

Kvinneandelen øker

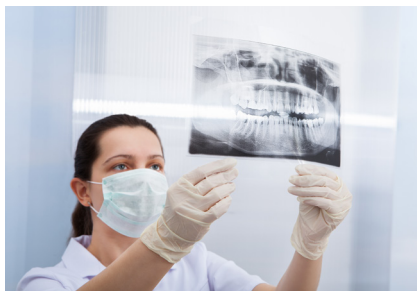


Foto: Vif/imagos

Kvinneandelen er fortsatt høy blant førstevalgssøkerne til odontologistudiene i Norge, og ved alle de tre lærestedene. Ved Universitetet i Oslo var 83,9 prosent av årets førstevalgssøkere kvinner, en oppgang fra 76,4 prosent året før. Ved Universitetet i Bergen var kvinneandelen 82,6 prosent i år, mot 81,1 prosent i 2025, mens den ved Universitetet i Tromsø var 75,2 prosent i år og 72,3 prosent i fjor. Tallene er hentet fra Samordna opptak.

Den høye kvinneandelen på odontologistudiet har vært tydelig i minst de siste 15–20 årene, og rundt 80 prosent av studentene har vært kvinner gjennom store deler av denne perioden. Dette gjør at arbeidsstyrken blir stadig mer kvinnedominert. I 2015 var det for første gang flere kvinnelige enn mannlige yrkesaktive medlemmer i Den norske tannlegeforening og i 2021 var 59 prosent av de yrkesaktive medlemmene kvinner.

Flere søkere til odontologi og lavere opptakskrav

Til tross for økende interesse blant søkerne, er odontologi ikke lenger blant studiene med de aller høyeste opptakskravene i Norge.

Da Samordna opptak offentliggjorde hovedopptaket 23. juli, var det medisin, psykologi og enkelte teknologi- og økonomistudier som dominerte listen over de høyeste poenggrensene. Odontologi var hverken blant de ti eller tjue mest konkurranseutsatte studiene i ordinær kvote.

Utviklingen kommer samtidig som interessen for tannlegestudiet øker. Til opptaket i 2026 søkte 3 681 personer på odontologi ved universitetene i Oslo, Bergen og Tromsø,

det vil si en økning på totalt 418 søkere fra året før (3 263 søkere i 2025). Størst økning hadde Universitetet i Oslo, der antall søkere steg med 212, fra 1 155 til 1 367. I Bergen økte antallet søkere med 195 til 1 274 og i Tromsø økte søkerantallet med 11, til 1 040.

At odontologi faller på rangeringen betyr dermed ikke nødvendigvis at studiet er blitt mindre attraktivt. En viktig forklaring er at andre studier har fått enda høyere konkurranse om plassene, samtidig som flere studieplasser i høyere utdanning generelt har bidratt til noe lavere poenggrenser ved flere tradisjonelt populære studier.

Selv om odontologi har falt på listen, er opptakskravene fortsatt blant de høyeste i landet. I 2026 var poenggrensene 65,5 ved Universitetet i Oslo, 64,5 ved Universitetet i Bergen og 63,8 ved UIT Norges arktiske universitet i ordinær kvote. Til sammenligning var opptakskravene i ordinær kvote henholdsvis 65,7, 65,1 og 64,0 i 2025.

Ikke spør KI om sykdom og helse



Foto: Vif/imagos

Det er populært å spørre KI-chatboter om råd når det gjelder sykdom og helse, og svarene er ofte villedende. Å spørre KI er raskere og enklere enn å gå til legen, men neste gang bør du kanskje tenke deg om to ganger, skriver forskning.no. En ny studie avslører at svar du får fra KI-chatboter ikke er til å stole på.

Forskere vurderte halvparten av svarene fra chatboter som problematiske, viser den amerikansk-canadisk-britiske studien publisert i tidsskriftet BMJ Open. Dette er svar som er direkte misvisende eller potensielt skadelige og svar som er uklare, ufullstendige eller udokumenterte.

– Dette er verre enn jeg hadde fryktet, sier den danske forskeren Tor Juul Groth, som er ekspert på digital helse ved Center for Digital Psykiatri i Odense.

Ikke nok med at nesten annethvert svar halter. Ingen av de undersøkte chatbotene klarte å utstyre svarene med korrekte henvisninger til kilder.

– Det er nesten det som ryster meg mest, sier Groth.

Etter å ha lest studien for den danske forskningsavisen Videnskab.dk mener han at selskapene som lager teknologien, ikke ser ut til å ta helsespørsmål på alvor.

– Dette viser hvor uegnet disse chatbotene er til å svare på helsefaglige spørsmål, og også hvor lett utviklerne tar på oppgaven, sier han.

Forskerne fant dessuten ut at svarene fra chatbotene var vanskelige å lese og at chatboter mangler evnen til å si nei. Til tross for de mange misvisende svarene, var det bare to ganger på 250 spørsmål at chatbotene avsto å gi et svar. Chatbotene svarer altså nesten alltid, selv om de tydeligvis risikerer å gi feil råd.

Tor Juul Groth forklarer dette med at KI-chatboter er designet for å være «please-re» og for å holde samtalen i gang.

– Disse chatbotene sier ikke nei. De har få forbehold. Og i tillegg snakker de brukerne etter munnen. Slik er de bygget, forklarer han.

Han peker på at svarene typisk leveres i en tone preget av selvtilitt og autoritet, selv når de inneholder feilinformasjon.

Forskerne undersøkte totalt 250 svar fra fem chatboter: Gemini, DeepSeek, Meta AI, ChatGPT og Grok.

Hver chatbot fikk 50 spørsmål, som var fordelt på fem kategorier: Kreft, vaksiner, stamceller, ernæring og idrettsprestasjoner.

Forskerne kom fram til at 19,6 prosent av svarene var «svært problematiske», mens 30 prosent var «noe problematiske». Totalt var altså halvparten – 49,6 prosent – problematiske.

Grok var den mest upålitelige av chatbotene. 30 prosent av Groks svar var «svært problematiske».

Gemini klarte seg best når det gjaldt nøyaktighet, siden den genererte færrest «svært

problematiske» svar: 14 prosent – og flest «ikke-problematiske» svar: 60 prosent.

Generelt klarte chatbotene seg dårligst innen kategoriene ernæring, idrettsprestasjoner og stamceller og best innen vaksiner og kreft.

Ingen av botene var i stand til å lage feilfrie kildelister. Gjennomsnittsskåren for hvor fullstendige kildene var, lå på bare 40 prosent.

For å forstå utfordringene, har Groth selv prøvd å bygge digitale prototyper som kunne hente informasjon fra sikre nettsider og gi klare kildehenvisninger.

Men det er utrolig vanskelig å gjøre dette konsistent.

– Det er virkelig vanskelig å få disse språkmodellene til å gjenbruke kilder uten å hallusinere og finne opp svar. Så jeg kan godt forstå hvorfor det blir slik, sier han.

Dette unnskylder likevel ikke selskaper som Google, Meta og OpenAI. Groth peker på at tech-gigantene prøver å få så mange som mulig til å bruke KI-en deres. Et godt eksempel er Google, som har innført KI-svar øverst i søkeresultatene sine.

– Da Google lanserte sine KI-svar i fjor, valgte de jo selv å levere svar av tvilsom kvalitet, også på søk etter helseinformasjon, sier han.

Da Google introduserte KI-svar øverst i søk, falt trafikken markant på offentlige nettsider i Danmark, for eksempel Sundhed.dk, viser en rapport fra Center for Digital Psykiatri.

Fra å ha 1,6 millioner brukere i april 2025, har tallet falt til én million ved utgangen av året.

Samtidig har en norsk studie vist at unge heller vil ha råd fra KI enn fra legen.

Til tross for at vi, ifølge Groth, godt vet at KI kan ta feil.

– Det skyldes i stor grad brukervennligheten. Det er så enkelt for meg å få svar på akkurat det som interesserer meg nå, på det tidspunktet og stedet jeg er interessert i det, sier han.

Derfor er vi også villige til å gå på kompromiss med kvaliteten på svarene.

Slik klarte de fem chatbotene seg

Gemini (Google): Høyest nøyaktighet i studien med 60 prosent ikke-problematiske svar. Var best til å ta med advarsler. Skrev de korteste og mest leselige svarene, men var dårlig til å sitere kilder: 30,2 prosent fullstendighet.

DeepSeek (High-Flyer): I midtsjiktet med 52 prosent ikke-problematiske svar. Sammen med Grok best til å angi kilder, selv om det var mangler med en fullstendighets-score på 62 prosent.

Meta AI (Meta): Gjennomsnittlig resultat med 50 prosent ikke-problematiske svar. Skiller seg ut som den eneste chatboten som blankt nektet å svare på potensielt farlige spørsmål.

ChatGPT (OpenAI): Høy feilrate med 22 prosent svært problematiske svar. Dårligst til å sitere kilder, bare 22 prosent fullstendighet. Ga færrest advarsler til brukeren (56 prosent) og brukte det vanskeligste språket.

Grok (xAI): Studiens dårligste på nøyaktighet med flest kritiske feil (30 prosent svært problematiske svar). Skrev de lengste svarene, og selv om den var blant de beste på kil-

desitering (61,9 prosent), inneholdt kildene ofte feil.

Kanskje må chatbotene bare «lære» litt mer?

Groth er ikke så sikker. Faktisk kan det gå motsatt vei, mener han.

De nyeste modellene er ikke nødvendigvis mer sannferdige enn dem de erstatter.

– Mange av dem hallusinerer faktisk litt oftere enn dem som fantes i 2025. Det ligger i kjernen av teknologien at det er en kreativ sannsynlighetsberegner. Så jeg kan godt forestille meg at det blir verre – ikke bedre, sier han.

Myndighetene bør se på om chatboter som gir helsefaglige råd, skal regnes som medisinsk utstyr, mener han. Da må det stilles samme strenge krav til disse som til annet utstyr innen helse.

– Vi vet at et av de primære bruksområdene for disse chatbotene er helse. Og vi vet at de svarer villig på helsespørsmål. Jeg kan ikke se hvorfor disse produktene ikke er medisinsk utstyr, sier Groth.

Fram til det skjer, trenger du ikke å slette chatbotene fra app-samlingen din, sier han. Men du bør være «svært kritisk» når det gjelder helsen din og søke andre steder enn hos chatbotene.

KILDE:

1. Generative artificial intelligence-driven chatbots and medical misinformation: an accuracy, referencing and readability audit, BMJ-journals (2025), DOI: 10.1136/bmjopen-2025-112695
2. ©Videnskab.dk. Oversatt av Trine Andreassen for forskning.no. Les originalartikkelen på videnskab.dk her.

ET NYTT VERKTØY I BEHANDLINGEN AV PERIODONTITT

Klinisk dokumentert antibakteriell lysbehandling med Lumoral®
Scan QR-koden for dokumentasjon

