Be My AI. AI’ens respons ved flere gentagne forsøg og manglende præcision i billedgenkendelsen resulterede i betydelig frustration: ”Så i forhold til at seende lige går forbi et billede her, og bruger lige nøjagtigt 5 sekunder på det der, og vi så har stået og været her 10 minutter nu...” Dette eksempel viser, hvordan systemets interaktionsflow kan skabe betydelige praktiske barrierer i situationer, hvor hurtig og præcis information er afgørende. Hvis denne opgave havde været vigtig, ville det være oplagt at ringe til en frivillig eller på anden måde opsøge menneskelig hjælp.

IP1’s frustration blev yderligere forstærket af systemets manglende evne til at justere billedfortolkningen: ”Det er irriterende, at det er det, den ikke vil give mig.” Dette illustrerer den asymmetri i agens, der opstår, når brugeren er afhængig af teknologien, men ikke har kontrol over dens fortolkning. Det er vigtigt at understrege, at denne ekskludering ikke er et bevidst designvalg fra udviklernes side, men en konsekvens af systemiske beslutninger i designet, som reelt skaber yderligere barrierer for brugeren Shakespeare, 2013.

Et forsvar for Be My AI’s design kunne være, at ”billede først”-modellen sigter mod umiddelbar enkelhed og brugervenlighed. Dette rationale kan have generel gyldighed for seende brugere, men min empiri understreger, at en sådan ”forenkling” for blinde brugere utilsigtet medfører øget kompleksitet. Brugere oplever manglende kontrol, behov for omfattende promptingstrategier og

en konstant kamp om fortolkningsretten.

Suchman (1987) hjælper os til at forstå dette gennem situeret handlingsteori. Hun påpeger, hvordan interaktioner altid er situerede i konkrete kontekster og opstår i mødet mellem brugerens intentioner og omgivelsernes affordances. Set gennem dette perspektiv kan vi forstå brugernes billedtagingsstrategier som forsøg på at genvinde kontrol over den situerede interaktion.

Når brugerne bevidst arrangerer objekter eller manipulerer kameravinkler, som observeret hos IP2 og IP3, omskriver de faktisk teknologiens affordances gennem praktisk handling. Dette er et eksempel på det Orlikowski (1992) beskriver som teknologiens dualitet, hvor praksis former teknologi samtidig med at være formet af den. Brugere bliver således ikke ”passive ofre” for designets begrænsninger, men aktive medskabere af nye interaktionsmuligheder (Herskovitz m.fl., 2023).

Adnin og Das (2024) dokumenterer lignende frustrationer, men fortolker dem primært som tekniske begrænsninger. Mine data peger derimod på, hvordan designvalg skaber strukturelle barrierer, der kræver kreative kompensationsstrategier. IP2’s objektmanipulation og IP3’s miljøkontrol er eksempler på dette.

Perspektivet udfordrer den designfilosofi, der i praksis lægger vægt på teknologisk enkelhed frem for brugerens selvbestemmelse. Analysen tyder på, at et enkelt design kan føre til oplevelsen af, at brugernes behov overses. Dette åbner for en diskussion om, hvordan tilgængelighed bedst forstås – ikke kun som kompensation for synstab, men også som understøttelse af brugernes behov for kontrol og agens.

7.3 Teknologihacking som bidrag til adaptiv brugeradfærd

Som svar på både tillidsudfordringer og intentionskommunikationsproblemer udvikler brugerne andre teknologiske praksisser. Set gennem Hutchins’ (1996) teori om, illustrerer *teknologihacking* som identificeret praksis, hvordan brugere udvikler strategier, der rækker ud over simpel teknisk tilpasning. Denne forståelse placerer teknologihacking som en del af et distribueret kognitivt system, hvor kognitive opgaver fordeles mellem bruger og teknologi baseret på situationens krav.

Mine data supplerer og nuancerer internationale fund om adaptiv brugeradfærd ved at pege på strategier mere komplekse end tidligere beskrevet. Herskovitz m.fl. (2023) introducere begrebet “DIY assistive technology”, men fokuserer primært på hardware. Mine resultater viser derimod, hvordan brugerne opererer som meddesignere, der aktivt reorganiserer det samlede teknologiske økosystem for den enkelte bruger. Dette illustrerer i praksis Hutchins’ idé om, at kognitive processer er situerede og distribuerede på tværs af mange elementer.

IP4’s strategi demonstrerer denne systemiske tilgang: ”Jeg ville bare gå ud af Be My AI og gå ind i et OCR-program.” Selvom dette umiddelbart fremstår som en logisk og praktisk løsning, er det netop et eksempel på et distribueret kognitivt system. Brugeren vurderer teknologiens forskellige styrker og svagheder og kombinerer dem målrettet afhængigt af opgavens krav. Set gennem Hutchins’ (1996) perspektiv repræsenterer dette en adaptiv reorganisering af det kognitive system, hvor brugeren fungerer som en ”systemarkitekt”, der fordeler kognitive opgaver mellem forskellige teknologiske komponenter.

Dette aktive valg af teknologier er særligt interessant, fordi det viser, hvordan brugerne ikke blot accepterer teknologiens begrænsninger, men aktivt kompenserer ved at kombinere forskellige værktøjer. Dette afspejler en form for metakognitiv kompetence, hvor brugerne udvikler en overordnet forståelse for, hvordan teknologier kan kombineres og optimeres for at overkomme konkrete begrænsninger.

Man kunne alternativt hævde, at denne hacking-praksis ikke er unik for denne brugergruppe, men snarere udtryk for generel teknologitilpasning, som mange erfarne brugere demonstrerer. En sådan fortolkning overser dog den særlige karakter af blinde brugeres tilpasningspraksis. I modsætning til almindelig teknologitilpasning handler den observerede praksis om fundamentale ændringer i teknologiens anvendelsesformål og om en systematisk omfordeling af kognitive opgaver, der opstår under betingelser med udpræget informationsasymmetri.

En anden mulig fortolkning er, at behovet for ”hacking” primært er et symptom på dårligt design og dermed en byrde for brugeren. Denne fortolkning undervurderer imidlertid den agens, innovation og dybe systemforståelse, brugerne udviser. Appen er i meget høj grad teknisk tilgængelig, og udviklet af personer med dyb forståelse for målgruppen. Observationerne viser, at det ikke kun handler om at overkomme begrænsninger, men om aktivt at forme og optimere et distribueret kognitivt system baseret på brugerens specialiserede viden om egne behov.

Analysen peger på, at kombinationen af afhængighed af ikke-visuel information, høje konsekvenser af fejl og systematisk brug af specialiserede værktøjer peger på en særlig form for adaptiv ekspertise. Denne praksis tyder på en form for adaptiv kompetence, hvor brugerne ikke kun behersker teknologien, men også har en dyb forståelse af egne behov og hvordan de kan mødes med de rigtige værktøjer.

J. G. Lee m.fl. (2024) beskriver lignende tilpasningsprocesser, men fokuserer typisk på individuel læring inden for enkelte teknologier. Mine data afdækker en mere kompleks praksis, hvor brugerne udvikler meta-kognitive kompetencer for *teknologiorkestration*. IP3’s forslag om ”hovedbårent kamera” demonstrerer ikke blot ønsketænkning, men systematisk refleksion over teknologiens fremtidige potentiale – hvad Ejning-Duun og Skovbjerg (2019) ville karakterisere som political inquiry, hvor brugeren aktivt engagerer sig i at forestille sig alternative teknologiske fremtider. Når IP3 taler om hovedbårne hjælpemidler flyttes baren endnu længere for teknologiorkestrationen, da der hermed kommer yderligere platforme i spil.

Dette udfordrer den deficit-orienterede tilgang, der dominerer tilgængelighedsforskning. I stedet for at fokusere på, hvordan brugere ”overvinder” teknologiske begrænsninger, peger undersøgelsen på brugerinnovation som et kernefokusområde i teknologiudviklingen. Brugernes teknologihacking illustrerer, hvordan distribueret kognition ikke blot handler om at tilpasse sig eksisterende systemer, men om at omforme systemernes kapaciteter gennem kreativ praksis. Som en note kan nævnes, at Meta, Ray-Ban og Be My Eyes arbejder på præcis det, IP3 taler ind i.¹

7.4 Tværgående mønstre

De tre hovedtemaer – kalibreret tillid, intentionskommunikation og teknologihacking – danner tilsammen et sammenhængende system, hvor hver del både påvirker og understøtter de andre.

¹<https://www.bemyeyes.com/be-my-eyes-smartglasses/>

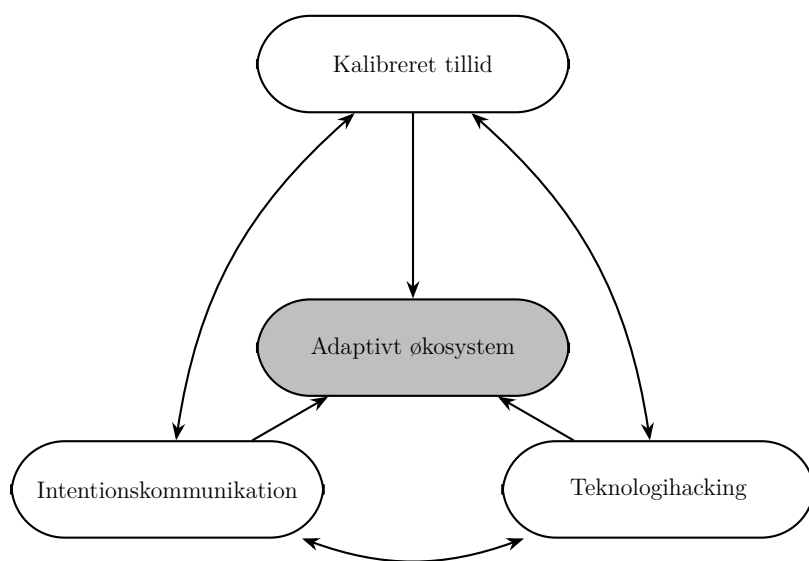
Det illustrerer dermed hvordan brugerne aktivt udvikler og tilpasser strategier for at håndtere teknologiens muligheder og begrænsninger.

Kalibreret tillid spiller en central rolle som et grundlag, der former brugernes valg og prioriteringer. Fx viser IP2, hvordan en bevidst vurdering af risici får hende til at vælge den hvide stok og følesansen frem for Be My AI i trafiksituationer. Denne praksis demonstrerer, hvordan tillidsvurderinger styrer, hvilke teknologier brugerne finder egnede i forskellige kontekster.

Samtidig påvirker tillidsstrategierne den måde, brugerne kommunikerer med teknologien på. Når IP4 stiller spørgsmål som “er du sikker?”, viser det en dynamisk proces, der former selve interaktionen med teknologien. Denne reflektive tilgang til at afklare og teste AI’ens svar skaber grobund for forbedret kommunikation - intentionskommunikation.

Intentionskommunikation fungerer således ikke kun som en teknisk tilpasning, men som en løbende dialog eller forhandling mellem brugeren og teknologien. Når brugerne gennem præcise prompter eller ændringer i billedkompositioner oplever, at AI’en giver mere præcise svar, justeres tilliden tilsvarende. Dette viser en cirkulær proces, hvor handling og refleksion gensidigt understøtter hinanden.

Teknologihacking fremstår i denne sammenhæng både som en reaktion på teknologiens begrænsninger og som en aktiv strategi for at udvide dens muligheder. IP4’s skift mellem Be My AI og OCR-programmer viser, hvordan brugerne kombinerer forskellige teknologier for at opnå bedre resultater og større kontrol over informationen. Denne praksis skaber nye måder at kommunikere intentioner på, som rækker ud over den enkelte teknologis standardfunktioner.



Figur 7.1: Illustration af samspillet mellem de tre hovedtemaer – kalibreret tillid, intentionskommunikation og teknologihacking – der tilsammen skaber et adaptivt økosystem. Diagrammet viser den cirkulære dynamik og gensidige afhængighed mellem elementerne.

Dette perspektiv kan illustreres som på figur 7.1. Brugernes strategier er ikke tilfældige tilpasninger, men en velovervejet og systematisk måde at skabe mening og handlemuligheder i teknologiske interaktioner. Set gennem Hutchins’ teori fremstår dette som en aktiv reorganisering af det teknologiske landskab, hvor tillid, kommunikation og teknologiorkestrering er tæt forbundne elementer i en samlet adaptiv praksis.

7.5 Hvordan danske fund supplerer og udfordrer internationale studier

Undersøgelsen viser tre mønstre hos danske brugere, der supplerer international forskning om AI-tillid (forskningsspørgsmål e).

Studier af Gonzalez Penuela m.fl. (2024) samt Alharbi m.fl. (2024) rapporterer primært om lav tillid og fokuserer på verifikation. Min undersøgelse af tillid udfordrer denne forståelse ved at vise systematisk variation inden for samme brugergruppe med fokus på kalibreret tillid for danske brugere, hvorfor binære tillidsmodeller udfordres. Hvor den internationale forskning fortolker skepsis som tegn på teknologiske mangler, fremtræder det i min undersøgelse, som velovervejede risikovurdering.

”Calibrated trust” blev oprindeligt udviklet af J. D. Lee og See (2004) i kontekster med høj teknologisk struktur og sikkerhedskritiske systemer - fx luftfart og moderne biler. Min undersøgelse viser, at kalibreret tillid også opstår i hverdagsteknologi, hvor brugeren ikke kan tjekke det, teknologien ved. Den kalibrerede tillid viser sig her ikke gennem teknisk monitorering, men gennem situeret erfaringsopbygning og kontekstuel risikovurdering.

Designbarrierer som strukturelhindring: Litteraturen behandler primært AI-fejl som tekniske problemer (Bisante m.fl., 2023; Hong & Kacorri, 2024), men undersøger sjældent designs rolle i brugeroplevelsen. Mit begreb ”paradoksalt ekskludering” bidrager med en analytisk dimension, der supplerer disability studies’ fokus på fysiske og sociale barrierer med designmæssige barriere strukturer.

(Adnin & Das, 2024) dokumenterer lignende frustrationer, men min undersøgelse specificerer hvordan ”billede først - intention bagefter”-strukturen systematisk flytter eller overfører agens fra bruger til system. Dette perspektiv mangler i den internationale litteratur, hvor teknologiske interaktionsmodeller, som jeg har læst litteraturen, opfattes som neutrale designvalg.

Brugerrinnovation som kilde: Mens international forskning som Herskovitz m.fl. (2023) og J. G. Lee m.fl. (2024) beskriver, hvordan brugere tilpasser sig teknologiske løsninger, viser min undersøgelse andre strategier. *Teknologihacking* fremstår ikke som en kompensation, men som aktiv skabelse af teknologisk værdi for brugeren. Dette udfordrer den internationale forsknings tendens til at fremstille brugere som passive modtagere af teknologi.

Særligt bemærkelsesværdigt er den bevidste brug af flere samtidige teknologier, som ikke er dokumenteret i internationale studier af assisterende AI. Dette kan skyldes særlige aspekter ved dansk eller måske skandinavisk tilgængelighedskultur eller høj digital kompetence, men kræver yderligere undersøgelse for at bekræfte kulturelle forskelle. Velfærdssamfund og høj levestandard kan tænkes at have en betydelig indflydelse.

Samlet set viser undersøgelsen, hvordan fokus kan flyttes fra teknologiens præstation til brugernes handlinger i praksis. Fundene viser, at tillid til AI og tilpasning ikke bare sker af sig selv, men skabes gennem brugernes kreative og innovative måder at bruge teknologien på i forskellige situationer.

7.6 Metodiske overvejelser

Den metodiske tilgang har haft stor betydning for mine resultater. Men det er samtidig vigtigt at reflektere kritisk over, hvordan metodiske valg har påvirket konklusionernes gyldighed og rækkevidde.

Observationen af informanterne i praksis med Be My AI gjorde det muligt at se strategier og mønstre, som de sandsynligvis ikke ville have nævnt eller beskrevet i interviews, der kun byggede på deres refleksioner (Holtzblatt & Beyer, 2017, s. 47).

IP2's rotation af objekter og IP3's opmærksomhed på miljøet omkring objekterne fremstod som naturlige handlinger under observationen, men ville ikke være blevet beskrevet som intentionskommunikationsstrategier uden det tilpassede kontekstuelle interview.

Dialogen undervejs muliggjorde løbende verification af observationer, hvilket har været særligt værdifuldt for forståelsen af tillidsaspekterne, da informanternes refleksioner over risiko gav noget andet og mere til analysen af deres konkrete handlinger.

Min dobbeltrolle som forsker og IKT-konsulent skabte både analytiske styrker og potentielle begrænsninger. Som insider havde jeg privilegeret adgang til informanter og dybere forståelse for målgruppen, hvilket muliggjorde nuanceret fortolkning af deres strategier. Samtidig kan min baggrund have påvirket informanternes svar i retning af at demonstrere kompetence eller give "korrekt" respons. Dette italesatte jeg i øvrigt inden interviewet, da informanterne havde dette i frisk erindring.

Rekrutteringen gennem mit netværk på IBOS resulterede i forholdvis teknologikyndige informanter med skærmlæserkompetence. Når mine informanter havde denne profil har det betydet, at jeg har kunne identificere strategier som teknologihacking og kalibreret tillid, men det kan samtidig have begrænset overførbareheden til brugere med mindre teknologierfaring.

Undersøgelsens kvalitative karakter med seks informanter (n=6) udelukker statistisk behandling data, men de identificerede mønstre og indsigter går igen på tværs af informanter, hvilket tyder på en vis grad af validitet.

7.7 Undersøgelsens analytiske bidrag og implikationer

Mine fund bidrager til forståelsen af tillid til AI-teknologi og udvikling af tilgængelighedsteknologi på tre centrale områder:

- Kalibreret tillid som situeret kompetence. Undersøgelsen viser, hvordan brugere udvikler sofistikerede tillidsstrategier baseret på funktionalitet, kontekst og konsekvenser frem for blot teknisk præcision (J. D. Lee & See, 2004). Dette udfordrer deterministiske tillidsmodeller og peger på tillid som en kompetence, brugere lærer at navigere i specifikke sammenhænge.
- Paradoksal ekskludering som designbarriere. Analysen peger på behovet for at udvide Shakespeare's (2013) relationelle model med større opmærksomhed på, hvordan design selv kan blive kilder til strukturelle barrierer. Begrebet paradoksal ekskludering demonstrerer,

hvordan assisterende teknologi utilsigtet kan skabe nye former for systematisk ekskludering, og understreger, at tilgængelig design kræver opmærksomhed på brugerautonomi.

- Brugerinnovation som kilde til viden. Når brugerne tilpasser og ”hacker” teknologien, skaber de ikke bare løsninger på problemer - de udvikler ny viden om, hvad teknologien egentlig kan bruges til (Hutchins, 1996; Orlikowski, 1992). Dette peger på, at brugernes kreative tilpasninger er en ressource frem for et tegn på, at designet ikke virker.

Analysens fund rejser spørgsmål, der peger ud over den nuværende undersøgelse. Hvordan giver disse tillidsmekanismer sig til kende i andre kontekster og blandt andre brugergrupper? Tre retninger kunne bidrage til at besvare dette:

- Komparative undersøgelser kunne vise, hvordan forskellige grupper af blinde brugere udvikler tillid - om det afhænger af faktorer som fx alder og erfaring med teknologi.
- Studier over længere tid kunne følge de samme brugere og se, hvordan deres måde at navigere tillid på ændrer sig efterhånden som teknologien bliver bedre og de selv får mere erfaring.
- Eksperimenter med nye designelementer kunne afprøve, om det hjælper brugerne, hvis AI'en er mere åben om sine usikkerheder og begrænsninger.

7.8 Implikationer for synsprofessionelle i Danmark

Med afsæt i min undersøgelses bidrag kan fundene vil jeg fremsætte nogle konkrete handlingsfor-slag for synsprofessionelle:

- Støt tillidskalibrering gennem progressiv træning. Undersøgelsen viser, at brugernes evne til situeret risikovurdering er afgørende for succesfuld anvendelse. Implementer strukturerede forløb, hvor brugere systematisk afprøver teknologien fra lavrisikosituationer til identifikation af højrisikokontekster. Dette hjælper brugere med at udvikle personlige strategier for, hvornår de kan stole på AI'en og hvornår de skal verificere gennem alternative metoder.
- Undervis i kompensation for designbegrænsninger. Undersøgelsen peger på, at Be My AI's ”billede først - intention bagefter”-struktur kan skabe paradoksalt ekskludering for målgruppen. Træn brugere i billedtagingsstrategier (kamerapositionering, objektmanipulation) og strategisk prompting, så de lærer at ”indkode” deres intention direkte i billedet. Fokuser på praktiske øvelser i forskellige anvendelsessituationer.
- Faciliter vidensdeling af brugerinnovation. Indsaml og formidl de kreative løsninger, som erfarne brugere har udviklet. Etabler praksisfællesskaber, hvor nye strategier kan deles - fx gennem korte sessioner i forbindelse med workshops eller netværk. Brugere demonstrerer deres personlige ”hacks” og tilpasninger. Anerkend disse løsninger som værdifuld ekspertise frem for afvigelser fra standard brug, og integrer de mest effektive strategier i vejledningspraksis.

Konklusion

Jeg har i denne masteropgave undersøgt, hvordan blinde brugere oplever og udvikler strategiske mønstre i interaktionen med generativ AI i billede-til-tekst-tjenesten Be My AI. Dette blev undersøgt ud fra problemformuleringen:

Hvordan oplever blinde brugere generativ AI i billede-til-tekst-tjenester som Be My AI, og hvilke strategiske mønstre udvikler de for at navigere i teknologiens funktionalitet og fejl?

Med afsæt i de fem forskningsspørgsmål fremhæves følgende hovedresultater:

Undersøgelsen viser, at de seks interviewede blinde brugeres oplevelser (forskningsspørgsmål a) spænder fra frustrationer over indlejrede begrænsninger – såsom Be My AI's 'billede først – intention bagefter'-design, der kan føre fra paradoksal ekskludering – til en kreativ, myndiggørende og selvhjulpne anvendelse i hverdagen.

For at håndtere de udfordringer og systemiske fejl (forskningsspørgsmål c) udvikler deltagerne et rigt repertoire af dels proaktive strategier, herunder bevidst billedkomposition og objektmanipulation for at forbedre inputtet til AI'en, og dels reaktive strategier i form af iterative promptingteknikker og kritisk informationsvurdering for at afkode og validere outputtet. De fungerer således som innovative teknologihackere, der aktivt former og orkestrerer deres personlige teknologiske økosystem.

De typer af fejl, der opstår (forskningsspørgsmål b) i interaktionen, omfatter faktuelle hallucinationer, kontekstuelle fejlfortolkninger af billedets indhold, kritiske udeladelser af væsentlige informationer samt relevansbaserede fejl, hvor AI'en utilsigtet fokuserer på for brugeren irrelevant information. Disse fejl har markant praktisk betydning for oplevelsen, da de i fraværet af visuel verifikation kan skabe usikkerhed, føre til fejltagelser og i visse situationer indebære potentielle risici, hvilket nødvendiggør deltageres årvågenhed og udvikling af valideringsmekanismer.

Centralt for deltageres navigation i teknologiens muligheder og begrænsninger er udviklingen af en kalibreret og situeret tillid (forskningsspørgsmål d). Fremfor en statisk eller binær opfattelse af tillid versus mistillid formes og justeres tillidsrelationen til Be My AI dynamisk gennem deres konkrete, praktiske erfaringer med både AI'ens styrker og dens iboende fejlbarlighed. Denne kalibrerede tillid er afgørende for deres risikovurdering og deraf følgende beslutninger om, hvornår, hvordan og til hvilke formål teknologien meningsfuldt kan anvendes. Den kombinerer både funktionelle, kontekstuelle og konsekvensbaserede vurderinger.

Undersøgelsens kvalitative fund vedrørende de seks danske informanter (forskningsspørgsmål e) supplerer og udfordrer eksisterende internationale studier ved at levere en dybdegående forståelse af deres agens, komplekse adaptive strategier og den situerede natur af deres interaktion

med generativ AI.

Konceptet om kalibreret tillid bidrager til at nuancere det ofte teknisk orienterede eller generelt fokuserede internationale perspektiv på tillidsdannelse, mens forståelsen af teknologihacking som en form for distribueret kognition og brugerinnovation giver et bredere indblik i, hvordan brugerne aktivt skaber værdi i samspil med teknologien.

Samlet set understreger disse resultater behovet for at designe teknologier, der anerkender og understøtter brugerens agens, kontekstuelle behov og det fundamentale ønske om meningsfuld kontrol, selvhjulpenhed og præcis intentionskommunikation, så AI reelt kan fungere som en myndiggørende og ligestillingsfremmende kraft for at øge selvhjulpenheden.

Bibliografi

- Adnin, R., & Das, M. (2024). "I look at it as the king of knowledge": How Blind People Use and Understand Generative AI Tools. *Proceedings of the 26th International ACM SIGACCESS Conference on Computers and Accessibility*, 1–14. <https://doi.org/10.1145/3663548.3675631>
- Alharbi, R., Brewer, R. N., & Schoenebeck, S. (2022). Understanding Emerging Obfuscation Technologies in Visual Description Services for Blind and Low Vision People [Place: New York, NY, USA Publisher: Association for Computing Machinery]. *Proc. ACM Hum.-Comput. Interact.*, 6. <https://doi.org/10.1145/3555570>
- Alharbi, R., Lor, P., Herskovitz, J., Schoenebeck, S., & Brewer, R. N. (2024). Misfitting with AI: How blind people verify and contest AI errors. *The 26th International ACM SIGACCESS Conference on Computers and Accessibility*, 1–17. <https://doi.org/10.1145/3663548.3675659>
- Be My Eyes. (2025). *Terms of service* [Be my eyes]. Hentet 31. maj 2025, fra <https://www.bemyeyes.com/terms-of-service/>
- Bendel, O. (2024). How Can Generative AI Enhance the Well-being of Blind? [Version Number: 1]. <https://doi.org/10.48550/ARXIV.2402.07919>
- Bisante, A., Dix, A., Panizzi, E., & Zeppleri, S. (2023). To err is AI. *Proceedings of the 15th Biannual Conference of the Italian SIGCHI Chapter*, 1–11. <https://doi.org/10.1145/3605390.3605414>
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://doi.org/10.1191/1478088706qp063oa>
- Brinkmann, S., & Lene Tanggaard (Red.). (2022). *Kvalitative metoder: en grundbog* (3. udgave). Hans Reitzel.
- Busaeed, S., Mehmood, R., & Katib, I. (2022). Requirements, Challenges, and Use of Digital Devices and Apps for Blind and Visually Impaired. *Preprints*, 1–20. Hentet 11. maj 2025, fra https://www.preprints.org/manuscript/202207.0068/download/final_file
- Dewey, J. (1988). Experience and education [1938]. I J. A. Boydston (Red.), *The later works of john dewey, 1925-1953* (s. 1–62, Bd. 13). Southern Illinois University Press.

- Ejsing-Duun, S., & Skovbjerg, H. M. (2019). Design as a mode of inquiry in design pedagogy and design thinking. *International Journal of Art & Design Education*, 38(2), 445–460. <https://doi.org/10.1111/jade.12214>
- Elkjær, B. (2025). *At analysere med John Dewey*. Nyt fra Samfundsvidenskaberne.
- Gonzalez Penuela, R. E., Collins, J., Bennett, C., & Azenkot, S. (2024). Investigating use cases of AI-powered scene description applications for blind and low vision people. *Proceedings of the CHI Conference on Human Factors in Computing Systems*, 1–21. <https://doi.org/10.1145/3613904.3642211>
- Herskovitz, J., Xu, A., Alharbi, R., & Guo, A. (2023). Hacking, switching, combining: Understanding and supporting DIY assistive technology design by blind people. *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*, 1–17. <https://doi.org/10.1145/3544548.3581249>
- Holtzblatt, K., & Beyer, H. (2017). *Contextual design : design for life* (Second edition.). Morgan Kaufmann.
- Hong, J., & Kacorri, H. (2024). Understanding how blind users handle object recognition errors: Strategies and challenges. *The 26th International ACM SIGACCESS Conference on Computers and Accessibility*, 1–15. <https://doi.org/10.1145/3663548.3675635>
- Hutchins, E. (1996). *Cognition in the wild*. The MIT Press. <https://doi.org/10.7551/mitpress/1881.001.0001>
- ibos.dk. (2025). *Tal og fakta om syn* [Instituttet for Blinde og Svagsynede]. Hentet 30. maj 2025, fra <https://ibos.dk/viden/mennesker-med-synshandicap-i-danmark/>
- Kameswaran, V., Robinson, J., Sambasivan, N., Aggarwal, G., & Morris, M. R. (2024). Help and The Social Construction of Access: A Case-Study from India [event-place: St. John’s, NL, Canada]. *Proceedings of the 26th International ACM SIGACCESS Conference on Computers and Accessibility*. <https://doi.org/10.1145/3663548.3675606>
- Lee, J. G., Lee, B., Choi, S., Seo, J., & Choe, E. K. (2024). Identify, adapt, persist: The journey of blind individuals with personal health technologies. *Proceedings of the ACM on Interactive, Mobile, Wearable and Ubiquitous Technologies*, 8(2), 1–21. <https://doi.org/10.1145/3659585>
- Lee, J. D., & See, K. A. (2004). Trust in Automation: Designing for Appropriate Reliance [Place: London, England Publisher: Oxford University Press OUP]. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- Madake, J., Bhatlawande, S., Solanke, A., & Shilaskar, S. (2023). A Qualitative and Quantitative Analysis of Research in Mobility Technologies for Visually Impaired People. *IEEE Access*, 11, 82496–82520. <https://doi.org/10.1109/ACCESS.2023.3291074>
- OpenAI. (2025). *ChatGPT* [GPT-4o]. Hentet 29. maj 2025, fra <https://chat.openai.com/chat>

- Orlikowski, W. J. (1992). The duality of technology: Rethinking the concept of technology in organizations. *Organization Science*, 3(3), 398–427. <https://doi.org/10.1287/orsc.3.3.398>
- Sahib, N. G., Stockman, T., Tombros, A., & Metatla, O. (2013 september). Participatory design with blind users: a scenario-based approach. I P. Kotzé, G. Marsden, G. Lindgaard, J. Wesson & M. Winckler (Red.), *Lecture notes in computer science* (s. 685–701, Bd. LNCS-8117). Springer. https://doi.org/10.1007/978-3-642-40483-2_48
- Schön, D. A. (2017, 2. marts). *The reflective practitioner* (0. udg.). Routledge. <https://doi.org/10.4324/9781315237473>
- Shakespeare, T. (2013). The Social Model of Disability. I *The Disability Studies Reader*. Taylor & Francis Group. <http://ebookcentral.proquest.com/lib/aalborguniv-ebooks/detail.action?docID=1125176>
- Stangl, A., Sadjó, E., Emami-Naeini, P., Wang, Y., Gurari, D., & Findlater, L. (2023). “Dump it, Destroy it, Send it to Data Heaven”: Blind People’s Expectations for Visual Privacy in Visual Assistance Technologies [event-place: Austin, TX, USA]. *Proceedings of the 20th International Web for All Conference*, 134–147. <https://doi.org/10.1145/3587281.3587296>
- Suchman, L. A. (1987). *Plans and situated actions: the problem of human-machine communication*. Cambridge University Press.
- Tzvetkova-Arsova, M., & Tomova, M. (2024). Use of Technology and Applications in Orientation and Mobility of Visually Impaired Persons in Bulgaria: A Contemporary Overview [event-place: Linz, Austria]. *Computers Helping People with Special Needs: 19th International Conference, ICCHP 2024, Linz, Austria, July 8–12, 2024, Proceedings, Part I*, 408–415. https://doi.org/10.1007/978-3-031-62846-7_49
- Zhao, K., Lai, R., Guo, B., Liu, L., He, L., & Zhao, Y. (2024). AI-Vision: A Three-Layer Accessible Image Exploration System for People with Visual Impairments in China [Place: New York, NY, USA Publisher: Association for Computing Machinery]. *Proc. ACM Interact. Mob. Wearable Ubiquitous Technol.*, 8(3). <https://doi.org/10.1145/3678537>

Bilag A: Søgestrategi

Tabel 9.1: Søgestrategier og resultater på tværs af databaser

Database	Søgekombination	Resultater / Relevante	Filter og udvælgelseskriterier
ACM Digital Library	[All: "be my eyes"] AND [[All: blind*] OR [All: "visual* impair*"] OR [All: "low vision"]]	92 / 28	Filtreret på E-publication date: 11/2022–03/2025. Artikler omhandlende "Be My Eyes", synshandicap, brugeroplevelser, fejlhåndtering, tillid eller AI-interaktion.
Scopus	TITLE-ABS-KEY(("be my eyes" OR "be my ai") AND (blind* OR "visual* impair*" OR "low vision") AND (error* OR fail* OR verif* OR trust* OR reliab*))	2 / 2	Begrænset til 2022–2025. Direkte fokus på "Be My Eyes"/"Be My AI", synshandicap og aspekter af fejl, verifikation eller tillid.
Google Scholar	("be my eyes" OR "be my ai") AND (blind OR "visual impairment" OR "low vision") AND (error OR fail OR verify OR trust OR reliable)	306 / 15	Afgrænset til 2022–2025. Artikler og preprints om "Be My Eyes"/"Be My AI", fejlstrategier og tillid, screenet for direkte relevans.
ERIC (EBSCOhost)	(TI "be my eyes" OR AB "be my eyes" OR TI "be my ai" OR AB "be my ai") AND (TI blind* OR AB blind* OR TI "visual* impair*" OR AB "visual* impair*" OR TI "low vision" OR AB "low vision")	0 / 0	Publiceringsdato 2022–2025. Ingen relevante artikler identificeret.
PsycINFO (EBSCOhost)	Samme søgestreng som ERIC	0 / 0	Publiceringsdato 2022–2025. Ingen relevante artikler identificeret.

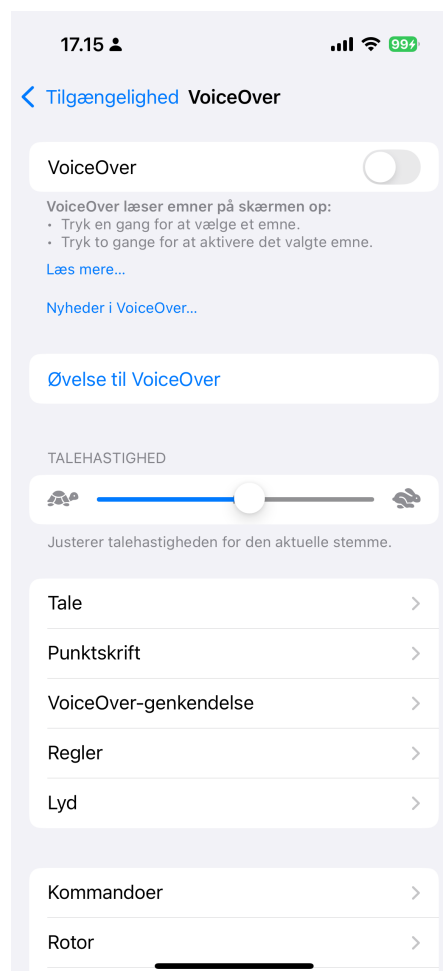
Bilag B: Skærmlæsere

Jeg har promptet ChatGPT: ”Skriv et afsnit om hvad en skærmlæser er og hvordan kan bruger den og hvorfor man bruger den.”.

ChatGPT: *En skærmlæser er et specialiseret stykke software, der gør det muligt for personer med synshandicap at få adgang til digital information og interagere med computere og mobile enheder. Skærmlæseren konverterer den visuelle information på skærmen til syntetisk tale eller punktskrift via en braille-enhed. Brugeren navigerer gennem indholdet ved hjælp af tastaturgenveje eller touch-bevægelser, som giver mulighed for at læse tekst, identificere billeder, forstå struktur og betjene interaktive elementer som knapper, formularfelter og links. Skærmlæsere er uundværlige for personer med nedsat eller ingen syn, da de muliggør uafhængig brug af digitale teknologier til alt fra arbejde og uddannelse til sociale og personlige aktiviteter. De understøtter inklusion og ligestilling ved at give adgang til information og kommunikation på lige fod med seende brugere (OpenAI, 2025).*

Det er svært at vise en skærmlæser i funktion, da det er mest af alt, er lyd. Når VoiceOver på iOS er i brug, er skærmlæseren i udgangspunktet ”gennemsigtig” da den virker som et lag oven det almindelige indhold, hvorfor der ikke noget at se. Mennesker med blindhed anvender desuden ofte ”skærmtæppe”, en indstilling der medfører, at skærmen altid er slukket. Dette sparer på strømmen og beskytter oplysningerne på skærmen for nysgerrige blikke, da en person med blindhed ikke nødvendigvis holder telefonen på samme måde som en seende.

Skærmlæsere skal betragtes som en brugergrænseflade på lige fod med et grafisk interface, da de fungerer som forbindelsen mellem bruger og computer. Hvor en almindelig brugergrænseflade viser information på skærmen gennem vinduer og knapper, omdanner skærmlæseren det samme indhold til tale og lyd. Den er ikke bare et program, der læser tekst op - det er en komplet måde at styre computeren på. Skærmlæseren har sine egne kommandoer og genveje, som giver brugeren fuld kontrol over, hvad der skal læses op og i hvilken rækkefølge. På den måde er skærmlæseren en selvstændig brugergrænseflade, der gør det muligt at bruge computeren uden at se på skærmen. Det viser, hvordan teknologi kan tilpasses forskellige behov uden at miste funktionalitet.



Figur 9.1: Indstillingerne til skærmlæseren *VoiceOver* på iOS - den eneste skærmlæser til iOS. Kilde: Eget foto.

Bilag C: Interviewguide

1. Introduktion (5-10 min)

- Præsenter dig selv og projektet
- Gennemgå samtykke og optagelse - husk fotos. To mikrofoner - optager til iPhone. Lydtjek!

2. Observation (60 min)

a) Fri brug (15 min)

- Bed deltageren om at vise dig, hvordan de normalt bruger Be My AI
- Observer uden at gribe ind
- Noter især reaktioner på første svar og opfølgende spørgsmål

b) Strukturerede opgaver på kontoret (45 min)

Opgave 1: Tekst på emballage

- Brug tilgængelige objekter: panodil, chokolade, fiskeolie etc.
- "Kan du tage et billede af denne [emballage] og fortælle mig, hvad der står på den?"
- Opfølgende spørgsmål ved førevarer: "Hvad med allergener?"

Opgave 2: Praktisk opgave

- Brug tilgængeligt tøj/tilbehør: "Tag et billede af min skjorte/jakke/ur og spørg AI om farve/stil"
- Alternativ: "Tag et billede af min arbejdsplads og bed AI beskrive opsætningen"

Opgave 3: Beskrive billede

- Brug billeder på væggene.
- Endnu enklere: "Tag et billede ud ad vinduet og få AI til at beskrive, hvad der er udenfor..."

Opgave 4: Køkken

- Hvor længe kan mælken holde sig?

3. Refleksiv dialog (30-45 min)

- "Hvordan oplever du, at AI'en svarer uden at kende din hensigt?"
- "Har du udviklet særlige måder at stille opfølgende spørgsmål på?"
- "Ville du foretrække at kunne forklare dit behov først?"
- "Hvordan adskiller det sig fra at få hjælp fra et menneske?"
- "Har du oplevet, at den har nægtet at svare?"
- "Hvordan påvirker det din tillid, at du har haft de oplevelser her i dag?"

4. Afrunding (5-10 min)

- Spørg om der er noget, de gerne vil tilføje
- Tak dem og forklar næste skridt

Observationsfokus

1. Deltagerens reaktion på AI'ens første svar
2. Hvordan de formulerer opfølgende spørgsmål
3. Strategier de bruger til at få AI'en til at forstå deres faktiske behov

Bilag D: Kodningsoversigt

Tabel 9.2 viser et repræsentativt udsnit af det samlede kodningsskema med eksempler på de mest fremtrædende koder i relation til de tre hovedtemaer. Tabellen præsenterer 43 ud af i alt 217 identificerede koder - udvalgt fordi de udgør opgavens analytiske kerne, mens de resterende enten overlapper med disse eller er for kontekstspecifikke til meningsfuld gengivelse uden omfattende forklaring.

Tabel 9.2: Udvalgte eksempler fra kodningsprocessen

Informant	Citat	Primær kode	Analytisk note/Begrundelse
<i>Hovedtema 1: Tillidsforhold og risikohåndtering i asymmetrisk information</i>			
IP1	"Det er jo altid i forhold til, jeg er altid en lille smule skeptisk i forhold til kunstig intelligens."	Tillid og verifikations-strategier	Udtrykker en grundlæggende skepsis overfor AI – kræver bekræftelse af output. Vigtig holdningsindikator.
IP1	"Jeg har været i nogle situationer, hvor jeg ved præcis, hvad det er for en information, jeg søger. Og så er det irriterende, at det er det, den ikke vil give mig."	Negativ oplevelse	Tydelig frustration over mismatch mellem hensigt og AI-output. Indikerer gentagen oplevelse.
IP1	"...giver det her mening, kan det virkelig være rigtigt, at jeg skal betale 135.000 kr. i husleje denne måned?"	Brugerens egne validerings-systemer	Beskriver, hvordan han bruger sin "snusfornuft" til at validere systemets svar. Viser en kontinuerlig kognitiv belastning ved brug af teknologien. Kritisk vurdering og tvivl.
IP1	"Jeg bruger den også til at bruge til at skrive brev osv. Der får jeg en hurtig. . . [...] Men man skal selv vurdere lidt. . . giver det her mening."	Tillid og verifikations-strategier	Udtrykker aktiv brug, men med forbehold for validitet. Afspejler tillidsbalancering. Emotionel skepsis. Brug under forbehold.

Tabel 9.2: Udvalgte eksempler fra kodningsprocessen

Informant	Citat	Primær kode	Analytisk note/Begrundelse
IP1	"Så tænker jeg, at i en del tilfælde kunne det være rigtig rart. [...] At sige, at nu tager jeg et billede, og det skal jeg bruge til... Altså, det du skal svare på, det er det her."	Hensigt	Udtrykker ønske om at kunne præcisere hensigt før billedtagning. Illustrerer forskellen mellem AI og menneskelig kontekstforståelse. Hensigtsstyring.
IP1	"Jeg tænkte egentlig, at jeg skulle sidde helt nonchalant og pege derhen, hvor cirka det er, eller at jeg skulle rejse mig og prøve at tage lidt oppefra."	Billedtagningsstrategier	Afvejer forskellige fysiske strategier for bedre framing og forståelighed.
IP1	"Så i forhold til at seende lige går forbi et billede her, og bruger lige nøjagtigt 5 sekunder på det der, og vi så har stået og været her 10 minutter nu..."	Sammenligning med menneskelig hjælp	Tydelig kontrast mellem effektivitet hos seende og AI-assisteret interaktion. Oplevet langsommelighed og begrænsning.
IP1	"...nu har jeg en datter, nu har hun så ikke en størrelse, hvor det er, men lad os nu sige, at man havde et barn på 4-5 år, og igen ude at lege, løber og kommer hjem og græder... Så er det irriterende, når den siger, at det der vil jeg ikke sige noget om, fordi der er et barn uden tøj..."	Design-bias og etiske grænser	Påpeger, hvordan systemets etiske regler kan blokere for praktisk anvendelse i dagligdagen. Afslører spændingen mellem beskyttende design og praktiske behov. Etiske dilemmaer og praktisk brug.
IP2	"Fordi jo mere konkret jeg er, jo bedre er den i stand til at forstå, hvad det er, jeg vil have den til."	Promptingstrategier	Viser udvikling i promptingstrategier – brugeren har lært at være mere specifik for bedre resultater. Kognitiv refleksion. Promptstrategier og interaktion.

Tabel 9.2: Udvalgte eksempler fra kodningsprocessen

Informant	Citat	Primær kode	Analytisk note/Begrundelse
IP2	"Hvis jeg nu havde været laktoseintolerant, så havde jeg måske troet, at så var det sikkert for mig at drikke den her mælk."	Risikovurdering	Eksemplificerer en potentiel farlig fejltolkning, der kunne have haft sundhedsmæssige konsekvenser. Hverdagsbrug og fødevarer.
IP2	"Så skal du ikke bruge den til at tage billeder af lyskrydset."	Risikovurdering	Citatet eksemplificerer en sikkerhedsmæssig begrænsning i brugen. Vigtigt for trafiksituationer. Visuel tilgang og billedforståelse.
IP2	"Det er jo den, der til at starte med udvælger, hvad den synes er vigtigt."	Design	Afslører et grundvilkår i interaktionen – modellen har initial styring af fokus. Begrænser brugerens handlefrihed i starten af interaktionen. Begrænsninger og potentiale i systemdesign.
IP2	"Jeg roterer den lige en gang mere, fordi det var det der med, at holdbarhedsdatoen vi var interesseret i."	Billedtagningsstrategier	Citatet viser bevidst manipulation af objektets vinkel for at opnå bedre information.
IP2	"Hvis jeg nu havde skrevet, at den er formet som en fodbold [...] så ville den stadigvæk have skrevet, tak for at rette mig."	Design	Demonstrerer hvordan modellen ukritisk accepterer rettelser – en designbeslutning, der har konsekvenser for tillid og læring. Begrænsninger og potentiale i systemdesign.

Tabel 9.2: Udvalgte eksempler fra kodningsprocessen

Informant	Citat	Primær kode	Analytisk note/Begrundelse
<i>Hovedtema 2: Brugerudviklede strategier for intentionskommunikation</i>			
IP3	"Jamen tilliden er fin, men nu er jeg jo også sådan lidt skeptisk-agtig. Så jeg tager ikke alting for gode varer, vel."	Risikovurdering	Beskriver sin tilgang til at validere AI'ens output og bevare kritisk sans. Kognitiv kritisk vurdering. Tillid og validering.
IP3	"Jeg er sgu ligeglad med, om bordet er rundt, eller lampen er moderne."	Hensigt	Markerer tydeligt, hvad der ikke er relevant for ham – tydeliggørelse af intention. Vurdering af relevans. Filtrering og fokus.
IP3	"Verbositeten er overvældende ofte."	Negativ oplevelse	Oplever informationsmængden som belastende – utilfredshed med output-stil. Emotionel overvældelse. Interaktionskvalitet.
IP3	"Jeg sørger for, at der er nogenlunde frit omkring det væsentlige."	Billedtagsnings-strategier	Forklarer sin metode til at sikre tydelighed i billedet ved at fjerne visuel støj. Strategi for fokusoptimering. Visuel tilgang og billedforståelse.
IP3	"Jamen, jeg er relativt kort for hovedet. Også fordi ellers kan man skrive sig synder sammen."	Prompting-strategier	Forklarer sin strategi med at holde opfølgende spørgsmål korte og præcise. Strategi for kommunikationsøkonomi. Effektiv interaktion.
IP3	"Ja, ja. Så må man spørge ind til det, ikke? Men det er jo også et... Altså meningen er, at det skal være enkelt at bruge, ikke?"	Begrænsninger i appens design	Kommenterer på appens design-filosofi og accepterer den trods begrænsninger. Metarefleksion om brugervenlighedsafvejning. Balance mellem enkelthed og funktionalitet.
IP3	"Kæmpe, kæmpe, kæmpe fordel."	Positivt overrasket	Udtrykker stor begejstring for forskellen Be My AI har gjort i forhold til tidligere praksis. Emotionel begejstring. Selvstændighed og empowerment.

Tabel 9.2: Udvalgte eksempler fra kodningsprocessen

Informant	Citat	Primær kode	Analytisk note/Begrundelse
IP3	"Det er jo ikke alting, der rager andre mennesker, man spørger om. Det er den ene ting. Den anden ting er, at det kan jo være på alle mulige mystiske tidspunkter af døgnet."	Sociale aspekter ved hjælpemiddel-teknologi	Fremhæver værdien af privatliv og selvstændighed som en fordel ved AI over menneskelig hjælp. Indikerer at teknologien har værdi ud over den rent funktionelle. Autonomi og værdighed.
IP3	"Man kunne måske have nogle forskellige prompts, hvor man fokuserede på det, der var centralt i billedet."	Prompting-strategier	Foreslår forskellige prompter til forskellige formål. Designforslag. Differentieret interaktion.
IP3	"Jeg synes måske, at det ville være rigtig dejligt, hvis vi fik noget hovedbårent kamera, så man fik en håndfri fidus."	Design	Foreslår hardware-forbedring der ville øge anvendeligheden betragteligt. Designforslag for ergonomisk innovation. Fysisk tilgængelighed og brugervenlighed.
IP4	"Det gør jeg jo i de fleste tilfælde. Jeg stoler på den."	Tillid og verifikations-strategier	Generel høj tillid, på trods af anerkendte fejlmuligheder. Tilkendegiver balanceret tillid med betinget accept. Tillidskalibrering.
IP4	"Man skal jo ikke stole blindt på noget eller nogen."	Tillid og verifikations-strategier	Formulerer en grundlæggende kritisk tilgang til al information. Generaliserbar strategi, ikke kun AI-specifik. Holdning om kritisk tænkning. Brugerens ansvar.
IP4	"Så spurgte jeg den, er du sikker? Nej, undskyld, det er en flæsketeg."	Prompting-strategier	Anvender efterspørgsel efter præcision som strategi. Eksempel på hvordan ompromptning fører til forbedret output. Strategi for validerende prompt. Iteration og forbedring.
IP4	"Jeg ville bare gå ud af Be My AI og gå ind i et OCR-program."	Sammenligning med menneskelig hjælp	Foretrækker anden teknologi (OCR) for at få nøjagtige resultater. Prioriterer præcision frem for bekvemmelighed. Strategi for teknologisk skift. Specialiserede værktøjer.

Tabel 9.2: Udvalgte eksempler fra kodningsprocessen

Informant	Citat	Primær kode	Analytisk note/Begrundelse
IP4	"Det var derfor min hjerne kortsluttede lidt, da jeg skulle spørge den ind til derude på billedet."	Prompting-strategier	Beskriver, hvordan det er svært at formulere præcist spørgsmål. Oplevet kognitiv belastning i promptsituationen. Emotionel forvirring. Sprog og kognition.
IP4	"Jeg synes, det skal være en mulighed. Helt klart."	Hensigt	Brugeren ønsker at kunne angive intention før billedtagning. Efterspørgsel om præintentional interaktion. Designforslag for forprompt. Intentionel styring.
IP4	"Så spurgte jeg den bare: er du sikker? Og så rettede den sig faktisk."	Prompting-strategier	Anvender valideringsstrategi for at få mere præcist output. Gentagelse og kontekst styrker AI's præcision. Strategi for kritisk forespørgsel. Interaktiv kontrol.
IP4	"Min største hallucination, den sagde, da den sagde min flæsketeg var et stykke franskbrød."	Fejlbeskrivelser/hallucinationer	Grotesk fejl i objektidentifikation illustrerer AI's begrænsninger. Fører til efterfølgende prompt og korrektion. Eksempel på grov fejl. Visuel usikkerhed.
IP4	"Så langt, så godt. Hold da op, det var lidt meget information."	Negativ oplevelse	Oplever outputtet som overvældende i mængde. Reflekterer verbositet som mulig barriere. Emotionel overbelastning. Informationsmængde og klarhed.
IP4	"Jeg er egentlig bare glad for, at den har beskrevet noget."	Forventningsjustering og tilfredshed	Udtrykker basal værdsættelse af beskrivelse fremfor intet. Signal om teknologisk ydmyghed og accept. Vurdering af grundlæggende værdi. Minimumseffekt.

Tabel 9.2: Udvalgte eksempler fra kodningsprocessen

Informant	Citat	Primær kode	Analytisk note/Begrundelse
<i>Hovedtema 3: Systemets muligheder og begrænsninger i hverdagspraksis</i>			
IP5	"Det er fint nok med en flot beskrivelse, men det er ikke det, jeg har brug for."	Hensigt	Efterspørger præcise svar frem for æstetisk output – klart udtryk for intention. Formålsstyret interaktion. Målstyret brug.
IP5	"Det ærgrer mig da. Især når man prøver at belære den. Så ærgrer det mig, at den så ikke lærer noget."	Negativ oplevelse	Bliver skuffet over AI'ens manglende læringssevne. Emotionel skuffelse. Forventningsbrud.
IP5	"Jeg kan godt mærke, at man skal jo selv lige træne sig selv i at stille det rigtige spørgsmål."	Prompting-strategier	Erkender behovet for at udvikle sine kommunikationsfærdigheder med AI. Læring sker gennem praksis. Kognitiv strategisk udvikling. Brugerlæring og tilpasning.
IP5	"Men jeg vil dog være varsom. Jeg er mere varsom nu"	Tillid og verifikations-strategier	Udvikler øget skepsis baseret på erfaring med fejl. Erfaring fører til mere kritisk tilgang. Emotionel forsigtighed. Langtidseffekter og tillidsjustering.
IP5	"Den nævner jo ikke alle ingredienser."	Begrænsninger i appens design	Observerer at vigtige fødevareoplysninger udelades. Strategi kræver opfølgning. Tilgængelighed og fødevareinformation.
IP5	"Hold da op. Det var da godt en lille cylindrisk [genstand]. Den havde jeg da fuldstændig overset."	Positivt overrasket	Bliver positivt overrasket over, at AI'en ser detaljer han overser. Styrker tillid til teknologi. Emotionel overraskelse. Visuel tilgang og billedforståelse.
IP6	"Det går jo faktisk hen og bliver en lille smule halvfærdigt, vil jeg jo mene."	Risikovurdering	Reflekterer over konsekvenser ved fejlagtige beskrivelser. Pointerer problem for brugere uden punktlæsningssevne. Sikkerhed og funktionel risiko. Følger af fejl.

Tabel 9.2: Udvalgte eksempler fra kodningsprocessen

Informant	Citat	Primær kode	Analytisk note/Begrundelse
IP6	"Sådan noget med vejvisning og sådan noget... det vil jeg sgu heller ikke rigtigt turde."	Risikovurdering	Udtryk for mistillid til AI ved navigation – potentielt farligt scenarie. Kritik af AI's pålidelighed i bevægelse i fysisk rum. Emotionel forsigtighed. Sikkerhed og bevægelsesfrihed.
IP6	"Altså ved det første billede, inden du overhovedet begynder at spørge ind til noget. Du får jo virkelig mange inputs."	Begrænsninger i appens design	Oplever informationsmængden som overvældende. Peger på behov for mere fokuseret output. Designkritik af informationsstruktur. Verbositet og overinformation.

Oversigt over primære koder og tværgående temaer

Tabel 9.3 viser sammenhængen mellem de mest fremtrædende primære koder og de tværgående temaer, der blev identificeret i analysen.

Tabel 9.3: Relation mellem primære koder og tværgående temaer

Primær kode	Relaterede tværgående temaer
Tillid og verifikationsstrategier	Pålidelighed og konsekvens; Troværdighed; Langtidseffekter og tillidsjustering
Risikovurdering	Sikkerhed og kritisk anvendelse; Sundhed og ansvar; Følger af fejl
Billedtagningsstrategier	Visuel tilgang og billedforståelse; Bruger-adaptation og læring; Visuel optimering
Promptingstrategier	Promptstrategier og interaktion; Brugerlæring og tilpasning; Effektiv kommunikation
Fejlbeskrivelser/hallucinationer	Usikkerhed i output; Visuel usikkerhed; Troværdighed og validitet
Hensigt	Hensigtsstyring; Målstyret brug; Mål og kontekst
Begrænsninger i appens design	Begrænsninger og potentiale i systemdesign; Systembegrænsninger og kontekstforståelse; Fokus og informationsmængde
Forventningsjustering og tilfredshed	Teknologiens potentiale og begrænsninger; Teknologisk progression; Brugerlæring og erkendelse
Sammenligning med menneskelig hjælp	Menneskelig støtte og kontekst; Effektivitet og tålmodighed; Grænser for maskinel vurdering

Bilag E: Samtykkeerklæring

Samtykke til deltagelse i masterprojekt om AI-baserede billede-til-tekst-tjenester for personer med nedsat syn

Kære informant

Som studerende på Masterstudiet i IKT og læring ved Aalborg Universitet, ønsker jeg at behandle dine personoplysninger til mit studieprojekt med titlen *“Brugeroplevelser med Generativ AI til billede-til-tekst for mennesker med nedsat syn: En kontekstuel undersøgelse”*.

Projektet har til formål at undersøge hvordan svagsynede og blinde brugere oplever generativ AI i billede-til-tekst-tjenester som fx Be My AI, og hvilke brugsmåder der viser sig, når systemet fungerer – og når det fejler.

Dataindsamling

I forbindelse med projektet indsamler jeg følgende data: - Dit navn - Skærmoptagelse af din interaktion med billede-til-tekst-tjenesten - Lydoptagelse af din stemme under interviewet - Fotos der tages under sessionen - Dit telefonnummer og email til kontaktformål

Dataindsamlingen vil foregå gennem: - **Skærmoptagelse:** For at dokumentere din faktiske interaktion med billede-til-tekst-tjenesten - **Lydoptagelse:** For at optage interviewet til senere transskribering - **Fotos:** For at dokumentere specifikke situationer under brugstesten

Jeg beder venligst om dit samtykke til at behandle dine personoplysninger til mit projekt.

Dataansvarlig: Rasmus Paasch (studienummer: 20221968) er, som studerende, dataansvarlig for behandling af dine oplysninger til projektet, uafhængigt af Aalborg Universitet.

Opbevaring og sikkerhed

Dine oplysninger bliver opbevaret sikkert, og jeg benytter dem udelukkende til ovenstående formål. For at beskytte dine følsomme personoplysninger opbevares alle data krypteret på Aalborg Universitets systemer, og adgangen er begrænset udelukkende til undertegnede som projektansvarlig.

Dine data vil ikke blive delt med tredjeparter udover min projektvejleder på AAU, som er underlagt samme fortrolighedsforpligtelser.

Sletning: Oplysninger i form af optagelser samt indhentede svar i spørgeskemaet slettes, når projektperioden slutter ultimo juni 2025.

Dine rettigheder

Du har altid ret til at trække dit samtykke tilbage. Ønsker du at trække dit samtykke tilbage, kan du kontakte mig på dx7p@kk.dk.

Databeskyttelsesforordningen giver dig en række rettigheder, som du finder beskrevet nedenfor.

Samtykke

Via mail: Svar blot på denne mail med “Jeg giver mit samtykke” eller lignende bekræftelse, så betragtes det som en gyldig underskrift på denne samtykkeerklæring.

Alternativt fysisk underskrift:

Dato: _____

Navn: _____

Underskrift: _____

Oplysningstekst - Sådan behandler jeg dine data

Jeg er dataansvarlig

Rasmus Paasch, studerende ved Aalborg Universitet og IKT-konsulent på IBOS.

Formålet med at behandle dine oplysninger

Projektet har til formål at undersøge hvordan blinde brugere oplever generativ AI i billede-til-tekst-tjenester som Be My AI, og hvilke brugsmåder der viser sig, når systemet fungerer – og når det fejler. Formålet med indsamling er at kunne analysere brugeroplevelsen og interaktionen med disse teknologier for blinde brugere.

For at kunne foretage analyser, har jeg brug for navn samt lyd- og skærmoptagelser. Jeg optager den fulde samtale og interaktion med produktet. Interviewet vil efterfølgende blive transskriberet i forbindelse med kodning.

Jeg behandler disse personoplysninger

Almindelige personoplysninger

Jf. databeskyttelsesforordningens art. 6, stk. 1, litra a.

Følgende almindelige personoplysninger behandles om dig: Navn, telefonnummer og email.

Følsomme personoplysninger

Jf. databeskyttelsesforordningens art. 9, stk. 2, litra a.

Da projektet specifikt fokuserer på blinde brugere, kan optagelser af din stemme og eventuelle billeder potentielt afsløre information om dit helbred (synstab), hvilket kategoriseres som følsomme personoplysninger under databeskyttelsesforordningen.

Følgende følsomme personoplysninger behandles om dig: Stemme (via lydoptagelse) og eventuelle billeder via skærmoptagelser eller fotos.

Sådan opbevarer jeg dine oplysninger

Jeg opbevarer dine personoplysninger, så længe det er nødvendigt i forhold til formålet med at indhente dit samtykke og i henhold til gældende lovgivning. Herefter sletter jeg dine personoplysninger.

Alle data slettes senest ultimo juni 2025.

Dine rettigheder

Når jeg behandler dine personoplysninger, har du ifølge databeskyttelsesforordningen flere rettigheder. Det betyder bl.a., at du har ret til sletning og dataportabilitet, og i visse tilfælde har du ret til indsigt, berigtigelse, begrænsning og til at gøre indsigelse mod min behandling af de omfattede personoplysninger.

Du har altid mulighed for at trække dit samtykke tilbage. Vær dog opmærksom på, at du ikke kan trække dit samtykke tilbage med tilbagevirkende kraft.

Læs mere om dine rettigheder: Datatilsynets hjemmeside

Vil du klage?

Du er altid velkommen til at kontakte mig med spørgsmål. Hvis du mener, at jeg ikke behandler dine oplysninger efter reglerne, så kontakt mig gerne.

Du har også altid mulighed for at klage til Datatilsynet på dt@datatilsynet.dk.

Jeg opfordrer dig dog til at kontakte mig først, da jeg vil gøre, hvad jeg kan, for at imødekomme din klage.

Videregivelse til og fra tredjepart inden for EU

Dine oplysninger (eller dele af dine oplysninger) kan blive videregivet til: Min personlige dataopbevaringsmappe på Aalborg Universitet (AAU) i forbindelse med opbevaring i projektperioden. Ingen andre tredjeparter vil få adgang til dine data.