

Helgelandssykehusets forskerskole

Plan for våren 2025

- | | |
|--------------------|--|
| 14.01.2025 | Hva er en mastergrad og hvordan kan jeg som ansatt ved Helgelandssykehuset ta en master? |
| 25.02.2025 | Hva kreves for å ta en doktorgrad? |
| 25.03.2025 | Regler for bruk av ChatGPT og andre store språkmodeller (LLM) i forskning |
| 29 .04.2025 | Akademisk CV |
| 17.06.2025 | Finnes det en «Fast Track» til suksess innen forskning? |

Helgelandssykehusets forskerskole

Tematiske presentasjoner



HVORDAN BLI EN FORSKER?

MÅNEDLIGE TEMATISKE PRESENTASJONER VED
HELGELANDSSYKEHUSET

Enkelte tirsdager 14:15-15:00 (ca hver tredje uke)

Helgelandssykehuset etablerte i februar 2024 en forskerskole for å motivere til mer forskning ved sykehuset. Forskerskolen, som har fått navnet HELGA, skal være et tilbud for alle ansatte ved Helgelandssykehuset – fra de som lurer på om forskning kan være noe for dem, kanskje har de et ønske om å ta en master- eller doktorgrad eller de har en ide til et mulig forskningsprosjekt, til de som ønsker å kvalifisere seg til en professor II eller seniorforskerstilling. Store deler av aktiviteten vil være digital med nettbaserte forelesninger, møter og seminarer.

Formålet med Helgelandssykehusets forskerskole (HELGA) er å stimulere til økt forskning på alle nivå ved Helgelandssykehuset. Vi ønsker med HELGA å heve kvaliteten og relevansen til forskningen samt legge forholdene bedre til rette for at ansatte ved Helgelandssykehuset kan drive med forskning og/eller kan ta master eller doktorgrad ved å knytte seg til master- eller doktorgradsprogrammene til fortrinnsvis UIT Norges arktiske universitet (UIT) og Nord universitet (Nord).

HELGA skal gjennom forskeropplæringen bidra til økt samarbeid mellom UIT, Nord og andre aktører i lokalsamfunnet. De store avstandene mellom Helgelandssykehusets ulike lokasjoner og til dele også til UIT og Nord, krever at aktiviteten i betydelig grad vil være nettbasert og/eller hybrid med en kombinasjon av fysisk og nettbasert deltakelse. Men det vil også være viktig å etablere fysiske møteplasser for ansatte, studenter og forskere tilknyttet forskerskolen.

Kliniker, midlertidige forskningsprosjekt eller ansatte/forskere ved HSYK som ønsker å være i forskningsavdelingen ved HSYK skal holde en eller flere av presentasjonene f.eks. i forbindelse med interne egne seminarer kan ta kontakt med:

Sina Mørkved, PhD Forsknings- og læringssjef Epost: sina.morkved@helgelandssykehuset.no Mobil: 92286055	Gunnar Hærtvigsen, PhD Professor Epost: gunnar.haertvigsen@helgelandssykehuset.no Mobil: 90657785
---	---

30-januar-2025

1

INNHOOLD

Hva enhver (forsker) bør vite om forskning	4
Hvordan starte, gjennomføre og avslutte et forskningsprosjekt?	5
Hva er en mastergrad og hvordan kan jeg som ansatt ved Helgelandssykehuset ta en master?	6
Hva kreves for å ta en doktorgrad (ph.d.)?	7
Hvordan skrive en masteroppgave, doktorgradsavhandling eller vitenskapelig rapport?	8
Veiledning av mastergrads- og doktorgradsstudenter	9
Hvordan evaluere en doktorgradsavhandling?	10
Hvordan skrive en vitenskapelig artikkel?	11
Rapporttidsskrifter («Predatory journals») og rapportkonferanser («Predatory conferences»)	12
Forskningsetiske regler og retningslinjer (forskningsetikk)	13
Vitenskapelig fusj i medisinsk forskning	14
Plagiering – kanskje det dumme du kan gjøre som forsker?	15
Hvordan presentere egen forskning på en vitenskapelig konferanse?	16
«Don't be such a scientist!» Hvordan fortelle de gode historiene?	17
20x20 – formidling på 6 minutter og 40 sekunder	18
Hvordan skrive et konferansetrukk?	19
Hvordan utforme en vitenskapelig poster?	20
Fagfellevurdering («Peer review») og fagfeltrapporter («Referee report»)	21
Søknad om finansiering av forskningsprosjekter	22
De Brunns forskning: Tverrfaglige regler for god forskning og vitenskapelig rapportering	23
Hvorfor skal jeg bruke en mentor for å utvikle meg som forsker?	24
Mentor-samtalen	25
Publiseringstradisjoner under endring – angår det meg?	26
Nettverksbygging: Fra lokal forskningsgruppe til internasjonalt forskernettverk	27
Finnes det en «Fast track» til suksess innen forskning?	28
Tjuvfilm «suksesskriterier» for en karriere innen forskning	29
Hva ville jeg truttet fram om jeg kun fikk tre til å gi ett råd om forskning?	30
Hvordan spille en professor II stilling (bi-stilling)?	31
Lokuttsetten for tippingskvalifikasjoner	32
Søknadskomite	33
Teorier om vitenskapelig aktivitet	34
Medforfatterkap («Co-authorship»)	35
Hvordan håndtere kulturelle forskjeller?	36
Akademisk CV	37
Kjennetegn	38

30-januar-2025

2

Akademisk virksomhet og universitetenes rolle	39
Hva gjør egentlig en docent eller professor?	40
Hvordan bli en PI (principal investigator)?	41
Hvordan bygge og vedlikeholde et godt omdømme?	42
Hvordan løse opp eller utsig en kollega?	43
«Publish or perish?» En professors vei mot «Publish»	44
Hvordan planlegge et forskningsopphold i utlandet?	45
Regler for et vellykket forsknings samarbeid	46
Regler for et vellykket tverrfaglig samarbeid	47
Regler for bruk av ChatGPT og andre store språkmodeller (LLM) i forskning	48
«Akademisk husarbeid»	49
Håndtering av forskningsdata fra helseforskning	50
Hva kjennetegner en god forskningsleder?	51
Hva gjør en forskningsgruppeleder?	52
Hva menes med formidling av forskning?	53
Vær strategisk, ikke kynisk	54
Finne din flytzone	55
Hva menes med formidling av forskning?	56
How to succeed as a foreigner in Norway?	57

30-januar-2025

3

Helgelandssykehusets forskerskole

Tematiske presentasjoner



HVORDAN BLI EN FORSKER?

MÅNEDLIGE TEMATISKE PRESENTASJONER VED
HELGELANDSSYKEHUSET

Enkelte tirsdager 14:15-15:00 (ca hver tredje uke)

Helgelandssykehuset startet i februar 2024 en forskerskole for å motivere og merke forskere i sykehuset og i kommunen. Sykehuset har et stort behov for forskning og utvikling, og vi ønsker å etablere et miljø der forskning og utvikling er en naturlig del av arbeidet. Vi ønsker å etablere et miljø der forskning og utvikling er en naturlig del av arbeidet. Vi ønsker å etablere et miljø der forskning og utvikling er en naturlig del av arbeidet.

Formålet med Helgelandssykehusets forskerskole (HELGA) er å stimulere til økt forskning på tvers ved Helgelandssykehuset. Vi ønsker med HELGA å heve kvaliteten og kompetansen til forskere og samtlige forskere, bidra til økt kunnskap og utvikling, og etablere et miljø der forskning og utvikling er en naturlig del av arbeidet.

HELGA skal gjennom forskningsprosjekter bidra til økt samarbeid mellom UVT, Nord og andre aktører i lokalsamfunnet. De store avstandene mellom Helgelandssykehusets ulike lokasjoner og til dele også til UVT og Nord, krever at aktiviteten i betydning grad vil være nettbasert og/eller hybrid med en kombinasjon av fysisk og nettbasert deltakelse. Men det vil også være viktig å etablere fysiske møteplasser for ansatte, studenter og forskere til å møte og diskutere forskning og utvikling.

Klinikk, medisin, og forskning vil i forskningsprosjektene bli etablert i forbindelse med interne og eksterne seminarer og konferanser.

Siss Mørkved, PhD
Forsknings- og Innovasjonssjef
Epost: siss.morkved@helgelandssykehuset.no
Mobil: 92286065

Gunnar Hartvigsen, PhD
Professor
Epost: gunnar.hartvigsen@helgelandssykehuset.no
Mobil: 90637785

30-januar-2025

INNHold

Hva enhver (forsker) bør vite om forskning	4
Hvordan starte, gjennomføre og avslutte et forskningsprosjekt?	5
Hva er en mastergrad og hvordan kan jeg som ansatt ved Helgelandssykehuset ta en master?	6
Hva kreves for å ta en doktorgrad (ph.d.)?	7
Hvordan skrive en masteroppgave, doktorgradsavhandling eller vitenskapelig rapport?	8
Veiledning av mastergrad- og doktorgradsstudenter	9
Hvordan evaluere en doktorgradsavhandling?	10
Hvordan skrive en vitenskapelig artikkel?	11
Rapporttidsskrifter («Predatory journals») og rapportkonferanser («Predatory conferences»)	12
Forskningsetiske regler og retningslinjer (forskningsetikk)	13
Vitenskapelig faglig/medisinsk forskning	14
Publisering – kanskje det dumme du kan gjøre som forsker	15
Hvordan skrive en vitenskapelig artikkel?	16
Hvordan skrive en vitenskapelig artikkel?	17
Hvordan skrive en vitenskapelig artikkel?	18
Hvordan skrive en vitenskapelig artikkel?	19
Hvordan skrive en vitenskapelig artikkel?	20
Fagfellevurdering («Peer review») og fagfeltrapporter («Referee report»)	21
Søknad om finansiering av forskningsprosjekter	22
Nettverksbygging: Fra lokal forskningsgruppe til internasjonalt forskernetverk	27
Finnes det en «Fast track» til suksess innen forskning?	28
Tjuvform og plagiatsaker i forskning	29
Hva vil jeg si om jeg har fått en opplysning?	30
Hva vil jeg si om jeg har fått en opplysning?	31
Hva vil jeg si om jeg har fått en opplysning?	32
Hva vil jeg si om jeg har fått en opplysning?	33
Hva vil jeg si om jeg har fått en opplysning?	34
Medforfatterkap (Co-authorship)	35
Hvordan håndtere kulturelle forskjeller?	36
Akademisk CV	37
Karriereplan	38

Akademisk virksomhet og universitetenes rolle	39
Hva gjør egentlig en docent eller professor?	40
Hvordan bli en PI (principal investigator)?	41
Hvordan bygge og vedlikeholde et godt undervisningsmiljø?	42
Hvordan løse opp eller utvalg i kullfisk?	43
"Publish or perish"? En professors vei mot "Publish"	44
Hvordan planlegge et forskningsopphold i utlandet?	45
Regler for et vellykket forsknings samarbeid	46
Regler for et vellykket tverrfaglig samarbeid	47
Regler for bruk av ChatGPT og andre store språkmodeller (LLM) i forskning	48
«Akademisk husarbeid»	49
Håndtering av forskningsdata fra helsetilforskning	50
Hva kjennetegner en god forskningsleder?	51
Hva kjennetegner en god forskningsleder?	52
Hva kjennetegner en god forskningsleder?	53
Hva kjennetegner en god forskningsleder?	54
Hva kjennetegner en god forskningsleder?	55
How to succeed as a foreigner in Norway?	56

30-januar-2025

Ta kontakt om du/dere ønsker at jeg skal arrangere et miniseminar med utgangspunkt i ett av temaene

Gunnar.Hartvigsen@helgelandssykehuset.no

Regler for bruk av ChatGPT og andre store språkmodeller (LLM) i forskning

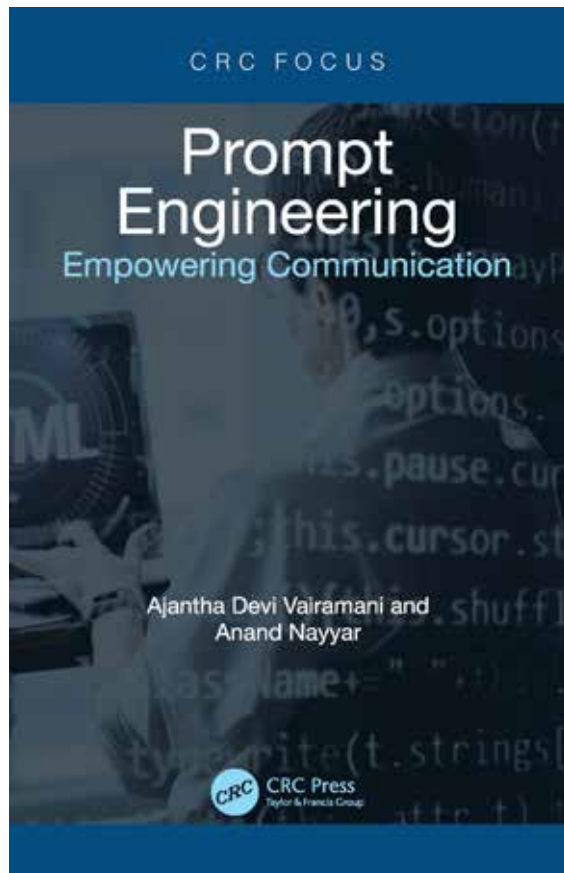
Professor Gunnar Hartvigsen

Helgelandssykehuset HF

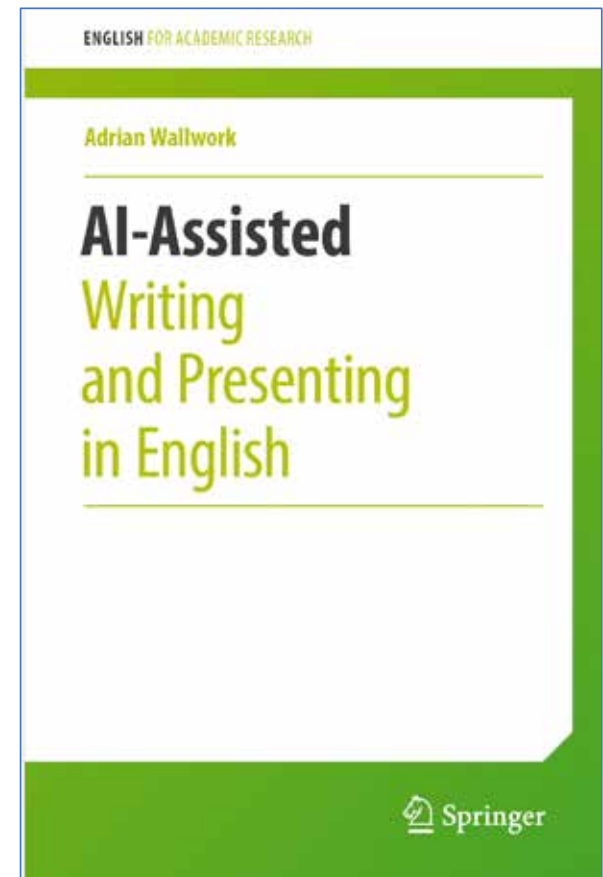
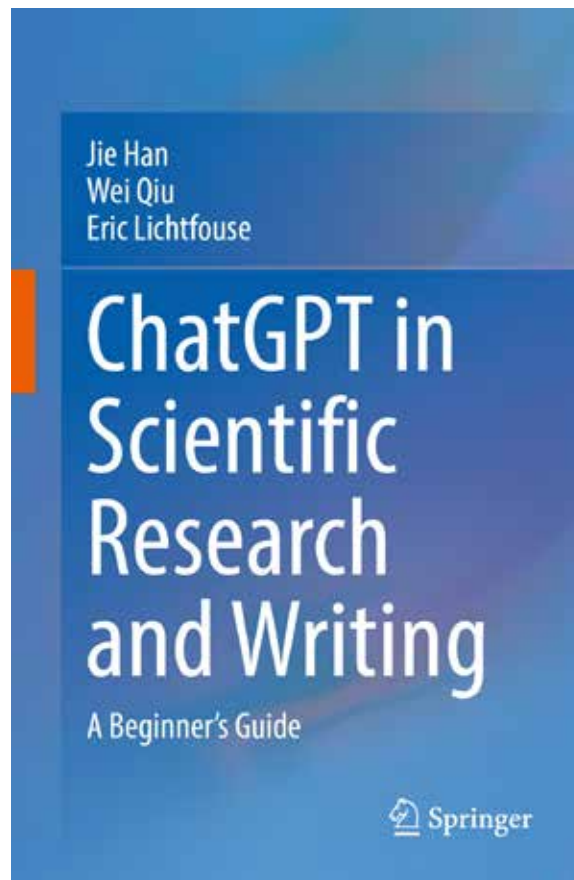
Universitetet i Tromsø – Norges arktiske universitet

25.03.2025

Det er hjelp å få ...



Eksempler på nyttige bøger




Klasse

[Hjemme](#)
[Tilbak](#)
[Skolest](#)
[Hjelp](#)

SI DOKUMENT, DEL 4

Slik kan KI hjelpe deg med språkvask, illustrasjoner og presentasjoner

I fjende del av KI skolen viser professor Morten Goodwin hvordan KI-verktøy kan brukes til å finne relevante artikler, til å lage illustrasjoner og presentasjoner.



Nettavisen Foto: Shutterstock / iStockphoto
 Illustrasjon: Morten Goodwin
 Illustrasjon: Morten Goodwin
 Illustrasjon: Morten Goodwin



[Netforum](#)
[Lærare](#)
[Studere](#)
[Netforum Help](#)

[Kthrono](#)

Hvordan gjøre kunstig intelligens til forsknings-assistent

I siste del av Kthrono KI-skole viser professor Morten Goodwin blant annet hvordan kunstig Intelligens (KI) kan hjelpe til med å utvikle hypoteser, lage forskningsdesign og brukersundersøkelser, og analysere datasett.



[Anvendelse av DeepQNs i medisin](#)
[Visualisering av forsterket læring med forklaringer](#)
[Kthrono: Enkle data-analytiske spørsmål](#)
[Kthrono: Enkle data-analytiske spørsmål](#)
[Kthrono: Enkle data-analytiske spørsmål](#)



Morten Goodwin

Professor i Kthrono

[Forskningsdesign](#)
[Datasett](#)
[Bruk](#)

[Forskningsdesign](#)
[Datasett](#)
[Bruk](#)

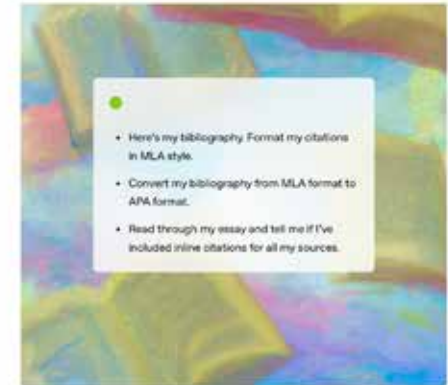
[Forskningsdesign](#)
[Datasett](#)
[Bruk](#)

A Student's Guide to Writing with ChatGPT

[Share](#)

1. Delegate citation grunt work to ChatGPT

AI excels at automating tedious, time-consuming tasks like formatting citations. Just remember to cross-check all source details against original materials for accuracy. By taking care of the grunt work, ChatGPT can free you to focus on the fun part: being creative, coming up with original ideas, and formulating an undeniable argument.

[Try it out ↗](#)

- Here's my bibliography. Format my citations in MLA style.
- Convert my bibliography from MLA format to APA format.
- Read through my essay and tell me if I've included inline citations for all my sources.



- What is Keynesian economics, and how does it differ from classical economic theory?
- Give me a broad overview of different theories of motivation.

2. Quickly get up to speed on a new topic

ChatGPT can jumpstart your research by providing a foundational understanding of a subject.

[Try it out ↗](#)

<https://openai.com/chatgpt/use-cases/student-writing-guide/>

Best AI Content Creators 2025

Experience the future of content creation with our curated selection of the best AI content creators. From generating persuasive copy to crafting engaging articles, these cutting-edge tools streamline your workflow, reduce costs, and help you deliver consistently high-quality content—without the heavy lifting. Discover how AI can revolutionize your strategy and take your brand's voice to the next level. Click on each of these leading services below for more info and special deals!

ADVERTISER DISCLOSURE

1	 HubSpot ★ Most Popular ★ - Use OpenAI's GPT model to boost your business - Ideate, create, and share content in seconds - Generate copy for blogs, social posts, emails, and more #1 Business hub to take your productivity to a new level - Used by over 194,000 customers in over 120 countries - Easily integrates with 1000+ apps and services - Get started free. No credit card required.	9.8 BEST OVERALL 	FREE! CLAIM OFFER www.hubspot.com
Excellent - 4.5  EDITORS' CHOICE  So easy to use It offers a user-friendly interface and versatile features for various writing tasks. James M. Jun 23rd, 2023  Must-Have Tool for SEO Content... Hubspot has been a game-changer for my SEO content writing. It's efficient, creative, and incredibly helpful, saving... Arjun B. Feb 23rd, 2023  Great Choice for content creators A great choice for anyone needing that extra push while writing, be it because you're facing a writer's block or... Josephine N. Mar 11th, 2023			
2	 ClickUp - Elevate your writing skills to new heights - Suggests grammar and readability improvements - Crafts emails, reports, blog posts, or creative content - Knowledge manager, project manager, and writer	9.5 GREAT 	FREE! CLAIM OFFER www.clickup.com
3	 Rytr - Generates engaging, human-like text - Supports various formats, from blogs to ads - Over 40+ templates to simplify your writing process - Named among the best AI writing tools by G2	9.3 VERY GOOD 	From 9,00 USD CLAIM OFFER
4	 ContentShake AI - AI Content Creator by SEMrush - Built by SEO experts for SEO writers - Get a full article optimized for your keywords - Full SEO article generation in 6+ languages	9.0 VERY GOOD 	From 60,00 USD CLAIM OFFER
5	 SURFER - Boost organic traffic with AI content optimization - Instant feedback on your on-page SEO performance - Optimize content for any language and market - Integrates with Google Docs, WordPress, Jasper, +	8.8 VERY GOOD 	From 79,00 USD CLAIM OFFER

Opplev fremtiden for innholdsskaping med vårt kuraterte utvalg av de beste AI-innholdsskaperne.

Fra å generere overbevisende tekster til å lage engasjerende artikler, effektiviserer disse banebrytende verktøyene arbeidsflyten din, reduserer kostnader og hjelper deg med å levere innhold av konsekvent høy kvalitet – uten tunge løft.

<https://software.fish/ai-writing-assistant/best-ai-writing-assistant>

Det er hjelp å få ...

FORSKERHÅNDBOKEN

KUNSTIG INTELLIGENS

Store Norske Leksikon hevder OpenAI har stjålet hele leksikonet deres

Serverloggene til Store Norske Leksikon (SNL) viser at roboter fra OpenAI har lastet ned hele leksikonet 30 millioner ganger.



Erik Bolstad er redaktør i Store Norske Leksikon. Foto: Kari Blom Haugstulen



Tekst:
Joar Hystad
joar.hystad@khrono.no

Publisert 24.03.2025 - 10:58



Redaktør i Store Norske Leksikon, Erik Bolstad, bekrefter overfor [TV 2](#) at roboter fra OpenAI skal ha utsatt SNL for et enormt svært og systematisk tyveri.

Store Norske Leksikon består av rundt 200.000 artikler, som alle er skrevet av eksperter og fagfolk, de aller fleste med tilhørighet innenfor universitets- og høgskolesektoren.

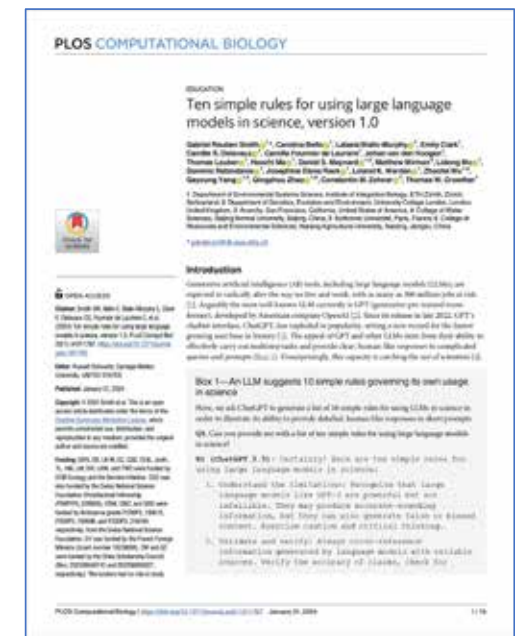
Da SNL gikk gjennom serverne sine i sommer, oppdaget de at OpenAI hadde lastet ned enorme mengder av stoffet deres, som er beskyttet av opphavsretten.

Ten simple rules for using large language models in science

Generative KI-systemer forventes å føre til store endringer i måten forskning utføres på. Disse er basert på store språkmodeller (Large Language Models – LLM).

Den mest kjente språkmodellen er GPT (generative pre-trained transformer, eller generativ forhåndstrent transformator på norsk) , utviklet av firmaet OpenAI.

Modellene gir ofte klare svar på kompliserte spørsmål og har allerede blitt svært populære innen flere bransjer.



Ten simple rules for using large language models in science



Men bruk av slike systemer er ikke utelukkende positivt. Smith m.fl. forteller at:

«Ved siden av store potensielle fordeler står store potensielle farer, og innenfor begge disse kategoriene er det sikkert bruk av generativ AI som vi (og andre) ikke har forutsett. Når vi utforsker disse nye verktøyene, må vi faktisk ikke miste av syne de uløste dilemmaene som fortsatt omgir generativ AI-teknologi. Noen av utfordringene som er relevante for forskning inkluderer bruk av uspesifiserte, potensielt partiske treningsdatasett for kommersielle modeller, en implisitt tilknytning til eksisterende vitenskapelige paradigmer, og et betydelig karbonavtrykk.»

Smith GR, Bello C, Bialic-Murphy L, Clark E, Delavaux CS, Fournier de Lauriere C, et al. (2024) Ten simple rules for using large language models in science, version 1.0. PLoS Comput Biol 20(1): e1011767. <https://doi.org/10.1371/journal.pcbi.1011767>

Ten simple rules for using large language models in science



- Smit m.fl. foreslår et sett med 10 enkle regler for bruk av LLM-er i forskning, hentet fra deres egen eksperimentering som forsiktig optimistiske miljøforskere som er nysgjerrige på nye verktøy for å effektivisere forskning.
- Listen er forankret i deres ekspertise som forskere og erfaring som sluttbrukere av LLM-er (GPT spesifikt), ikke som AI-utviklere.

Smith GR, Bello C, Bialic-Murphy L, Clark E, Delavaux CS, Fournier de Lauriere C, et al. (2024) Ten simple rules for using large language models in science, version 1.0. PLoS Comput Biol 20(1): e1011767. <https://doi.org/10.1371/journal.pcbi.1011767>

Box 1—An LLM suggests 10 simple rules governing its own usage in science

Here, we ask ChatGPT to generate a list of 10 simple rules for using LLMs in science in order to illustrate its ability to provide detailed, human-like responses to short prompts.

Q1: Can you provide me with a list of ten simple rules for using large language models in science?

R1 (ChatGPT 3.5): Certainly! Here are ten simple rules for using large language models in science:

1. Understand the limitations: Recognize that large language models like GPT-3 are powerful but not infallible. They may produce accurate-sounding information, but they can also generate false or biased content. Exercise caution and critical thinking.
2. Validate and verify: Always cross-reference information generated by language models with reliable sources. Verify the accuracy of claims, check for consistency, and corroborate the findings through empirical evidence or expert opinions.

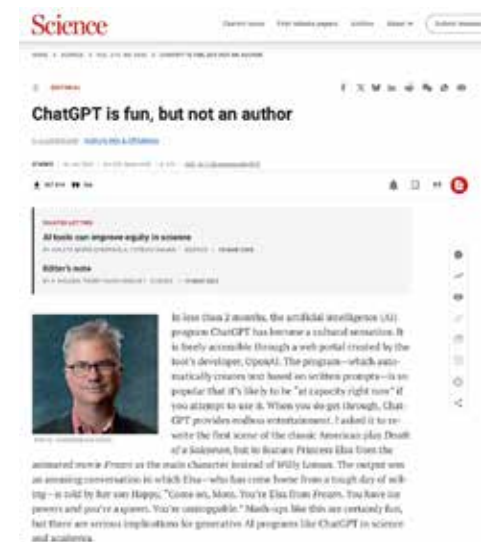
Ten simple rules for using large language models in science



1. Følg reglene til tidsskriftet som artikkelen skal publiseres i.
2. Skisser relevante risikoer før LLM-bruk.
3. Unngå plagiering.
4. Respekter konfidensialitet.
5. Bekreft sannheten til innhold generert av en LLM.
6. Bruk en LLM for å utføre et mer inkluderende datasøk.
7. Bruk en LLM for å oppsummere innhold.
8. Bruk en LLM til å avgrense skriftlig engelsk i formell og uformell kommunikasjon.
9. Bruk en LLM for å forbedre vitenskapelig koding.
10. Bruk en LLM for å starte din vitenskapelige skriveprosess.

1. Følg reglene til tidsskriftet som artikkelen skal publiseres i

- Det er viktig å konsultere og følge en oppdatert versjon av reglene for måltidsskriftet før du bruker en LLM til forskning.
- Etter hvert som disse verktøyene øker i popularitet, vil tidsskrifter sannsynligvis gi eksplisitte retningslinjer for hva de anser som akseptabel eller uakseptabel bruk i innsendte manuskripter.
- Flere tidsskrifter har allerede gitt uttalelser om dette emnet, men disse retningslinjene kan oppdateres og endres etter hvert som forskere i økende grad eksperimenterer med LLM-er og oppdager nye bruksområder.



[DOI: 10.1126/science.adg7879](https://doi.org/10.1126/science.adg7879)

2. Skisser relevante risikoer før LLM-bruk

- Fordi LLM-bruk ikke bare kan føre til nytte, men også til skade, bør forskere som ønsker å bruke en LLM i prosjektet sitt også først skissere relevante risikoer. Basert på en vurdering av risikoen ved et bestemt brukstilfelle, kan det velges en tilnærming som maksimerer nytte og minimerer potensiell skade.
- For noen risikoer kan avbøtende tiltak være enkelt. For eksempel kan feilaktig LLM-generert innhold fanges opp og korrigeres ved grundig og kritisk korrekturlesing (regel 5).
- Imidlertid er andre veldokumenterte LLM-risikoer for komplekse til å være mottagelige for fullstendig avbøtende tiltak av en enkelt forsker eller team.
- I disse tilfellene kan en forsker i stedet kvalitativt evaluere i hvilken grad prosjektet deres kan forverre problemet og gjøre modifikasjoner for å begrense den spesifikke negative effekten av arbeidet deres.

3. Unngå plagiering

- Fordi generativ AI byr på nye utfordringer for juridiske rammeverk rundt opphavsrett og immaterielle rettigheter, må forskere proaktivt sikre at deres LLM-bruk ikke anses som plagiat i sin relevante kontekst.
- Faktisk, utover gjennomsliktig uetisk bruk av LLM-er, for eksempel for å bekjempe programvare for plagiatdeteksjon, kan subtile problemer rundt plagiering også oppstå i daglig bruk av LLM-er i forskning.
- Så lenge en LLM ikke regnes som en juridisk person, vil ikke-attribuert bruk av teksten de produserer ikke se ut til å være plagiat under de fleste gjeldende juridiske definisjoner.
- Det kan imidlertid være uakseptabelt av andre grunner – for eksempel direkte inkorporering av tekst fra GPT i et manuskript bryter for øyeblikket med noen tidsskriftregler, noe som gjør dette til en uakseptabel bruk (regel 1).

4. Respekter konfidensialitet

- I motsetning til bruk av standard analyseverktøy, representerer deling av konfidensielle data eller informasjon med en LLM et potensielt kontraktsbrudd og må unngås med mindre eksplisitt tillatelse er innhentet.
- Dette er viktig å merke seg fordi forskere som ønsker å feilsøke datakode (regel 6), oppsummere innhold (regel 7) eller forbedre manuskriptskriving (regel 10) kanskje ønsker å dele kode, data eller skriving med en LLM.

5. Bekreft sannheten til innhold generert av en LLM

- Innhold generert av en LLM må alltid dobbeltsjekkes av en fagekspert for faktanøyaktighet.
- LLM-er er velkjente for å «hallusinere» – med andre ord å gi tilsynelatende riktige, men falske svar – spesielt som svar på komplekse spørsmål.
- Problematisk nok kan modellens veltalenhet føre til at brukere overser falske utsagn eller referanser.
- Dette problemet er et eksempel på «halo-effekten», en velkjent kognitiv skjevhet der et individs spesifikke positive egenskaper (f.eks. veltalenhet) feilaktig tolkes som bevis på en generell «godhet» som strekker seg utover disse spesifikke egenskapene.
- Når det gjelder en LLM, kan kvaliteten på presentasjonen ubevisst antas å samsvare med innholdet.

Box 2—Examples of LLM hallucination

Here, we request factual, referenced information from ChatGPT. In response, we receive apparently convincing answers which actually contain subtle, potentially misleading inaccuracies.

Q1: Provide a referenced estimate of the amount of vegetation biomass on planet Earth.

R1 (ChatGPT3.5): As of my knowledge cutoff in September 2021, estimating the exact amount of vegetation biomass on planet Earth is challenging due to the vastness and complexity of ecosystems. However, scientists have made various estimations using different methods and data sources. It's important to note that these estimates can vary depending on the specific criteria and assumptions used in the calculations.

One study published in the journal Nature in 2020 provides an estimate of global vegetation biomass based on satellite observations and ground measurements. The research suggests that the Earth's total vegetation biomass is approximately 450 billion metric tons (gigatons). This includes all the vegetation on land, including forests, grasslands, shrublands, and other types of vegetation.

Here is the reference for the study:

Here is the reference for the study:

Zhu, Z., et al. (2020). Greening of the Earth and its drivers. *Nature*, 597(7871), 405-409. doi: [10.1038/s41586-020-2474-6](https://doi.org/10.1038/s41586-020-2474-6).

Please keep in mind that new studies and advancements in remote sensing technologies may have provided more accurate estimates since my last update. It is advisable to consult recent scientific literature for the most up-to-date information on the subject.

Inaccuracies: This is a good example of a seemingly correct answer with some critical mistakes. The paper “Greening of the Earth and its drivers” was indeed published by Zhu and colleagues [23], but in *Nature Climate Change*, 2016 (doi: [10.1038/nclimate3004](https://doi.org/10.1038/nclimate3004)), not in *Nature*. Moreover, in Zhu and colleagues (2016), no estimate was provided of Earth’s total vegetation biomass. The estimate of 450 Gt C can instead be traced to Bar-On and colleagues (2018) [24].

6. Bruk en LLM for å utføre et mer inkluderende datasøk

- En LLM kan hjelpe forskere med datainnsamling for vitenskapelige metaanalyser eller synteser ved å samle "grå litteratur"-ressurser – rapporter utgitt av statlige og ikke-statlige organisasjoner i stedet for vitenskapelige tidsskrifter.
- Dette er nyttig fordi selv om det finnes etablerte metoder for systematisk å søke i vitenskapelig litteratur ved hjelp av databaser som Web of Science, kan disse tilnærmingene neglisjere grå litteratur.
- Denne oversikten er spesielt problematisk i anvendte felt som bevaring og økologisk restaurering, der innsikten til utøvere på bakken er av sentral betydning.

Box 3—Performing an LLM-assisted grey literature search

Here, we use ChatGPT to generate a preliminary list of organisations that work in some aspect of forest restoration and which may have published reports online detailing the impact of their work.

Q1: Please generate a list of forest restoration organizations that publish data on the biodiversity and carbon outcomes of their work.

R1: Certainly! Here is a list of forest restoration organizations that publish data on the biodiversity and carbon outcomes of their work:

1. **The Nature Conservancy:** The Nature Conservancy is a global environmental organization that works to protect and restore critical habitats, including forests. They publish regular reports on their conservation efforts, including data on biodiversity and carbon sequestration.

2. **Trees for the Future:** Trees for the Future is a non-profit organization that works with local communities to plant trees and restore degraded lands. They publish annual reports that detail the impact of their work, including data on biodiversity and carbon sequestration.

Q2: Have these organizations published reports on their work online?

R2: Yes, these organizations typically have published reports on their work available online. While I don't have real-time access to the internet, I can provide you with general information about where to find these reports. You can visit the websites of these organizations and look for sections such as "Publications," "Reports," or "Impact" to find detailed information about their forest restoration projects, including data on biodiversity and carbon outcomes.

For example:

1. The Nature Conservancy: Visit their website at [nature.org](https://www.nature.org) and explore their "Our Work" or "Publications" sections.
2. Trees for the Future: Visit their website at [trees.org](https://www.trees.org) and look for sections like "Impact" or "Resources" for access to their reports.
3. Rainforest Alliance: Visit their website at [rainforest-alliance.org](https://www.rainforest-alliance.org) and search for sections such as "Our Work" or "Publications" to find their reports.
4. World Wildlife Fund: Visit their website at [wwf.org](https://www.wwf.org) and

Q3: Please turn this into a table that facilitates finding and extracting data from the reports for these organizations, and 5 other relevant organizations.

R3: Certainly! Here is a table that includes the organizations mentioned earlier, as well as five additional forest restoration organizations, along with relevant information that can facilitate finding and extracting data from their reports:

[Table not shown for brevity, see Q4 for follow up that adds column on geographical scope]

Note that this table is not exhaustive and there may be additional relevant sections and types of reports available on each organization's website. However, this should provide a good starting point for finding and extracting data from their reports.

Q4: Please add a column to this table that describes the geographic scope of each organization.

R4: Certainly! Here is an updated table that includes the organizations mentioned earlier, along with five additional forest restoration organizations, relevant sections on their website, types of reports they publish, and the geographic scope of their work:

Note that the geographic scope of each organization's work may vary over time, and this information is accurate as of the knowledge cutoff date for my training.

All organisations identified were relevant and many publish data or summaries of their impacts in annual reports somewhere on their websites. An extended list of additional organisations was provided with 1 additional prompt, and it is easy to export the list in a standard format (e.g.,.csv). However, the list does contain several inaccuracies (see **Rule 5**). For example, the "Relevant Sections" column often listed website sections that did not actually exist.

Organization Name	Website	Relevant Sections	Types of Reports	Geographical Scope
The Nature Conservancy	https://www.nature.org/	Our Work, Publications	Annual reports, Impact reports, Biodiversity and Carbon reports	Global
Trees for the Future	https://trees.org/	Impact, Resources	Annual reports, Impact reports, Biodiversity and Carbon reports	Africa, Asia, Latin America, and the Caribbean
Rainforest Alliance	https://www.rainforest-alliance.org/	Our Work, Publications	Annual reports, Impact reports, Biodiversity and Carbon reports	Global
World Wildlife Fund	https://www.worldwildlife.org/	Conservation, Publications	Annual reports, Impact reports, Biodiversity and Carbon reports	Global
Forest Stewardship Council	https://fsc.org/	Resources, Publications	Certification reports, Biodiversity and Carbon reports	Global
Eden Reforestation Projects	https://edenprojects.org/	Impact, Stories	Annual reports, Impact reports, Biodiversity and Carbon reports	Africa, Asia, Haiti, Madagascar, and Central America
The Arbor Day Foundation	https://www.arborday.org/	Programs, Reports	Annual reports, Impact reports, Biodiversity and Carbon reports	United States and internationally
American Forests	https://www.americanforests.org/	Our Work, Publications	Annual reports, Impact reports, Biodiversity and Carbon reports	United States and internationally
Trees, Water & People	https://www.treeswaterpeople.org/	Our Programs, Reports	Annual reports, Impact reports, Biodiversity and Carbon reports	Latin America and the United States
Conservation International	https://www.conservation.org/	Our Work, Publications	Annual reports, Impact reports, Biodiversity and Carbon reports	Global

<https://doi.org/10.1371/journal.pcbi.1011767.t001>

Smith GR, Bello C, Bialic-Murphy L, Clark E, Delavaux CS, et al. (2024) Ten simple rules for using large language models in science, version 1.0. PLOS Computational Biology 20(1): e1011767. <https://doi.org/10.1371/journal.pcbi.1011767>
<https://journals.plos.org/ploscompbiol/article?id=10.1371/journal.pcbi.1011767>

PLOS COMPUTATIONAL BIOLOGY

7. Bruk en LLM for å oppsummere innhold

- For å hjelpe forskere med å utnytte tiden sin mer effektivt, kan en LLM raskt generere korte sammendrag basert på langt inndatainnhold.
- Dette er nyttig fordi det som forsker er viktig å lese ny litteratur regelmessig, men å balansere denne oppgaven med pliktene til vår egen forskning, undervisning og veiledning kan være utfordrende.
- Likevel kan forskere som ikke klarer å finne en god balanse og ikke ofte undersøker ny litteratur, gå glipp av funn som kan hjelpe til med å informere og forme deres egne forskningsprogrammer, eller som vil være relevante for klasser de underviser.

Box 4—Using GPT to summarise scientific literature

Here, we use GPT-4 (via NewBing) to summarise a recently published open-access scientific paper.

Q1: I want you to act as a scientific paper summarizer. I will give you the URL of a scientific article. I want you to reply with a bold title for each section including the general information of this paper, the background of this research, the addressed question and hypothesis, the key findings, and major contributions. The summary of each section should be logical, concise, clear, and informative. Make sure the summary includes most of the key information and can grip the audience. The URL is:

<https://www.nature.com/articles/s41467-022-30888-2>.

R1 (GPT-4 via NewBing): General information: The article is titled “Global relationships in tree functional traits” and it was published in Nature Communications on 08 June 2022. The authors are Daniel S. Maynard, Lalasia Bialic-Murphy, Constantin M. Zohner, Colin Averill, Johan van den Hoogen, Haozhi Ma, Lidong Mo, Gabriel Reuben Smith, Alicia T. R. Acosta, Isabelle Aubin, Erika Berenguer, Coline C. F. Boonman, Jane A. Catford, Bruno E. L. Cerabolini, Arildo S. Dias, Andrés González-Melo, Peter Hietz, Christopher H. Lusk, Akira S. Mori, Ülo Niinemets, Valério D. Pillar, Bruno X. Pinho, Julieta A. Rosell, Frank M. Schurr and Thomas W. Crowther¹.

Background: Trees are subject to unique physiological, mechanical and ecological pressures due to their massive energetic investments in woody support structures. Functional traits are the physiological and morphological features that determine how trees interact with, influence and respond to their environment. Understanding the patterns and trade-offs in tree functional traits is challenging because of the diversity and complexity of tree form and function.

Question and hypothesis: The authors aim to quantify the multidimensional relationships in tree trait expression across 18 functional traits encompassing leaf, seed, bark, wood, crown and root characteristics. They hypothesize that these traits reflect two main axes: one related to leaf economics (i.e., the trade-off between resource acquisition and conservation) and the other related to tree size and competition for light (i.e., the trade-off between growth and survival). They also expect that these axes show environmental convergence across temperature, moisture and elevation gradients.

Key findings: The authors find that nearly half of trait variation is captured by the two hypothesized axes¹. They also find that these axes exhibit correlated responses to environmental factors¹. Moreover, they identify eight distinct clusters of traits that capture the full dimensionality of trait space. These clusters reflect different aspects of tree form and function such as drought tolerance, shade tolerance, dispersal ability and nutrient cycling.

Major contributions: The authors identify a core set of traits needed to quantify global patterns in functional biodiversity. They also contribute to the fundamental understanding of the functioning of forests worldwide by revealing how tree traits shape their water, nutrient and light economies¹. They suggest that their trait framework can be used to improve forest management and conservation strategies.

8. Bruk en LLM til å avgrense skriftlig engelsk i formell og uformell kommunikasjon (1/2)

- En LLM kan legge til rette for klarere og mer idiomatisk bruk av engelsk, *lingua franca* i det internasjonale vitenskapelige samfunnet.
- Dette kan være spesielt nyttig for forskere som snakker engelsk som andrespråk, som kan møte unike hindringer i det engelskdominerte miljøet av global vitenskap.
- Dessverre har fagfellevurdering en tendens til å være til ulempe for forfattere fra utviklingsland og ikke-engelsktalende land, og forskere som snakker engelsk som andrespråk får noen ganger uhøflige, nedslående kommentarer fra fageksperter.

8. Bruk en LLM til å avgrense skriftlig engelsk i formell og uformell kommunikasjon (2/2)

- Sammen med tidsskriftadopsjon av dobbeltblind vurderingspraksis, kan en LLM bidra til å løse dette problemet ved å redigere manuskripter og annen kommunikasjon for å sikre at engelsk brukes idiomatisk, noe som reduserer mulighetene for skjevhet fra referanser.
- Det er viktig at før du bruker en LLM på denne måten, må retningslinjer for tidsskrift konsulteres (regel 1). Hvis en slik bruk ikke er eksplisitt tillatt i henhold til tidsskriftets retningslinjer, foreslår vi at modellens rolle som manuskriptredaktør bør kommuniseres direkte til tidsskriftredaktører.

9. Bruk en LLM for å forbedre vitenskapelig koding

- En LLM kan hjelpe forskere med å skrive kode [datakode] raskere og mer effektivt, ettersom LLM-er utmerker seg ikke bare i naturlig språkbehandling, men også i programmeringsoppgaver på tvers av forskjellige programmeringsspråk.
- Modeller kan brukes til å generere innledende kodebiter eller for å forbedre strukturen, effektiviteten og lesbarheten til eksisterende kode.
- Dette kan for eksempel gjøres ved å (i) ordne kode i logiske seksjoner som funksjoner, klasser og moduler; (ii) refaktorere kode ved å foreslå bedre navnekonvensjoner, eliminere kodeduplisering og forbedre den generelle kodelesbarheten; (iii) foreslå alternative algoritmer eller datastrukturer som kan forbedre effektiviteten og ytelsen til koden; eller (iv) skrive tydeligere kodemerknader.

10. Bruk en LLM for å starte din vitenskapelige skriveprosess

- Forskere kan overvinne skrivesperre og "blank page-angst" (writer's block - jernteppe) ved å be en LLM om å produsere en disposisjon for deres vitenskapelige artikkel.
- Dette kan være nyttig fordi skriving er et stort hinder for mange forskere, med noen av de største hindringene som utgjør en utfordring i begynnelsen av et prosjekts skrivefase.
- I tillegg til å gi et utgangspunkt, kan en LLM hjelpe til med å organisere underseksjoner, tilby relevante eksempler og finjustere den generelle flyten i manuskriptet.
- På denne måten kan en LLM fungere for forskere som en rask og billig korrekturleser og tekstredigerer.

Box 6—Using GPT to produce a tailored paper outline

Q1: Give me a 4-heading structure for a research article. Topic: Climate change effects on biodiversity in the tropics Context: Ecology Language: Academic Tone: Formal.

R1: I. Introduction

A. Background on climate change and its global impact

B. Importance of studying climate change effects on biodiversity in the tropics

C. Research question and objectives

II. Literature Review

A. Overview of climate change impacts on biodiversity worldwide

B. Key studies and findings on climate change effects on biodiversity in the tropics

C. Knowledge gaps and research limitations

III. Methodology

III. Methodology

A. Study area and data collection

B. Variables and measurements

C. Statistical or analytical methods used

IV. Results and Discussion

A. Presentation of findings related to climate change effects on biodiversity in the tropics

B. Discussion of the results in the context of existing literature

C. Implications and potential future directions

V. Conclusion

A. Summary of the main findings

B. Significance of the research and its contribution to the field

C. Recommendations for future research and policy implications

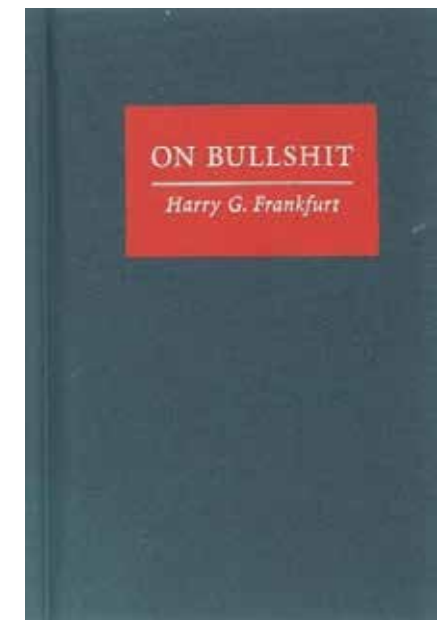
Kritikere sier:

ChatGPT is bullshit

ChatGPT is bullshit [sludder]

- I den senere tid har det vært stor interesse for store språkmodeller: maskinlæringssystemer som produserer menneskelignende tekst og dialog.
- Anvendelser av disse systemene har vært plaget av vedvarende unøyaktigheter i produksjonen; disse kalles ofte "AI-hallusinasjoner".
- Vi argumenterer for at disse usannhetene, og den generelle aktiviteten til store språkmodeller, bedre forstås som bullshit i den forstand som ble utforsket av Frankfurt (On Bullshit, Princeton, 2005): modellene er likegyldige til sannheten i resultatene deres.

Hicks, M.T., Humphries, J. & Slater, J. ChatGPT is bullshit. *Ethics Inf Technol* 26, 38 (2024). <https://doi.org/10.1007/s10676-024-09775-5>



ChatGPT is bullshit



- Vi skiller mellom to typer bullshit, som vi kaller «**hard**» og «**myk**» (soft) bullshit, der førstnevnte krever et aktivt forsøk på å lure leseren eller lytteren med hensyn til virksomhetens natur, og sistnevnte bare krever mangel på omsorg for sannheten.
- Vi argumenterer for at i det minste er resultatene til LLM-er som ChatGPT **myk** bullshit: bullshit – det vil si tale eller tekst produsert uten bekymring for sannheten – som er produsert uten noen intensjon om å villedde publikum om fortellerens holdning til sannheten.
- Vi antyder også, mer kontroversielt, at ChatGPT faktisk kan produsere **hard** bullshit: hvis vi ser på det som å ha intensjoner (for eksempel i kraft av hvordan det er utformet), så kvalifiserer det faktum at det er designet for å gi inntrykk av bekymring for sannhet det som et forsøk på å villedde publikum om sine mål, mål eller agenda.

Hicks, M.T., Humphries, J. & Slater, J. ChatGPT is bullshit. *Ethics Inf Technol* 26, 38 (2024).
<https://doi.org/10.1007/s10676-024-09775-5>

ChatGPT is bullshit



- Store språkmodeller, og andre AI-modeller som ChatGPT, gjør betydelig mindre enn hva menneskelige hjerner gjør, og det er ikke klart om de gjør det de gjør på samme måte som oss.
- Den mest åpenbare forskjellen mellom en LLM og et menneskesinn involverer målene til systemet.
- Mennesker har en rekke mål og atferd, hvorav de fleste er utenomspråklige: vi har grunnleggende fysiske ønsker, for ting som mat og næring; vi har sosiale mål og relasjoner; vi har prosjekter; og vi skaper fysiske objekter.
- Store språkmodeller tar ganske enkelt sikte på å gjenskape menneskelig tale eller skrift. Dette betyr at deres primære mål, i den grad de har en, er å produsere menneskelignende tekst. De gjør det ved å estimere sannsynligheten for at et bestemt ord vil dukke opp neste gang, gitt teksten som har kommet før.

Hicks, M.T., Humphries, J. & Slater, J. ChatGPT is bullshit. *Ethics Inf Technol* 26, 38 (2024).
<https://doi.org/10.1007/s10676-024-09775-5>

ChatGPT is bullshit



- Eksempler på dette er allerede mange, for eksempel utarbeidet en advokat nylig sin prosedyre [brief] (utvalg av saksdokumenter som skal legges frem i retten) ved hjelp av ChatGPT og oppdaget til sin ergrelse at de fleste av de siterte sakene ikke var ekte; som dommer P. Kevin Castel sa det, produserte ChatGPT en tekst fylt med "falske rettsavgjørelser, med falske sitater og falske interne sitater".
- På samme måte, da informatikkforskere testet ChatGPTs evne til å hjelpe til med akademisk skriving, fant de ut at den var i stand til å produsere overraskende omfattende og noen ganger til og med nøyaktig tekst om biologiske emner gitt de riktige instruksjonene.
- Men da de ble bedt om å fremlegge bevis for sine påstander, «ga de fem referanser som dateres til tidlig på 2000-tallet. Ingen av de oppgitte papirtitlene eksisterte, og alle oppgitte PubMed-ID-er (PMID-er) var av forskjellige ikke-relaterte artikler».

Hicks, M.T., Humphries, J. & Slater, J. ChatGPT is bullshit. *Ethics Inf Technol* 26, 38 (2024). <https://doi.org/10.1007/s10676-024-09775-5>

ChatGPT is bullshit



- Disse feilene kan "snøball": når språkmodellen blir bedt om å gi bevis for eller en dypere forklaring på en falsk påstand, sjekker den sjelden seg selv; i stedet produserer den selvsikkert flere falske, men normalt klingende påstander.
- Nøyaktighetsproblemet for LLM-er og andre generative AI-er blir ofte referert til som problemet med "AI-hallusinasjoner": chatboten ser ut til å hallusinere kilder og fakta som ikke eksisterer.
- Disse unøyaktighetene omtales som "hallusinasjoner" i både tekniske og populære sammenhenger.

Entusiastene sier:

**ChatGPT kan anvendes
til kvalitetskontroll**



Å være fagfelle på artikler tar mye tid og er ubetalt jobb. Tidsskrifter melder at det blir stadig vanskeligere å finne fagfeller til å vurdere artikler.

Kan kunstig intelligens være til hjelp?

Nå spytter den britiske regjeringen i 4,8 millioner pund (68 millioner kroner) til KI-prosjekter som skal effektivisere forskningsprosesser, skriver [The Telegraph](#).

I flere prosjekter skal KI bidra til fagfelleevaluering.

Tanken er at dette skal frigjøre tid for forskerne, sier forskningsminister Sir Patrick Vallance, ifølge The Guardian.

Fagfellene skal ikke erstattes av taleroboter, men de skal brukes som hjelp i prosessen.



Nå skal KI gjøre fagfelleevaluering

- University of Sheffield skal teste om KI-modeller som ChatGPT kan brukes til kvalitetskontroll av vitenskapelige artikler. Det betyr at KI kan analysere innholdet og gi tilbakemeldinger, men ikke ta endelige beslutninger.
- University of Bath skal teste en totrinns fagfelleevalueringsprosess, der KI vil analysere utfallet fagfelleevalueringen. Forskerne skal bruke KI til å undersøke hvordan de utfører fagfelleevaluering og om det kan ha noen konsekvenser for forskningskvaliteten.
- Universitetene samarbeider med storforlaget Sage Publishing og vitenskapsakademiet Royal Society om disse prosjektene.
- I tillegg skal blant annet University of Aberdeen utvikle et KI-verktøy som skal hjelpe forskere med å avdekke dårlige eller misvisende studier, slik at disse ikke brukes i videre forskning.
- University College London samarbeider med den britiske regjeringen for å bruke KI i analyser av hvordan forskningsresultater brukes i politiske beslutninger.



Khrono's KI-skole

<https://www.khrono.no>

[illegible][illegible]

NB! All tekst fra KI-skolen er
kopiert fra originalen

FORSKERHÅNDBOKEN

KI-SKOLEN, DEL 1

Dette kan du bruke KI til i hverdagen

Hvordan kan du bruke kunstig intelligens (KI) til å forenkle og fornye hverdagen? Det får du svar på i første del av Khronos KI-skole med professor Morten Goodwin.



Følg Khronos KI-skole som har som mål å vise hvordan kunstig intelligens kan forenkle og fornye den akademiske hverdagen. Illustrasjonen er laget av ChatGPT, som fikk følgende instruks: Tegn en robot som både er en forsker og en underviser på et universitet. La illustrasjonen være inspirert av gammel sovjetisk propaganda, men uten sovjetiske symboler. Illustrasjon: ChatGPT



Professor
Morten Goodwin
Universitetet i Agder

KI kan bli en assistent for forskningsarbeid. Her er noen konkrete eksempler:

- **Automatisk språkvask:** Språkmodeller kan forbedre og polere akademisk tekst, slik at du kan fokusere på innholdet.
- **Dataanalyse:** KI kan analysere store datasett, identifisere mønstre og til og med generere hypoteser basert på trender.
- **Tekstgenerering:** Sliter du med å skrive introduksjoner eller sammendrag? KI kan lage utkast som gir deg et springbrett.

Å «snakke» med språkmodeller skjer gjennom tekstlige instruksjoner, kjent som *prompts* eller *instrukser*. En instruks er ganske enkelt beskjeden du gir til en språkmodell, som spørsmålet du stiller eller instruksjonen du gir.

Måten du formulerer den på, påvirker direkte kvaliteten på svarene du mottar.

Dette ligner på hvordan vi kommuniserer oss mennesker imellom, for eksempel fra en elev til en lærer. Et dårlig eller lite gjennomtenkt spørsmål fører ofte til et middelmådig svar. Still et klart og presist spørsmål, og sjansen for et godt svar øker betraktelig.

Khrono

NyheterDebuttStudentNytt om folk+

KI-SKOLEN, DEL 2

Slik «snakker» du med kunstig intelligens

Hvordan hvordan går man egentlig fram når man skal «snakke» med en kunstig intelligens (KI)? Nøkkelen ligger i gode prompts, eller instrukser, forklarer professor Morten Goodwin i andre del av Khronos KI-skole.



Illustrasjonen er laget av ChatGPT, som fikk følgende instruks: Tegn en illustrasjon av en forsker som snakker med kunstig intelligens. Det hele skal se veldig profesjonelt ut. Det skal være formler, studenter, diplomer på veggen og slikt. Bildet skal være et superhelt tegneserieaktig preg. Illustrasjonen skal være i 16:9

Oppfølgende instruks: Det skal være en helt vanlig forsker som snakker med kunstig intelligens. Illustrasjon: ChatGPT

Professor
Morten Goodwin
Universitetet i Agder

Publisert 22.01.2025 - 16:18 Sist oppdatert 09.03.2025 - 19:17



Gi nok kontekst

- Tenk deg at du har et datasett innen klimaendringer og vil undersøke hvilke faktorer som kan ha størst betydning for temperaturøkning. Du kan be ChatGPT foreslå flere mulige hypoteser, basert på eksisterende litteratur.
- For eksempel:

Jeg forsker på klimaendringer i arktiske områder. Her er en kort oversikt over data fra de siste ti årene, inkludert temperatur, istykkelse og nedbørsmengder [legg ved de viktigste tallene]. Kan du foreslå minst tre konkrete forskningshypoteser som kan forklare endringene i gjennomsnittlig temperatur? Begrunn hver hypotese med relevant vitenskapsteori.

Gi nok kontekst

- Et annet eksempel kan være for sammendrag av en artikkel. KI kan hjelpe med å oppsummere lange tekster eller sammensette «state of the art» på et felt.
- Mangelfull instruks:

Oppsummer denne artikkelen om hjertesykdom.

Gi nok kontekst

- Hvorfor mangelfull? Fordi instruksen er for vag, og det er lite kontekst. Du får kanskje en generalisert oppsummering.
- God instruks:

Denne artikkelen om hjertesykdom diskuterer risikofaktorer og forebyggende tiltak. Lag en punktvis oppsummering på maks 200 ord, der du tydelig angir hvordan funnene kan overføres til forebyggende helsearbeid blant eldre pasienter.

Instrukser på norsk eller engelsk?

- Ved å bruke engelsk som instruksjonsspråk, utnytter du modellens tilgang til omfattende tekstkilder, noe som gjør det spesielt effektivt for komplekse spørsmål og oppgaver som krever grundige analyser og dypdykk i detaljene.
- På norsk er datagrunnlaget mindre omfattende, noe som kan føre til enklere svar med færre nyanser.
- ChatGPT presterer merkbart bedre på engelsk enn på norsk, bedre på bokmål enn på nynorsk, og betydelig svakere på samiske språk. Dette skyldes hovedsakelig mengden tilgjengelig treningsdata: Jo mer tekst modellen har tilgang til, desto mer presise, detaljerte og nyanserte blir svarene.
- Instrukser på norsk gir derimot ofte svar som er bedre tilpasset norske forhold, enten det handler om norsk historie eller spesifikke språklige uttrykk.

KI-SKOLEN, DEL 3

Her er de viktigste og nyttigste KI-verktøyene

Bør du bruke ChatGPT, Claude eller Gemini — eller DeepSeek? Hvilke andre KI-verktøy finnes, og hva er de forskjellige verktøyene best på? Det forklarer professor Morten Goodwin i tredje del av Khronos KI-skole.



Illustrasjonen er laget av ChatGPT, som fikk følgende instruks: Tegn en forsker som har mange AI-verktøy tilgjengelig. Superhelt tegneserietegning. 16:9. Oppdatert instruks: La det være en kvinnelig superhelt. Illustrasjon: ChatGPT



Professor Universitetet i Agder
Morten Goodwin
morten.goodwin@uia.no

Publisert 06.02.2025 - 19:35 Sist oppdatert 09.02.2025 - 16:43



Det finnes et bredt spekter av språkmodeller, og felles for dem alle er at de bygger på samtaleprinsippet: Du kommuniserer med dem ved å skrive instruksjoner, eller såkalte «prompts», og får svar tilbake. [ChatGPT](#) er språkmodellem som utvilsomt har fått mest oppmerksomhet, men det finnes mange andre modeller som hver har sine unike fordeler.

Konkurransen mellom språkmodeller er knallhard, med alternativer som [Anthropic Claude](#), [Google Gemini](#) og den rykende ferske [DeepSeek R1](#) er blant de mest kjente. Til tross for at de fungerer på lignende vis, er kampen om å være den beste intens.

KI-verktøy

Språkmodeller:

- [ChatGPT](#): En avansert språkmodell fra OpenAI som brukes til alt fra tekstforståelse til analyse.
- [Claude](#): En språkmodell fra Anthropic kjent for sin sikkerhet og evne til lange kontekster.
- [Gemini](#): En språkmodell fra Google med sterke multimodale evner.
- [DeepSeek R1](#) : En kinesiskutviklet open-source KI-modell.
- [Grok](#): Elon Musks KI-modell som prioriterer frihet og lite innholdsbegrensninger.
- [CoPilot](#): Flere språkmodeller fra Microsoft blant annet for assistanse i arbeidsoppgaver.
- [NORA.LLM](#): Norsk KI-modell fra UiO utviklet med vekt på skandinavisk språk.
- [NORLLM](#): NTNU sin bidrag til norskutviklet KI med fokus på forskning og utdanning.
- [LM Studio](#): En plattform for å kjøre språkmodeller lokalt på egen maskin.

KI-verktøy

RAG-modeller:

- [Perplexity](#): Kombinerer KI og sanntidssøk for mer pålitelige svar.
- [Storm](#): Stanford-utviklet KI som genererer artikler basert på henting og bearbeiding av data.

Innsikt i artikler:

- [Notebook LM](#) : KI fra Google for dyp analyse av tekst og forskningsartikler.
- [Humata](#): Analyserer dokumenter og trekker ut nøkkelinformasjon.
- [Afforai](#) : KI-verktøy for rask innsikt i store tekstmengder.

Microsoft – personlig KI-assistent

- [Microsoft CoPilot](#) er ikke én enkelt språkmodell, men et samspill av flere modeller, mange av dem utviklet i samarbeid med og basert på ChatGPT. Når du bruker CoPilot, er det i praksis samme teknologi som driver ChatGPT. Programvareutviklere er allerede kjent med GitHub CoPilot, som har blitt en uunnværlig hjelp for å skrive programkode mer effektivt.
- En av de store fordelene med Microsoft CoPilot er hvordan den kan integreres sømløst i brukerens system. Med riktig oppsett kan CoPilot få tilgang til e-postene dine, lære seg din skrivestil og svare på vegne av deg, direkte i Outlook.
- Den kan også gi deg avansert støtte i andre programmer som Excel, Word og PowerPoint.
- For eksempel kan den hjelpe deg med å generere komplekse formler i Excel, skrive profesjonelle dokumenter i Word eller lage overbevisende presentasjoner i PowerPoint, alt med en personlig tilpasning basert på dine vaner og preferanser.

Chatbot Arena LLM Leaderboard: Community-driven Evaluation for Best LLM and AI chatbots

[Discord](#) | [Twitter](#) | [小红书](#) | [Blog](#) | [GitHub](#) | [Paper](#) | [Dataset](#) | [Kaggle Competition](#)

Chatbot Arena is an open platform for crowdsourced AI benchmarking, developed by researchers at UC Berkeley [SkyLab](#) and [LMArena](#). With over 1,000,000 user votes, the platform ranks best LLM and AI chatbots using the Bradley-Terry model to generate live leaderboards. For technical details, check out our [paper](#).

Chatbot Arena thrives on community engagement — cast your vote to help improve AI evaluation!

New Launch! Chat with online LLMs in the [Search Arena tab](#)!

[Language](#)
[Overview](#)
[Price Analysis](#)
[WebDev Arena](#)
[Vision](#)
[Text-to-Image](#)
[Copilot Arena](#)
[Arena-Hard-Auto](#)

Total #models: 218. Total #votes: 2,801,990. Last updated: 2025-03-20.

Code to recreate leaderboard tables and plots in this [notebook](#). You can contribute your vote at [lmarena.ai](#)

Category

Overall

Apply filter

☐ Style Control

☐ Show Deprecated

Overall Questions

#models: 218 (100%) #votes: 2,801,990 (100%)

Rank* (UB)	Rank (StyleCtrl)	Model	Arena Score	95% CI	Votes	Organization	License
1	3	Grok-3-Preview-02-24	1404	+5/-6	9871	xAI	Proprietary
1	1	GPT-4.5-Preview	1398	+5/-6	9871	OpenAI	Proprietary
3	6	Gemini-2.0-Flash-Thinking-Exp-01-21	1382	+5/-4	22004	Google	Proprietary
3	3	Gemini-2.0-Pro-Exp-02-05	1379	+4/-4	19963	Google	Proprietary
3	2	ChatGPT-4o-latest (2025-01-29)	1374	+4/-5	21831	OpenAI	Proprietary
6	4	DeepSeek-R1	1360	+7/-7	12224	DeepSeek	MIT
6	11	Gemini-2.0-Flash-001	1355	+6/-6	17689	Google	Proprietary
6	3	o1-2024-12-17	1351	+4/-3	24350	OpenAI	Proprietary
8	11	Gemma-3-27B-it	1341	+8/-6	6248	Google	Gemma
9	11	Qwen2.5-Max	1340	+5/-4	16493	Alibaba	Proprietary
9	7	o1-preview	1335	+4/-4	33186	OpenAI	Proprietary

Retrieve Augment Generate (RAG)

- En av de største utfordringene med språkmodeller, uansett hvilken du velger, er fenomenet kjent som hallusinasjoner. Det skjer når modellen begynner å generere tekst som virker overbevisende, men som ikke stemmer overens med virkeligheten. Modellen «tror» den leverer sann informasjon, men fortsetter i et sammensurium av feil og oppspinn.
- Slike hallusinasjoner fører til en rekke merkverdige utslag i språkmodeller, men det vi forskere merker mest, er hvordan den fabrikkerer fakta, dikter opp referanser og til og med konstruerer falske sitater og kilder.
- **I akademiske kontekster, hvor nøyaktighet og etterprøvbarhet er avgjørende, er slike hallusinasjoner et betydelig problem.**

Med instruksen «Give me a list of recent papers on deep learning and skin cancer megalomania, and compare it to the expertise of medical doctors» får vi følgende svar:

Recent research has shown significant advancements in the application of deep learning for skin cancer diagnosis, often comparing the performance of artificial intelligence (AI) systems with that of medical professionals. Here's a list of recent papers on this topic:

1. **"Artificial Intelligence is Better than Dermatologists at Diagnosing Skin Cancer"** (2025)

- Published in *Annals of Oncology*
- Key findings: A deep learning convolutional neural network (CNN) outperformed 58 international dermatologists in detecting melanomas and correctly identifying benign moles

2. **"Enhanced skin cancer diagnosis using optimized CNN architecture"** (2024)

- Published in a medical journal
- Key findings: A sophisticated CNN model achieved 97.78% accuracy in skin lesion classification, demonstrating its potential as a robust tool for early detection and classification of skin cancer)

3. **"Experience And Dermoscopy Improve Skin Cancer Diagnosis Accuracy"** (2024)

- Published in *EMJ Dermatology*
- Key findings: While not focused on AI, this study provides a benchmark for human performance, showing that experienced dermatologists using dermoscopy achieved higher accuracy in diagnosing skin cancer compared to primary care physicians

4. **"Human-AI interaction in skin cancer diagnosis: a systematic review"** (2024)

- Published in *Nature*
- Key findings: AI assistance improved diagnostic performance for all clinicians, with the largest improvement seen among non-dermatologists

5. **"The Role of Artificial Intelligence in the Diagnosis of Melanoma"** (2024)

Merk at jeg feilaktig skrev *megalomania* (overdreven tro på egen makt) istedenfor *melanoma* (type hudkreft som oppstår i melanocytene) i instruksene. Perplexity forstod allikevel at det bare hudkreft og ikke stormannsgalskap jeg var ute etter.

Den store fordelene med RAG-modeller er at de lar deg som forsker gå inn og kvalitetssikre innholdet, omtrent på samme måte som du ville lest en godt skrevet oversiktsartikkel innenfor ditt fagfelt. Du kan eksplisitt spørre modellen om hvilke deler av teksten som kommer fra hvilke kilder, og dermed vurdere påliteligheten til hver enkelt del av svaret.

Perplexity gir deg også muligheten til å organisere litteraturen du har funnet. Vi kan fortsette dialogen og spørre den om å sette opp forskningen om hudkreftanalyse i en oversiktlig tabell.

Instruks:

Write it as a table with the accuracy of the algorithm and how many patients were part of each study.

Perplexity svarer:

Here's a table summarizing the accuracy of skin cancer detection algorithms and the number of patients involved in each study:

Study	Accuracy	Number of Patients / Images
Integrated Deep Learning Model [1]	92.27 %	Not specified (HAM10000 dataset used)
Hybrid Dynamic Bayesian Deep Learning [1]	88.95 % - 90.96 %	Two datasets, sizes not specified
Deep Learning on Dermoscopic Images [2]	89.3% ± 1.1% (4 conditions) 94.5% ± 0.9% (benign vs. malignant)	57,536 dermoscopic images
Hybrid U-Net and MobileNet-V33	98.86 %	Not specified
Stanford Medicine AI Study [4]	81.1 % sensitivity 86.1 % specificity	67,000+ evaluations
VGG-16 Model [5]	88 %	Not specified
CNN on ISIC 2018 Dataset [5]	83.2 %	3,533 images
Malignant Melanoma Classification [6]	80.06 %	Not specified

[1] <https://arxiv.org/html/2410.14489v1>

[2] https://jbpe.sums.ac.ir/article_48805.html

[3] <https://www.nature.com/articles/s41598-024-54212-8>

[4] <https://med.stanford.edu/news/all-news/2024/04/ai-skin-diagnosis.html>

[5] <https://pmc.ncbi.nlm.nih.gov/articles/PMC10252190/>

[6] <https://pmc.ncbi.nlm.nih.gov/articles/PMC10571810/>

Retrieve Augment Generate (RAG)

- En annen interessant funksjon er muligheten til å laste opp egne artikler og få modellen til å finne relevant litteratur. Dette bringer oss raskt nærmere en algoritme som kan utarbeide litteraturgjennomgangen i en vitenskapelig artikkel, en oppgave som hittil har vært både tidkrevende og møysommelig.
- Kanskje er dette et steg mot en forskningshverdag der mer tid kan brukes på analyse og diskusjon av egne resultater, mens modellen håndterer sorteringen av faglitteraturen.
- [Storm](#) fra Stanford er en annen RAG-modell som kan generere imponerende Wikipedia-lignende artikler om helt nye temaer.
- Med instruksjonen «Write an article about AI in the public sector, specifically what has been tested in practice» søker den først etter informasjon, og så produserer den en lettest og velformulert tekst om kunstig intelligens i offentlig sektor, en artikkel som like gjerne kunne vært hentet fra Wikipedia.

Raskere innsikt i kompliserte artikler

- KI-verktøy kan også brukes til å forstå og få oversikt over komplekse artikler. [Google Notebook LM](#) er et kraftig hjelpemiddel som bruker Gemini-språkmodellen for forskere og undervisere som arbeider med store mengder tekst og krevende temaer.
- En forsker som analyserer en omfattende samling vitenskapelige artikler, kan laste opp dokumentene til Notebook LM og få automatiske sammendrag av hver enkelt artikkel. Dette sparer tid og gir en rask oversikt over hovedfunnene.
- En lege som gjennomgår medisinske forskningsrapporter, kan få KI til å trekke ut nøkkelpunkter om behandlingsmetoder, bivirkninger og statistiske resultater, slik at ny kunnskap raskere kan omsettes til praksis.

Automatisk fagfellevurdering

- Notebook LM gjør det også mulig å få en automatisk fagfellevurdering, som fremhever både styrker og svakheter ved artikkelen.
- Når jeg skriver en vitenskapelig artikkel, kan jeg laste den opp og umiddelbart få innsikt i svakheter og forbedringspunkter. Jeg trenger ikke lenger vente i månedsvis på tilbakemeldinger fra fagfeller, vurderingen skjer øyeblikkelig.
- Dette er en ørliten revolusjon i forskningsprosessen. I stedet for å sende inn en artikkel og vente på en dom, kan forskere raskt identifisere forbedringsområder og styrke argumentasjonen før innsendelse. Resultatet er en mer dynamisk og effektiv måte å kvalitetssikre forskning på.

Genrer en podkast

- Det mest fascinerende med Notebook LM er kanskje at du kan få den til å lage en egen podkast basert på en tekst. Hvis du foretrekker å lære gjennom lyd i stedet for å lese, kan du be verktøyet generere en dialog mellom to syntetiske personer som diskuterer artikkelen.
- Resultatet blir en overraskende god tilpasset, innsiktsfull podkast, ofte langt mer relevant enn mange av de generiske podkastene som finnes der ute.

Analyse og oversikt

- Mange vil hevde med rette av den virkelige styrken til Notebook LM er at den på en effektiv måte gir oversikt over et helt fagfelt ved å analysere og sammenstille informasjon fra flere artikler samtidig. Du kan laste opp en samling forskningsartikler, og verktøyet vil automatisk generere sammendrag, trekke ut nøkkelbegreper med forklaringer og identifisere hovedtrender i forskningen.
- Dette gjør det enklere å se sammenhenger mellom ulike studier, forstå hvordan feltet har utviklet seg, og oppdage viktige forskningsspørsmål.
- I stedet for å lese hundrevis av sider manuelt, kan du raskt få innsikt i de mest sentrale argumentene, metodene og funnene, noe som sparer tid og gir et mer helhetlig perspektiv på temaet du undersøker.

Det finnes en rekke dedikerte verktøy for å søke etter artikler, avhengigheter og bøker.

Et eksempel er [Keenious](#), en modell som fungerer nesten som en avansert søkemotor for faglitteratur.

 **Khrono**

Nyheter Debatt Student Nyttomfolk +

KI-SKOLEN, DEL 4

Slik kan KI hjelpe deg med språkvask, illustrasjoner og presentasjoner

I fjerde del av KI-skolen viser professor Morten Goodwin hvordan KI-verktøy kan brukes til å finne relevante artikler, til språkvask og til å lage illustrasjoner og presentasjoner.



Illustrasjon fra ChatGPT med følgende instruks: Tegn en trollmann i klassisk fantasy-stil, iført hatt og silkekappe, som representerer en forsker med tilgang til mange KI-verktøy. Rundt ham svever magiske symboler og glødende ikoner som representerer språkvask og illustrasjonsverktøy. Format: 16:9. Illustrasjon: ChatGPT

 **Professor Universitetet i Agder**
Morten Goodwin
morten.goodwin@uia.no

Publisert 16.02.2025 - 18.47



Med instruksen: «*List the recent papers on deep learning for improved education, tested in real-world settings*» kommer Keenious med denne litteraturlisten:

Based on 100 % of your document

1000 results

Article

Open Access

Deep Learning applied to Learning Analytics and Educational Data Mining: A Systematic Literature Review

- 2017 Anais do ... Simpósio Brasileiro de Informática na Educação/Anais do Simpósio Brasileiro de Informática na Educação
- Orlando Bisacchi Coelho et al.
- 40 Citations

Review

Open Access

A Systematic Review of Deep Learning Approaches to Educational Data Mining

Lete etter vitenskapelig innhold

- Denne funksjonaliteten ligner på «beslektede artikler» i Google Scholar og andre litteraturlatabaser, men her benyttes en KI-modell for å gjøre søket enda mer presist og målrettet.
- Alternativer til Keenious inkluderer KI-verktøy som [Elicit](#), som ikke bare søker etter litteratur, men også trekker ut nøkkelinformasjon, og [Connected Papers](#), som går et skritt videre ved å visualisere relaterte artikler i en interaktiv graf du kan navigere i.

Ki-verktøy

Litteratursøk:

- [Keenious](#): KI-verktøy for å finne relevante akademiske artikler.
- [Elicit](#): KI som trekker ut kjerneinformasjon fra forskningsartikler.
- [Connected Papers](#): Visualiserer sammenhenger mellom akademiske publikasjoner.
- [Scite](#): Hjelper forskere med å finne, evaluere og sitere pålitelige artikler.

Språkvask:

- [Grammarly](#): KI-assistent for grammatikk og stilforbedring.
- [Writefull](#): KI-verktøy for akademisk språkvask.
- [DeepL Write](#): Språkvask med fokus på naturlige setningsstrukturer.
- [Quillbot](#): KI for omformulering og tekstforbedring.

KI-verktøy

Illustrasjoner og lysbilder:

- [DALL·E](#): OpenAIs verktøy for bildegenerering.
- [Midjourney](#): Lager realistiske og kunstneriske bilder via Discord.
- [Napkin.ai](#): KI-verktøy for generering av diagrammer og konsepter.

Presentasjoner:

- [Gamma](#): Lager lysbildeserier.

Andre:

- [ChatPDF](#): Gjør PDF-dokumenter interaktive med KI-drevet spørsmålsstøtte.
- [Quizlet AI](#): KI-genererte flashcards for effektiv læring.
- [Scribe](#): Automatiserer opprettelse av guider og dokumentasjon.
- [Eduaide.AI](#): Genererer undervisningsmaterieell og læringsopplegg.

Scite er mer enn bare et verktøy for litteratursøk. De omtaler seg selv som en KI-assistent for forskning. Her kan du stille spørsmål om alt fra å finne relevant litteratur til å designe et eksperiment.

Hvis du for eksempel gir instruksjonen: «*Create a suggested experimental setup for a psychological study to test how colleagues' choice of tools affects your choice*», vil Scite ikke bare foreslå relevante artikler, men også hjelpe til med å strukturere et eksperiment basert på eksisterende forskning.

Scite svarer:

To investigate how colleagues' choice of tools affects individual decision-making in a workplace setting, a well-structured experimental setup is essential. This study can be framed within the context of social influence and decision-making processes, drawing on existing literature that highlights the impact of workplace dynamics on individual choices.

Experimental Design

Participants: The study should involve a diverse group of employees from various departments within an organization. A sample size of at least 100 participants would enhance the statistical power of the findings. Participants should be randomly assigned to different groups to mitigate selection bias.

Procedure: The experiment will consist of two phases: a baseline assessment and an intervention phase.

1. **Baseline Assessment:** Initially, participants will complete a survey assessing their current tool usage and decision-making styles. This survey should include questions about their preferences for specific tools, their perceived effectiveness, and their confidence in using these tools. The survey will also measure individual traits such as

Tekstforbedringer og «live» språkvask

- Språkmodeller kan utføre språkvask ved at du enkelt limer inn teksten og ber dem rydde opp. I tillegg finnes det flere verktøy som er spesialutviklet for nettopp dette, særlig på engelsk.
- Foreløpig finnes det imidlertid ingen KI-baserte språkvaskverktøy som fungerer godt nok på norsk til å kunne måle seg med de engelskspråklige.
- [Grammarly](#) er et eksempel på en KI-basert språkassistent som gir direkte forslag til forbedringer. Den integreres sømløst med Word, Google Docs, Outlook og de fleste andre teksteditorer, slik at feil markeres automatisk, for eksempel med en rød strek i Word.

- 
- Illustrasjoner er nyttige, enten som bakgrunn på et lysbilde eller som blikkfang i en rapport. To av de mest kjente verktøyene for å lage illustrasjoner er DALL·E og [Midjourney](#).
 - DALL·E er tilgjengelig gjennom den betalte versjonen av ChatGPT, og du kan enkelt gi kommandoer som: «Tegn en illustrasjon av ...».
- Illustrasjonen som står helt først i denne artikkelen, og de tilsvarende hovedbildene i de andre KI-skoleartiklene, er laget med ChatGPT/ DALL·E og en lignende instruks.

I ChatGPT kan vi for eksempel komme med følgende instruks:

Tegn en illustrasjon av en forsker som jobber med å utvikle moderne kunstig intelligens i 16:9 format.

Og du har en perfekt bakgrunn for ditt lysbilde:



Hvis resultatet ikke er helt som ønsket, kan du endre illustrasjonen ved hjelp av en teknikk kalt *inpainting*. Dette går ut på at du markerer området du vil justere og gir en ny instruks.



For eksempel tegnet jeg over roboten i bildet og skrev: «Få roboten til å se mer ut som en fabrikkrobot». Modellen justerte deretter illustrasjonen og erstattet den opprinnelige roboten med en mer industriell variant:



Lage presentasjoner

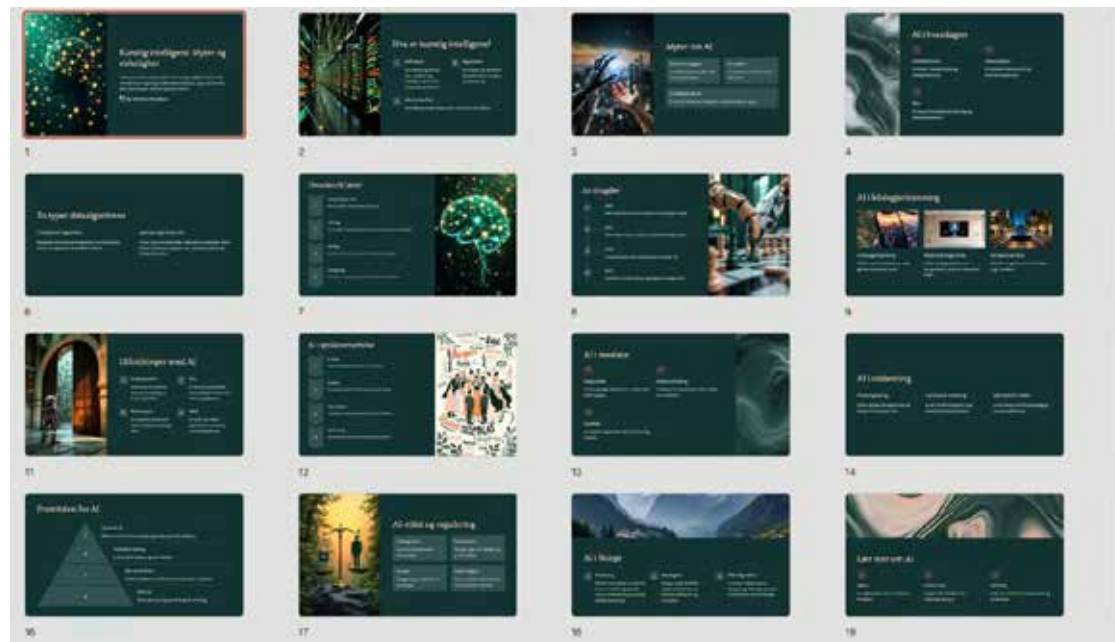
- KI har for lengst bevist sin nytteverdi innen tekst, illustrasjoner og figurer, men den kan også ta et steg videre og lage hele presentasjoner.
- [Gamma](#) er et av de mest effektive verktøyene for å generere nettsider, presentasjoner og innlegg for sosiale medier. Den største fordelen ligger likevel i hvor raskt den kan skape et førsteutkast til lysbilder.
- Du kan enten gi en kort instruks og la verktøyet generere innhold fra bunnen av, eller laste opp eget materiale, for eksempel en PDF-fil, og la KI videreutvikle det.
- Sistnevnte gir ofte mer presise, relevante og gjennomarbeidede resultater. Med noen få klikk kan du ha en grovskisse til en presentasjon, kanskje basert på ett kapittel i læreboken, klar til bruk i en forelesning.

Lage presentasjoner

- Som et eksempel laster jeg opp første kapittel av boken «Myten om maskinene» fra 2020, der jeg selv står som forfatter.
- Jeg instruerer Gamma til å lage en presentasjon på norsk, der innholdet fra kapitlet forkortes — det gir jo liten mening å plassere all teksten på lysbildene. Tekstmengden per lysbilde skal være moderat, og totalt ønsker jeg 20 lysbilder. Stilmessig velger jeg en mørk bakgrunn med grønne nyanser.
- Dette fungerer som en instruks, på samme måte som med ChatGPT, men her brukes KI til å generere lysbilder med tekst, figurer og bilder.
- Resultatet er en strukturert presentasjon der hovedpoengene fra kapitlet er trukket frem.

Lage presentasjoner

- Når presentasjonen er ferdig, kan jeg tilpasse innholdet etter behov. Det mest praktiske for meg er å eksportere den til PowerPoint, slik at jeg kan bruke de verktøyene jeg foretrekker for å gjøre justeringer og finpussing.



Se opp for
hallusinasjoner,
falske referanser
og studenter og
forskere som
jukser.

 **Khrono**

NyheterDebatStudentNytt om folk+

KI-SKOLEN DEL 5

Slik kan du bli lurt av kunstig intelligens

Se opp for hallusinasjoner, falske referanser og studenter og forskere som jukser. I femte del av Khronos KI-skole tar professor Morten Goodwin for seg noen av de store utfordringene med kunstig intelligens — og viser hvordan vi kan håndtere dem.



Illustrasjonen er laget av ChatGPT, som fikk følgende instruks: Lag en sovjetisk propagandainspirert illustrasjon av en etisk kunstig intelligens som veier rett mot galt. Den skal holde en vekt i sine mekaniske hender, symbolsk balanserende rettferdighet. Dramatiske lysstråler skal fremheve dens moralske autoritet.

Illustrasjon: ChatGPT



Professor Universitetet i Agder
Morten Goodwin
morten.goodwin@iuh.no

Publisert 25.02.2025 - 15:57 Sist oppdatert 01.03.2025 - 12:31



Kunstig intelligens: Juks eller hjelpemiddel?

- For å bruke KI til å styrke akademia, må vi bekjempe juks med KI. Akademia har alltid hatt uredelige deltagere, enten det gjelder studenter, forskere eller etablerte akademikere. Juks er ikke nytt.
- Historikeren Herodot, ofte kalt «historiens far,» ble anklaget for å dikte opp hendelser og kilder for å gjøre fortellingene sine mer dramatiske.
- Forskjellen nå, når KI er mer tilgjengelig enn fagbøker, er at juks ikke bare er lettere å gjennomføre, det er også langt vanskeligere å avsløre og enklere å skjule som originalt arbeid.
- **Om fusk er mer utbredt enn før, vet jeg ikke, men én ting er sikkert: Det har aldri vært lettere å kamuflere.**
- Tradisjonelt har akademisk uredelighet vært mulig å avsløre. Hvis en forsker kopierer resultater fra en annen artikkel, kan fagfeller, redaktører eller andre lesere finne den opprinnelige kilden og påpeke plagiering. Slik juks er sporbar, setningene kan sammenlignes, kildegrunnlaget avdekkes, og den ansvarlige kan stilles til ansvar for sitt akademiske brudd.

Forskere som fusker

- Gjør forskere slikt uredelig arbeid?
- Forskere fra Tyskland og USA har avdekket en [eksplosiv økning i KI-genererte ord](#) siden ChatGPTs lansering. Ved å analysere PubMed, en sentral søkemotor for biomedisinsk forskning, fant de at ordet «delves» (utforsker) i 2024 ble brukt 28 ganger oftere i abstrakter enn før ChatGPTs inntog. «Showcasing» (vise fram) og «underscores» (understreker) forekom over ti ganger så hyppig som tidligere.
- Ord som ChatGPT ofte bruker, havner altså oftere i akademiske artikler. Konklusjonen må være at forskerne bak artiklene i PubMed får hjelp av ChatGPT til å skrive abstraktene. Det ingen grunn til å tro at forskning som dukker opp i PubMed skiller seg fra forskning i andre disipliner.

Når algoritmene dikter opp sannheten

- Kunstig intelligens kan formulere seg overbevisende, men det betyr ikke at den alltid har rett. Et av de største problemene i språkmodeller er fenomenet kjent som KI-hallusinasjon, når en språkmodell genererer feilinformasjon som høres plausibel ut. Dette er spesielt bekymringsfullt i akademia, hvor presisjon og etterprøvbarehet er avgjørende.
- Hva skjer egentlig når språkmodeller hallusinerer?
- En språkmodell som ChatGPT fungerer ved å forutsi det mest sannsynlige neste token i en setning. Et token kan være et enkelt tegn, en bokstavkombinasjon eller et helt ord, avhengig av konteksten.
- Modellen har ingen innebygd forståelse av sannhet eller fakta. Den analyserer enorme mengder tekst og forsøker å generere et mest mulig sannsynlig svar basert på mønstrene den har lært.

Når algoritmene dikter opp sannheten

- Men når modellen mangler tilstrekkelige data, eller når den prøver å svare på noe den ikke har dekning for, fyller den inn hullene med noe som kunne vært sant. Det er dette som kalles hallusinasjon. Språkmodellen gjetter, og den gjør det med en selvsikkerhet som gjør det vanskelig å skille mellom fakta og fiksjon.
- Begrepet hallusinasjon kan virke misvisende, siden det ikke fungerer på samme måte som menneskelige hallusinasjoner. Mennesker ser eller hører ting som ikke er der, mens en KI skaper feilaktig informasjon ut av statistiske mønstre. Likevel er resultatet det samme: en virkelighetsforvrengning som fremstår troverdig.

Når algoritmene dikter opp sannheten

Hvordan oppdager vi hallusinasjoner?

- Kildekritikk: Kontroller alltid om referanser faktisk eksisterer.
- Flerkildeverifisering: Sammenlign KI-generert informasjon med pålitelige kilder.
- Nyere modeller: I nyere språkmodeller modeller er problemet ytterligere redusert. Pålitelige tall mangler, men erfaring tilsier at hallusinasjoner skjer langt sjeldnere enn før.
- Bruk av RAG-modeller: Modeller som kombinerer KI med faktiske databaser (for eksempel [Perplexity](#)) kan redusere risikoen for feil.

Fordommer

- Når en språkmodell mates med stereotyp eller misvisende informasjon, blir disse skjevhetene en naturlig del av svarene den gir. Majoritetens perspektiv dominerer, mens nyanser og unntak forsvinner. Dermed blir sykepleiere kvinner, professorer menn, sjefer middelaldrende hvite menn, datanerder unge gutter med briller og barnehagelærere smilende kvinner.
- Modellen gjenskaper ikke nødvendigvis verden slik den er eller burde være, men slik den oftest er beskrevet i treningsdataene sine.
- I akademia kan dette bety at forskning som baserer seg på språkmodeller, ikke bare gjenspeiler eksisterende fordommer, men også viderefører dem.

Fordommer

Konsekvenser for academia:

- **Feilslutninger:** Når treningsdata er ufullstendige eller skjevfordelte, vil KI trekke konklusjoner som preger disse virkelighetene.
- **Skjevheter i datagenerering:** Hvis vi bruker KI til datagenerering, kan skjevheter forsterke eksisterende skjevheter, spesielt i samfunnsfag, medisin og naturvitenskap.
- **Redusert mangfold i forskningen:** Når KI primært bygger på tidligere data, kan det bidra til en selvforsterkende sirkel hvor etablerte perspektiver dominerer.
- **Spredning av stereotypier:** KI kan videreføre kulturelle og sosiale fordommer i pensum og undervisningsmaterieell.

KI-SKOLEN DEL 6

Hvordan gjøre kunstig intelligens til forsknings-assistent

I siste del av Khronos KI-skole viser professor Morten Goodwin blant annet hvordan kunstig intelligens (KI) kan hjelpe til med å utvikle hypoteser, lage forskningsdesign og brukerundersøkelser, og analysere datasett.



Illustrasjon med ChatGPT med instruksjon: Visualiser en forsker med ett bein i fortiden og ett i fremtiden. Bildet skal symbolisere overgangen fra gammel til ny forskning. Noen deler av bildet skal ha et gammeldags preg, mens andre skal være moderne og futuristiske. Format: 16:9. Illustrasjon: ChatGPT



Professor Universitetet i Agder
Morten Goodwin
morten.goodwin@uia.no

Teknologien kan effektivisere rutineoppgaver, frigjøre tid til kjerneaktiviteter og samtidig heve kvaliteten på det som produseres.

En stor del av arbeidsdagen i akademia går med til oppgaver som ikke er direkte knyttet til forskning og undervisning. Møter, e-poster, administrativt arbeid og byråkratiske prosesser stjeler tid og etterlater mindre rom for fordypning og faglig nyskaping.

KI kan avlaste oss på mange av disse områdene, men i denne delen retter vi blikket mot det som virkelig betyr noe — akademias kjerneoppgaver.

Effektivitet versus kvalitet

- Språkmodeller kan hjelpe oss med å skrive flere forskningsartikler og forskningsrådssøknader, fagfelleverdure flere publikasjoner og undervise i flere fag, rett og slett fordi de kan overta deler av arbeidet.
- Men dette er ikke nødvendigvis en utvikling vi ønsker.
- Ingen ønsker at akademisk arbeid skal reduseres til en form for masseproduksjon, en slags «fast food science». I stedet ønsker vi å bruke KI til å heve kvaliteten på forskningen, styrke undervisningen, øke den analytiske dybden og sikre at arbeidet gir grundigere og mer presise resultater.

La KI foreslå hypoteser

- Hypoteser er avgjørende fordi de danner grunnlaget for systematisk utforskning og vitenskapelig forståelse. Uten klare hypoteser famler vi forskere i blinde, mens gode hypoteser gir retning og struktur til undersøkelsene.
- Her kan kunstig intelligens være til stor hjelp. En språkmodell kan foreslå hypoteser basert på eksisterende kunnskap og mønstre i data. På denne måten kan KI bidra til å avdekke nye sammenhenger og inspirere forskere til å stille enda skarpere spørsmål.

Eksempelvis kan vi spørre ChatGPT direkte om å komme med forslag til hypoteser med følgende instruks:

Generer tre mulige hypoteser om hvordan bruk av digitale læringsplattformer påvirker studentengasjement i høyere utdanning.

ChatGPT svarer:

Her er tre mulige hypoteser om hvordan bruk av digitale læringsplattformer påvirker studentengasjement i høyere utdanning:

Økt interaktivitet øker engasjementet: Bruk av digitale læringsplattformer med interaktive funksjoner (f.eks. diskusjonsfora, quiz, og gamification) øker studentenes engasjement ved å gjøre læringsprosessen mer dynamisk og deltakende.

Fleksibilitet forbedrer selvregulert læring: Digitale læringsplattformer gir studentene større fleksibilitet i læringsprosessen, noe som kan føre til økt engasjement ved å muliggjøre individuell tilpasning og selvregulert læring.

Manglende sosial interaksjon kan redusere engasjementet: Selv om digitale læringsplattformer kan forbedre tilgjengelighet og fleksibilitet, kan de også redusere studentengasjement dersom de erstatter ansikt-til-ansikt interaksjoner og fører til mindre sosial tilknytning mellom studenter og undervisere.

Disse hypotesene kan testes gjennom kvantitative og kvalitative metoder, for eksempel spørreundersøkelser, eksperimenter eller intervjuer med studenter og undervisere.

Dette kan være et nyttig førsteutkast hvis du skal starte et nytt arbeid om hvordan digitale læringsplattformer påvirker studenter.

Hypoteser bygget på forskningslitteratur

- Som forskere har vi inngående kjennskap til de mest innflytelsesrike arbeidene innen vårt felt og en klar forståelse av hvilke spørsmål som bør utforskes.
- Ved å veilede språkmodellen med nøye utvalgte kilder og relevant kontekst, kan vi ikke bare forankre hypoteseforslagene i vår egen ekspertise, men vi har også muligheten til å utvide perspektivet ved å utforske mulighetene som ligger i den nyeste forskningen.

Vi kan gi ChatGPT følgende instruks:

Her laster jeg opp tre forskningsartikler som omhandler digitale læringsplattformer og studentengasjement:

«Synthesis of student engagement with digital technologies: a systematic review of the literature»

«Engagement in Online Learning: Student Attitudes and Behavior During COVID-19»

og

«Impact of online learning on student's performance and engagement: a systematic review»

På bakgrunn av funnene i disse studiene, formuler nye hypoteser som ikke allerede er besvart i artiklene, men som kan bidra til å utvikle forskningsfeltet videre.

ChatGPT svarer:

Her er noen potensielle hypoteser basert på funnene i de tre forskningsartiklene, men som tar sikte på å utvide forskningsfeltet:

1. Individuell motivasjon og teknologistøtte

Hypotese: «Personlig tilpassede læringsanbefalinger ved hjelp av kunstig intelligens vil øke studentens kognitive og atferdsmessige engasjement i digitale læringsplattformer.»

Begrunnelse: Flere studier har påpekt utfordringer med studentengasjement i digitale læringsmiljøer. Ved å bruke AI-baserte adaptive læringsverktøy kan man undersøke om økt personalisering kan redusere passiv deltakelse og øke motivasjon.

2. Sosial tilknytning og peer-engasjement

Hypotese: «Bruk av gamification-elementer som belønningssystemer og konkurranser vil øke student-til-student-interaksjon i digitale læringsmiljøer.»

Begrunnelse: Studier viser at sosial interaksjon er en kritisk faktor for engasjement, men at mange studenter føler seg isolerte i digitale læringsmiljøer. Dette kan testes ved å undersøke hvordan konkurransepregede elementer påvirker samarbeidsvilje og faglig interaksjon.

3. Effekten av blandet læring på langtidseffekter

Hypotese: «Studenter som kombinerer asynkrone og synkrone læringsaktiviteter vil vise høyere læringsutbytte og langsiktig kunnskapsretensjon sammenlignet med de som kun deltar i én læringsmodus.»

Begrunnelse: Flere studier viser at synkron og asynkron læring har ulike fordeler. En longitudinell studie kan undersøke om en kombinasjon av de to gir bedre læringsutbytte over tid.

4. Teknologibarrierer og engasjement

Hypotese: «Studenter med begrenset tilgang til høyhastighetsinternett vil ha lavere affektivt og kognitivt engasjement i digitale læringsplattformer, uavhengig av andre

Her bruker ChatGPT faktisk forskning som grunnlag, der du som forsker selv velger hvilke artikler som skal danne utgangspunkt for hypotesene.

I stedet for at modellen baserer seg på generelle kilder, har du full kontroll over prosessen og sikrer at den bygger videre på relevant og pålitelig forskning.

Dette gir langt større presisjon og transparens i hypoteseforslagene, ettersom språkmodellen ikke bare genererer ideer, men også begrunner dem ut fra de valgte kildene.

Datadrevne hypoteser

Vi kan ta dette et steg videre ved å la språkmodellen foreslå hypoteser direkte fra et datasett, enten et du har samlet selv eller et som andre forskere har satt sammen.

- På denne måten kan KI bidra til å avdekke mønstre og sammenhenger som ellers kunne gått ubemerket hen.

La oss se på et eksempel:

Her er data fra en faktisk brukerundersøkelse fra «Student Engagement and Empowerment» samlet i Washington 2014–2019.

På bakgrunn av funnene i disse studiene, formuler nye hypoteser basert på disse undersøkelsene, som kan bidra til å utvikle forskningsfeltet videre for nye brukerundersøkelser.

Dataene er lastet ned fra **Student Engagement and Empowerment (SEE) Project, Washington, 2014–2019 (ICPSR 37659)**.

Datasettet inneholder brukerundersøkelser med svar fra en bred gruppe studenter, både i form av avkrysningsspørsmål og utdypende beskrivelser av deres opplevelser av hvor nyttige og engasjerende ulike deler av studiet var.

ChatGPT svarer:

Basert på funnene fra «Student Engagement and Empowerment (SEE) Project, Washington, 2014-2019», her er noen nye hypoteser som kan bidra til å utvikle forskningsfeltet videre:

1. Effekten av personlig relevans på læringsengasjement

Hypotese: «Studenter som får læreplaninnhold tilpasset deres personlige interesser og bakgrunn vil vise høyere grad av kognitiv og atferdsmessig engasjement enn de som følger en standardisert læreplan.»

Begrunnelse: Flere studenter i undersøkelsen uttrykte at videoer og PowerPoint-presentasjoner var verdifulle når de følte relevante, men mange følte at de var for generelle.

2. Langtidseffekten av veksttankegang i studentmotivasjon

Hypotese: «Studenter som introduseres for veksttankegang (growth mindset) tidlig i skoleløpet vil vise en større økning i akademisk utholdenhet og selvregulering enn de som først eksponeres for konseptet senere.»

Begrunnelse: Dataene viser at noen studenter responderte positivt på veksttankegangs-videoer, men effekten varierte avhengig av når i skoleløpet de fikk tilgang til denne kunnskapen.

3. Læring gjennom deling og refleksjon

Hypotese: «Å skrive refleksjonsbrev til yngre studenter om læringsprosessen sin vil føre til dypere kognitiv bearbeiding og økt akademisk selvtillit sammenlignet med å kun motta informasjon passivt.»

Begrunnelse: En betydelig andel studenter uttrykte at brevskrivningen hjalp dem til å internalisere læring, men noen fant den mindre nyttig fordi de skrev til andre og ikke til seg selv.

4. Forbindelsen mellom sosial tilhørighet og akademisk prestasjon

Hypotese: «Studenter som deltar i strukturerte gruppediskusjoner om sosial tilhørighet og skolemiljø vil vise høyere grad av akademisk engasjement enn de som

Å la ChatGPT analysere dataene dine og følge instruksjonene dine gir et solid grunnlag for videre utforskning. Det kan brukes til mye, blant annet å definere hypoteser basert på faktiske datakilder, avdekke mønstre som ellers kunne gått ubemerket hen, eller foreslå nye retninger for eksperimenter og analyser.

Forskningsdesign med KI

- Dette kan videreutvikles ved å bruke språkmodeller til å utforme et forskningsdesign.
- Gjennom målrettede spørsmål kan modellen foreslå metodikk, identifisere relevante variabler og anbefale egnede analyseverktøy, noe som styrker både struktur og gjennomføring av studien.

Instruks:

Du skal nå bygge videre på den første hypotesen «Effekten av personlig relevans på læringsengasjement». Sett opp et klart og presist forskningsdesign for studien, med både kvantitative og kvalitative metoder. Inkluder en beskrivelse av forskningsspørsmål, metodevalg, utvalg, datainnsamling, analysemetoder og etiske hensyn.

Spesifiser hvordan brukerundersøkelser skal gjennomføres på fem videregående skoler og hvordan dybdeintervjuer skal struktureres.

Forskningsdesign: Effekten av personlig relevans på læringsengasjement

1. Introduksjon

Formålet med denne studien er å undersøke hvordan personlig relevant læringsinnhold påvirker studenters kognitive, atferdsmessige og affektive engasjement i videregående skole. Studien benytter en mixed-methods-tilnærming, som kombinerer kvantitative spørreundersøkelser og kvalitative dybdeintervjuer.

2. Forskningsspørsmål

- Hvordan påvirker personlig tilpasset læringsinnhold studenters motivasjon og deltakelse i undervisningen?
- Hvilke forskjeller i læringsengasjement finnes mellom elever som får personlig tilpasset innhold og de som får standardisert innhold?
- Hvordan opplever elever personlig relevans i ulike fag, og hvordan påvirker dette deres engasjement?
- Hvilke faktorer påvirker hvorvidt personlig tilpasset undervisning gir økt engasjement?

3. Metodevalg

Studien gjennomføres som et kvasi-eksperiment på fem videregående skoler, hvor en gruppe elever får personlig tilpasset læringsinnhold og en kontrollgruppe får standardisert undervisning. Kvantitative data samles inn via spørreundersøkelser, mens kvalitative data samles inn gjennom dybdeintervjuer.

Metodetriangulering:

- Kvantitativ metode: Spørreundersøkelser før og etter intervensjonen
- Kvalitativ metode: Dybdeintervjuer med utvalgte elever og lærere

4. Utvalg

- Fem videregående skoler deltar i studien
- Totalt utvalg: ~500 elever fordelt på ulike fagretninger
- Eksperimentgruppe (n = 250): Elever får undervisning med personlig tilpasset innhold
- Kontrollgruppe (n = 250): Elever får ordinært undervisningsopplegg
- Lærere (n = 10): Intervjues for å forstå deres perspektiv på læringsengasjement

Inklusjonskriterier:

- Elever på Vg1 og Vg2 fra ulike studieprogram
- Villighet til å delta i både spørreundersøkelser og mulige intervjuer

5. Datainnsamling

5.1 Kvantitativ metode: Spørreundersøkelser

- Gjennomføres digitalt på skolene
- To måletidspunkter:
 - Pre-test (uke 0, før intervensjonen)
 - Post-test (uke 8, etter intervensjonen)

Variabler og måleinstrumenter:

Variabel	Måleverktøy (Likert-skala 1-5)
Kognitivt engasjement	Interesse for fagstoff, refleksjon over læring
Atferdsmessig engasjement	Deltakelse i diskusjoner, innsatsnivå
Affektivt engasjement	Følelsesmessig tilknytning til faget
Opplevd relevans	Hvor nært læringsinnholdet oppleves til egne interesser og fremtidig karriere

Statistiske analyser:

- Sammenligning av pre- og post-test i eksperiment- og kontrollgruppen (t-test, ANOVA)
- Korrelasjonsanalyse mellom opplevd relevans og engasjement

5.2 Kvalitativ metode: Dybdeintervjuer

- Semistrukturerte intervjuer gjennomføres med et utvalg av 20 elever (10 fra eksperimentgruppen, 10 fra kontrollgruppen)
- Fem lærere intervjues for å få innsikt i hvordan undervisningsopplegget påvirker studentenes engasjement

Intervjuguide:

- For elever:
 - Hvordan opplever du relevansen av læringsinnholdet i ulike fag?
 - På hvilke måter påvirker personlig tilpasset undervisning din motivasjon?
 - Hvordan sammenligner du læringsopplevelsen mellom ulike undervisningsformer?
 - Er det spesielle undervisningsmetoder du finner mer engasjerende enn andre?
- For lærere:
 - Hvordan tilpasser du undervisningen for å gjøre innholdet relevant for elevene?
 - Hvordan opplever du at personlig relevans påvirker elevenes engasjement?
 - Hvilke utfordringer møter du med tilpasset undervisning?

Analyse av intervjuene:

- Tematisk analyse for å identifisere mønstre og forskjeller mellom elevgruppene
- Sammenligning av erfaringer mellom elever som fikk tilpasset innhold og de som fikk standardisert innhold

6. Etiske hensyn

- Informert samtykke fra alle deltakere
- Foreldresamtykke for elever under 18 år
- Anonymisering av data for å beskytte deltakernes personvern
- Godkjenning fra relevante instanser (f.eks. Norak senter for forskningsdata)

7. Forventede resultater og betydning

- Forventet funn: Elever i eksperimentgruppen vil rapportere høyere engasjement enn de i kontrollgruppen
- Studien kan gi praktiske implikasjoner for undervisningsutvikling, ved å vise hvordan tilpasset læring kan øke motivasjon og prestasjoner
- Resultatene kan informere utviklingen av digitale læringsplattformer og pedagogiske strategier

Dette forskningsdesignet sikrer en robust tilnærming for å undersøke hvordan personlig relevans påvirker studentengasjement. Er det noen aspekter du ønsker å tilpasse eller utdype?

Modellen gir metodeforslag, utvalgsstrategier og relevante analysemetoder. Selv om dette er en nyttig start, er det åpenbart at det alene ikke er tilstrekkelig for å sette i gang studien.

Forskningsdesign med KI

- Igjen, målet er ikke å la ChatGPT gjøre forskningen for oss, men å bruke den som et verktøy for idéutvikling og strukturering.
- Forskningsprosessen krever kritisk refleksjon og metodisk presisjon, noe forskeren selv må sikre. KI kan fungere som en sparringpartner og bidra til å styrke forskningsdesignet, men det er opp til forskeren å forme og gjennomføre studien.
- Vi kan derfor fortsette dialogen med ChatGPT, få den til å justere metodene og komme med alternative forslag.

Lage brukerundersøkelser

Vi kan fortsette dialogen med ChatGPT innenfor det samme forskningssporet, men denne gangen med fokus på utforming av brukerundersøkelsen.

- Språkmodellen kan bistå med å formulere relevante spørsmål ved hjelp av følgende instruks:

Kom med eksempler på spørsmål vi bør stille i brukerundersøkelsen. Pass på at den måler flere aspekter av engasjement: studentenes egne refleksjoner, faktisk deltagelse og følelsesmessig tilknytning.

ChatGPT er, som alltid, hjelpsom:

Her er et forslag til spørsmål for brukerundersøkelsen, organisert etter tema og tilpasset en Likert-skala (1–5) der det er relevant. Disse spørsmålene vil bidra til å måle studentenes kognitive, atferdsmessige og affektive engasjement, samt deres opplevelse av personlig relevans i undervisningen.

1. Bakgrunnsspørsmål

(Disse hjelper med å kategorisere data og identifisere forskjeller mellom grupper.)

- Hvilket studieprogram går du på? (Studieforberedende / Yrkesfaglig / Annet)
- Hvilket klassetrinn er du på? (Vg1 / Vg2 / Vg3)
- Hvor interessert er du generelt i skolearbeid? (1 = Ikke interessert, 5 = Svært interessert)

2. Opplevd personlig relevans

(Måler i hvilken grad læringsinnholdet oppleves som relevant for den enkelte student.)

- Undervisningen jeg mottar er relevant for mine personlige interesser. (1 = Helt uenig, 5 = Helt enig)
- Jeg ser sammenhenger mellom det vi lærer og mine fremtidsplaner. (1–5)
- Læringsaktivitetene vi gjør i klassen er tilpasset mine behov og preferanser. (1–5)
- Jeg opplever at lærerne gjør en innsats for å knytte undervisningen til virkelige situasjoner jeg kan relatere meg til. (1–5)

3. Kognitivt engasjement

(Måler hvor aktivt studentene reflekterer over læringen.)

- Jeg tenker ofte over hvordan det jeg lærer kan brukes i praksis. (1–5)
- Når jeg lærer noe nytt, prøver jeg å finne ut hvordan det henger sammen med tidligere kunnskap. (1–5)
- Jeg stiller ofte spørsmål til læreren for å forstå faget bedre. (1–5)
- Jeg tar initiativ til å finne ekstra informasjon om temaer vi lærer om. (1–5)

4. Atferdsmessig engasjement

(Måler deltakelse i læringsaktiviteter.)

- Jeg deltar aktivt i diskusjoner og gruppearbeid i timene. (1–5)
- Jeg fullfører de fleste av leksene og oppgavene som blir gitt. (1–5)
- Jeg føler meg motivert til å delta i aktiviteter som krever problemløsning eller kreativ tenkning. (1–5)
- Jeg prøver å holde meg engasjert, selv om undervisningen ikke er direkte interessant for meg. (1–5)

5. Affektivt engasjement

(Måler følelsesmessig tilknytning til læring og skole.)

- Jeg gleder meg til timene der vi jobber med temaer jeg synes er relevante. (1–5)
- Jeg føler meg engasjert og motivert når undervisningen kobles til mine interesser. (1–5)
- Jeg opplever at læring er meningsfullt for meg. (1–5)
- Jeg føler meg stolt når jeg mestrer et tema vi har jobbet med. (1–5)

6. Spørsmål for sammenligning av eksperiment- og kontrollgruppen

(Disse brukes til å sammenligne elever som har fått personlig tilpasset undervisning med de som har fått standardundervisning.)

- Jeg opplever at undervisningen i dette faget er mer engasjerende enn i andre fag. (1–5)
- Jeg opplever at læringsmetodene i dette faget hjelper meg å forstå stoffet bedre. (1–5)
- Jeg ville foretrukket at flere fag brukte samme undervisningsmetoder som dette. (1–5)

7. Åpne spørsmål

(Gir mulighet for utdypende svar og kvalitativ innsikt.)

- Hva slags undervisningsmetoder gjør deg mest engasjert i læringen?
- Kan du gi eksempler på når undervisningen har følt seg særlig relevant for deg?
- Hva kan læreren gjøre for å gjøre undervisningen mer relevant for dine interesser og fremtidsplaner?
- Er det noe i denne studien du ønsker å gi tilbakemelding om?

Tilpasning for før- og etterundersøkelse:

- Før-test: Spør om studentens generelle opplevelse av undervisning og læringsengasjement før intervensjonen.
- Etter-test: Samme spørsmål, men med fokus på hvordan opplevelsen har endret seg etter intervensjonen.

Datainnsamling og analyse

- Distribusjon: Spørreundersøkelsen sendes digitalt til elever på fem videregående skoler via skolens læringsplattform eller e-post.
- Analyse:
 - Deskriptiv statistikk for å identifisere generelle trender.
 - T-test eller ANOVA for å sammenligne eksperiment- og kontrollgruppens resultater.
 - Korrelasjonsanalyse mellom opplevd relevans og engasjement.
 - Kvalitativ tematisk analyse av åpne spørsmål for å identifisere mønstre i svarene.

Disse spørsmålene sikrer en grundig evaluering av hvordan personlig relevans i undervisningen påvirker læringsengasjement. Vil du ha tilpasninger eller flere forslag?

Like verdifullt som å be ChatGPT om direkte forslag, er det å bruke modellen som en sparringpartner etter at du selv har utformet et utkast til brukerundersøkelsen.

Lage brukerundersøkelser

- Du kan laste opp spørsmålene dine og få modellen til å identifisere svakheter og forbedringsområder, foreslå alternative formuleringer og justeringer, knytte spørsmålene til relevant forskningslitteratur, vurdere om forskningsdesignet dekker alle nødvendige aspekter, påpeke eventuelle skjevheter eller manglende perspektiver og mer.
- En slik prosess kan bidra til å styrke undersøkelsens kvalitet og sikre at den er både metodisk solid og forskningsmessig relevant.
- På samme måte kan språkmodellen brukes i andre deler av forskningsprosessen, ikke bare som en sparringpartner for hypoteser og undersøkelser, men den kan foreslå analysemetoder, strukturere intervjuguider eller gi innspill til hvordan forskningsfunn kan formidles på en klar og engasjerende måte.

KI-hjelp til dataanalyse

- En viktig del av både kvalitativ og kvantitativ forskning er dataanalyse. Kvantitative forskere, i særdeleshet, bruker mye tid på å analysere, visualisere og tolke data gjennom statistiske metoder for å avdekke mønstre og sammenhenger. Dette gjelder alt fra eksperimentelle analyser av algoritmer, som jeg forsker på, til enhver annen form for dataanalyse.
- Her kan også språkmodeller være til stor hjelp. De kan bistå med å analysere store datamengder, identifisere mønstre, oppsummere sentrale funn og til og med foreslå alternative analysemetoder. På denne måten kan KI ikke bare effektivisere arbeidsflyten, men også bidra til dypere innsikt i forskningsdataene.

KI-hjelp til dataanalyse

- I eksemplene som følger, analyseres vin gjennom kjemiske data samlet fra én italiensk region og tre ulike druesorter. Selve vinen er ikke hovedfokuset — det er analyseteknikkene og hvordan en språkmodell kan støtte forskningen.
- Ulike fagfelt benytter forskjellige metoder, og her demonstreres en regresjonsanalyse kombinert med visualisering og en enkel maskinlæringsmodell. Maskinlæringsmodellen forutsier hvilken vingård en vin stammer fra basert på dens kjemiske sammensetning.
- Språkmodeller kan naturligvis håndtere mange ulike statistiske metoder, men målet her er ikke å dekke alle muligheter — det er å vise en konkret analyse i praksis. Forskere kan bruke språkmodeller til å teste nærmest hvilken som helst dataanalyseteknikk, fra enkle statistiske metoder til avanserte maskinlæringsmodeller. Her presenteres ett eksempel på hvordan slike verktøy kan brukes i praksis.

KI-hjelp til dataanalyse

- Vi tar utgangspunkt i en vinstudie presentert i «Aeberhard, S., Coomans, D., & De Vel, O. (1994). [*Comparative analysis of statistical pattern recognition methods in high dimensional settings. Pattern Recognition*](#), 27(8), 1065—1077.
- Dataene kan lastes ned fra [UC Irvine Machine Learning Repository](#).

Først etablerer vi konteksten for ChatGPT ved å introdusere vinstudien og gi en presis instruks:

Her er et dataset over vin.

Attributtene er:

1. Alcohol
2. Malic acid
3. Ash
4. Alcalinity of ash
5. Magnesium
6. Total phenols
7. Flavanoids
8. Nonflavanoid phenols
9. Proanthocyanins
10. Color intensity
11. Hue
12. OD280/OD315 of diluted wines
13. Proline

Alle spørsmål jeg stiller fra nå av, skal være for å analysere dette. Forstår du?

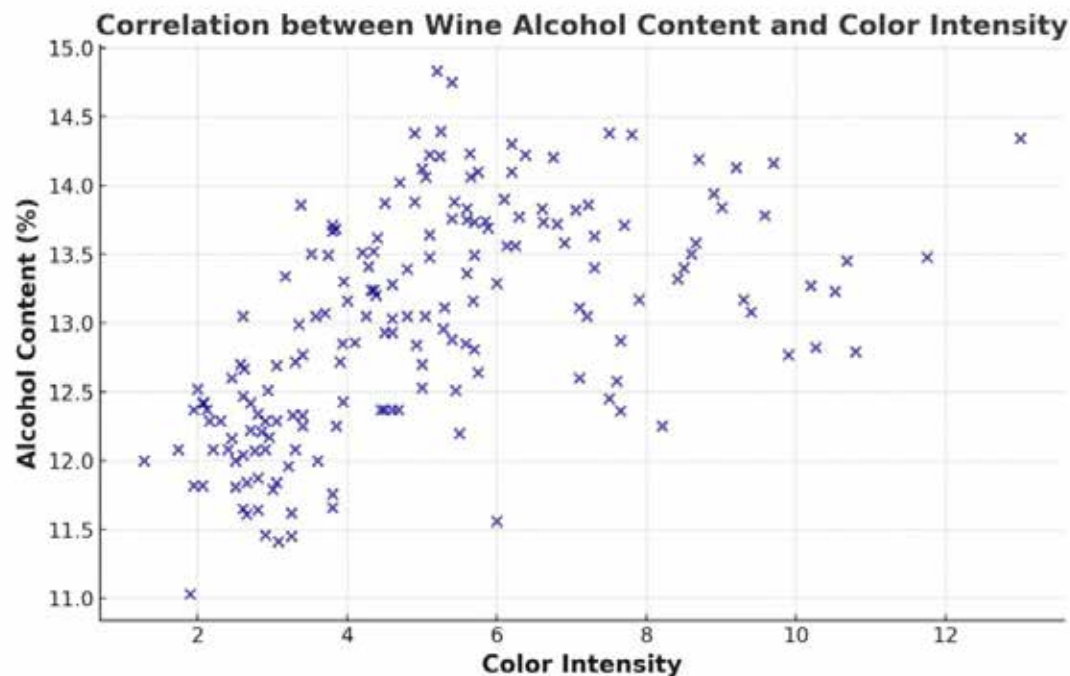
ChatGPT2 svarer:

Ja, jeg forstår! Jeg kan analysere datasetet ditt på forskjellige måter. Ønsker du en generell oversikt først, for eksempel en beskrivelse av dataene, statistiske oppsummeringer eller visualiseringer? Eller har du spesifikke spørsmål du vil undersøke?

La oss si at vi har en forskningshypotese om at det er en sammenheng mellom farge og alkoholprosent. Du kan vi for eksempel først se på dataene, og be språkmodellen lage en graf med følgende instruks:

Lag en graf som sammenligner Alcohol med Color intensity. På x-aksen skal det være Color intensity på y-aksen skal det være Alcohol. Gjør det så vitenskapelig som mulig. Det skal være en del av en forskningsartikkel.

ChatGPT koder da et lite python-script for å gjøre analysen, kjører sin egen kode, og presenterer resultatene i form av visualiseringer, som i dette eksempel er et dataplott:



Det er viktig å forstå at det ikke er ChatGPT som utfører selve analysen. Den skriver koden for oss, og det er denne koden som gjennomfører beregningene.

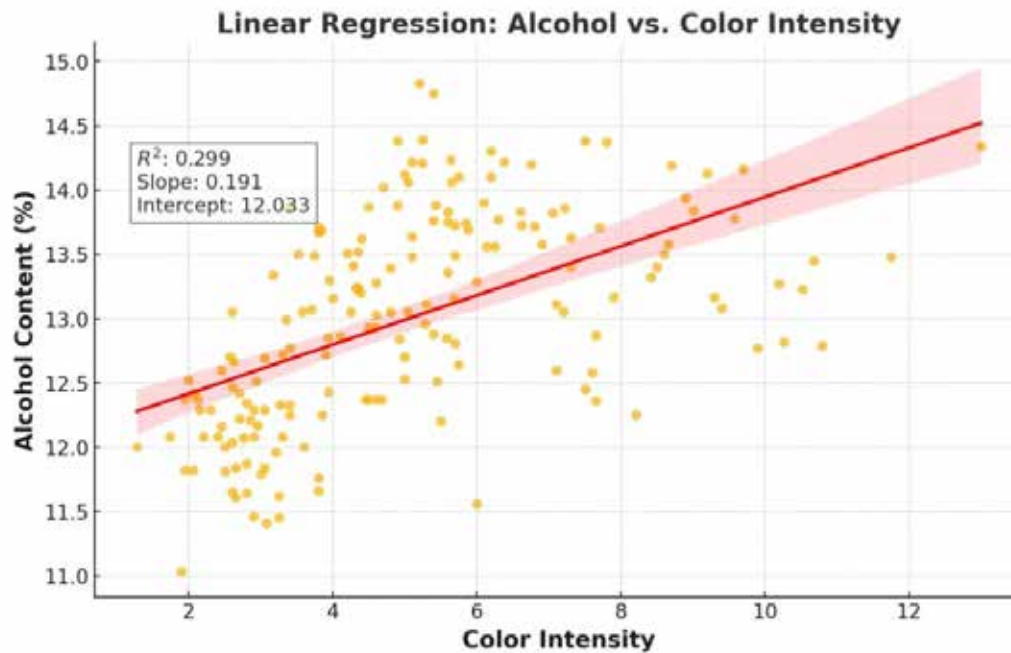
Denne åpne Python-koden, som like gjerne ChatGPT kunne skrevet i Matlab, R eller Excel, kan granskes i detalj av oss akademikere, kvalitetssikres og justeres ved behov. Eventuelle feil eller «hallusinasjoner» oppstår i selve koden, ikke i analysen.

- 
- Så lenge vi kvalitetssikrer koden før vi bruker den, kan vi derfor ha høy tillit til at analysen er korrekt. Faren for hallusinasjoner er minimal fordi det er den faktiske Python-koden som utfører analysen, ikke ChatGPTs språkmodell.
 - Vi kan ta analysen et steg videre ved å bruke regresjonsanalyse, en metode som ofte benyttes for å avdekke mønstre, kartlegge sammenhenger mellom variabler og forutsi fremtidige verdier basert på historiske data.
 - I dette tilfellet kan vi undersøke om det finnes en lineær sammenheng mellom farge og alkoholprosent. Med andre ord: Har mørkere viner en tendens til å inneholde mer alkohol?

Instruks:

Kjør en lineær regresjonsanalyse og plott den lineære regresjonen i grafen.

ChatGPT fortsetter å kode, og grafen koden spytter ut, er:



Samtidig svarer ChatGPT:

Her er resultatene fra den lineære regresjonsanalysen mellom Alcohol Content (%) og Color Intensity, samt det oppdaterte plottet med regresjonslinjen.

Regresjonsresultater:

- Regresjonslikning: $\text{Alcohol} = 12.033 + 0.191 \times \text{Color Intensity}$
- R^2 -verdi: 0.299 → Omtrent 29.9 % av variasjonen i alkoholinnhold kan forklares av fargeintensitet.
- p-verdi for Color Intensity: < 0.0001 → Sterkt signifikant sammenheng.
- F-statistikk: 74.90 → Indikerer en sterk sammenheng mellom variablene.
- Durbin-Watson-verdi: 1.009 → Liten sjanse for seriell autokorrelasjon i restene.

Tolkning:

- Det er en statistisk signifikant sammenheng mellom Color Intensity og Alcohol Content, men sammenhengen er moderat ($R^2 = 0.299$).
- En økning på 1 enhet i Color Intensity er assosiert med en økning på 0.191 % i Alcohol Content.
- Likevel, siden R^2 ikke er veldig høy, er det andre faktorer som også påvirker alkoholinnholdet i vin.

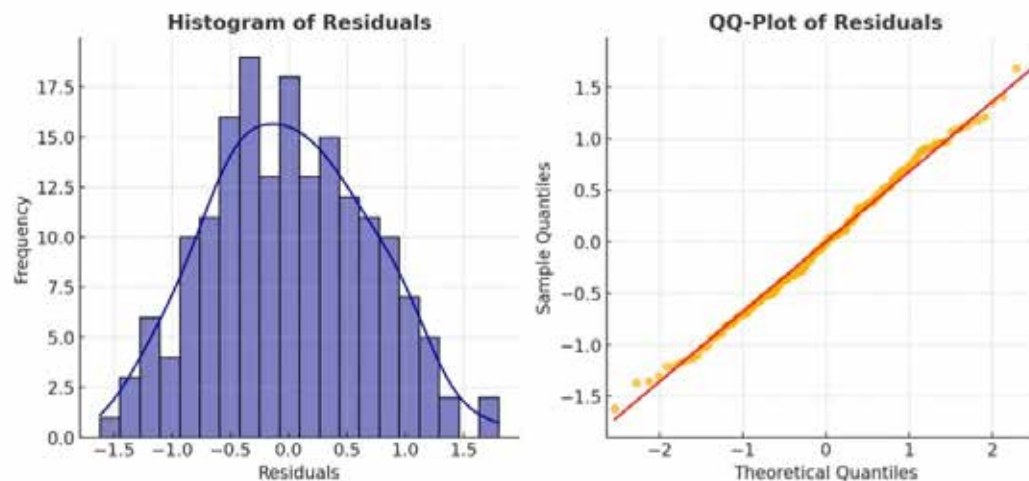
[ChatGPT oppgir en lang Python-kode, som ikke gjengis her]

- 
- Ifølge ChatGPTs analyse finnes det en viss sammenheng mellom vinens farge og alkoholprosent, men den er ikke særlig sterk. For å sjekke hvor pålitelig denne sammenhengen er, kan vi se nærmere på residualene — altså forskjellen mellom den faktiske alkoholprosenten i vinene og den prosentandelen modellen vår forutsier.
 - Enklere sagt: Vi sammenligner de faktiske målingene med den røde linjen ChatGPT nettopp tegnet, og ser hvor mye de avviker.
 - Hvis disse avvikene fordeler seg jevnt rundt null, betyr det at modellen treffer ganske godt. Men hvis vi ser et tydelig mønster, for eksempel at modellen konsekvent bommer i én retning, kan det være et tegn på at vi mangler viktige faktorer eller at sammenhengen ikke er så enkel som en rett linje.
 - Kort sagt, vi sjekker om modellen holder mål, eller om vi må finjustere den for å få et mer realistisk bilde.

Vi kan derfor spørre:

Kan du sjekke om det er «normalfordeling av residualer»

Språkmodellen plottes:



Og svarer:

Her er analysen av residualene for å sjekke om de er normalfordelte:

Visualiseringer:

1. Histogram av residualer: Viser en omtrent klokkeformet fordeling, noe som indikerer en tilnærmet normalfordeling.
2. QQ-Plot: Punktene ligger nær den diagonale linjen, noe som støtter at residualene følger en normalfordeling.

Shapiro-Wilk test for normalitet:

- Teststatistikk: 0.993
- p-verdi: 0.610 (større enn 0.05)

Tolkning:

- p-verdi > 0.05 betyr at vi ikke kan forkaste nullhypotesen om at residualene er normalfordelte.
- Både den visuelle inspeksjonen og Shapiro-Wilk testen tyder på at residualene er tilnærmet normalfordelte, noe som er en god indikasjon på at den lineære regresjonsmodellen er passende.

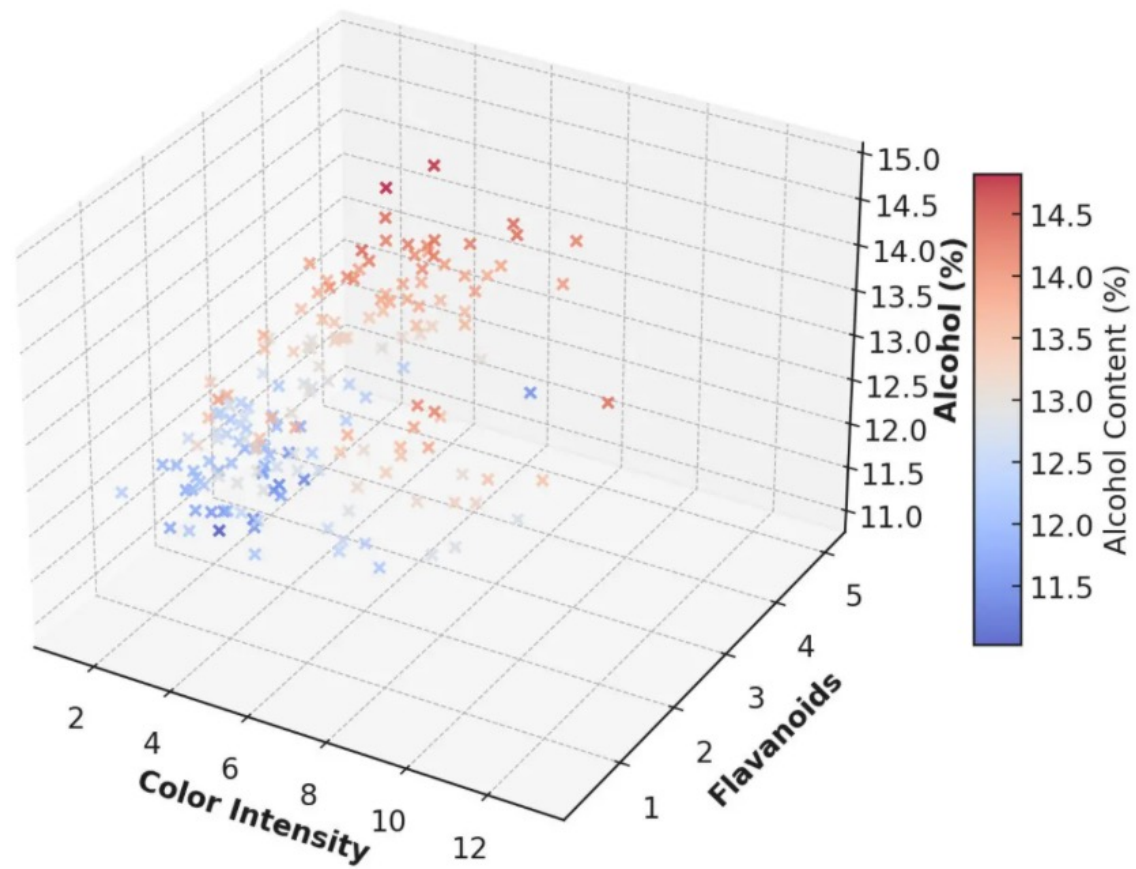
[ChatGPT oppgir en lang Python-kode, som ikke gjengis her]

- 
- Vi må selvfølgelig kvalitetssikre resultatene, men fordelene er at vi kan åpne Python-koden og, så lenge vi forstår den, undersøke i detalj hvordan modellen fungerer.
 - I tillegg kan vi følge med på plott og visualiseringer for å bekrefte at modellen faktisk leverer meningsfulle resultater.
 - Og vi kan gå langt utover enkle 2D-plott, for eksempel med følgende instruks:

Kom med et 3D-Scatterplot som utforsker Alcohol, Color Intensity og Flavanoids.

ChatGPT svarer med denne figuren:

3D Scatterplot of Alcohol, Color Intensity, and Flavanoids



Potensielle etiske
problemer?

Den utbredte bruken av Chat Generative Pre-trained Transformer (kjent som ChatGPT) og annen ny teknologi som drives av generativ kunstig intelligens (GenAI) har trukket oppmerksomhet til de potensielle etiske problemene de kan forårsake, spesielt i applikasjoner med høy innsats som helsevesen, men etiske diskusjoner har ennå ikke blitt oversatt til operasjonelle løsninger.

Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist

Yilan Ning*, Seilvelut Teraszyewski, Yueping Shang, Julian Sawicki, Vishvaarth Nagaraj, Di Mao, Mayel Mertens, David Shu Wei Tang, Jasmine Chai Ling Ong, Mingxuan Liu, Jikwen Cao, Michael Duran, Roger Vaughan, Marcus Ong Hock Ong, Joseph Joo Yiu Sung, Eric J Tapil, Nan Lai

The widespread use of Chat Generative Pre-trained Transformer (known as ChatGPT) and other emerging technology that is powered by generative artificial intelligence (GenAI) has drawn attention to the potential ethical issues they can cause, especially in high-stakes applications such as health care, but ethical discussions have not yet been translated into operationalisable solutions. Furthermore, ongoing ethical discussions often neglect other types of GenAI that have been used to synthesise data (eg, images) for research and practical purposes, which resolve some ethical issues and expose others. We did a scoping review of the ethical discussions on GenAI in health care to comprehensively analyse gaps in the research. To reduce the gaps, we have developed a checklist for comprehensive assessment and evaluation of ethical discussions in GenAI research. The checklist can be integrated into peer review and publication systems to enhance GenAI research and might be useful for ethics-related disclosures for GenAI-powered products and health-care applications of such products and beyond.

Introduction

In the past few years, development of Chat Generative Pre-trained Transformer (ChatGPT) and other chatbots powered by large language models (LLMs) have drawn the attention of the general public, researchers, and stakeholders to the fast developing technology of GenAI. GenAI differentiates from general AI technology due to its capacity to generate realistic content, such as text (eg, ChatGPT and Gemini), images (eg, Midjourney and DALL-E), and videos (eg, Sora).¹ The promising capability of GenAI models are also being explored for their medical applications, such as for mental health,² breast cancer,³ dietary care,⁴ and medical education.⁵

However, medical applications of GenAI raise additional ethical concerns due to the ongoing discussions on general AI. For example, the large data volume and user community of LLMs raise the stakes of privacy breaches, as exemplified by a brief ban of ChatGPT in Italy to address and clarify privacy concerns.⁶ GenAI-generated data might leak personal information if sufficient details were captured that almost reproduce real training samples.⁷ As GenAI is trained to generate realistic outputs based on patterns learnt from data and can create hallucinations (ie, incorrect and misleading outputs), ensuring factual accuracy for medical applications requires great effort.⁸ Stochasticity of GenAI output, vulnerability to prompt injections, and opacity of input data complicate the assessment of GenAI for clinical fitness, which is predominantly based on predictive performance in the current AI context.⁹ These factors lead to active discussions on, and urgent calls for, new guidelines, regulations, and legislations on GenAI,¹⁰⁻¹² especially for health care and medical education.¹³⁻¹⁵

Emerging GenAI-powered technology has been quickly rolled out to a huge and diverse user community, which explains why there are enthusiastic discussions on its ethical implications on social media and within the research community. However, other GenAI

methods, such as the generative adversarial network (GAN),¹⁶ have not been considered in current ethical discussions, even though these methods have been used in medical research and are subject to similar ethical considerations. Notably, GAN (and related GenAI methods)¹⁷ are often used as solutions to address some ethical concerns in medical research (eg, to protect privacy by masking images and videos,^{18,19} or by creating synthetic data).²⁰⁻²² A closer investigation of existing discussions on GenAI that cause and mitigate ethical concerns would provide a more holistic view on this important topic, which could ultimately help identify actionable points for GenAI research. In this scoping review, we aimed to understand the responses of the research community to potential ethical issues of GenAI in health care. Informed by these findings, we developed an ethics checklist to operationalise existing ethical guidelines, which can be used by journals, institutional review boards, funders, and regulators to promote responsible GenAI research in health care and beyond.

Methods

Our systematic scoping review on ethical discussions associated with GenAI in health care followed the PRISMA extension for scoping reviews guidelines.²³

Search strategy and selection criteria

We searched PubMed, Embase (Excerpta Medica Database), Ovid, Web of Science, and Scopus for articles written in English with a set of search terms associated with three main concepts: 'AI ethics', 'generative AI', and 'healthcare' that were published between Jan 1, 2013, and July 25, 2023. Detailed definitions of the three concepts, search terms, and search strategy are described in the appendix (p 1).

We excluded articles that were not within the domain of health care, did not apply GenAI, or did not discuss ethical issues in relation to the use of GenAI in



Current Digital Health 2024; 4:e448-55

Published Online

September 12, 2024

https://doi.org/10.1016/j.dighe.2024.09.002

© 2024 The Author(s)

*Contributed equally

Center for Quantitative

Medicine (Y Ning PhD),

Teknisk og Kvalitet, Y Ning PhD,

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Stavanger, Norway (J Tapil PhD),

Generative artificial intelligence and ethical considerations in health care

- Videre neglisjerer pågående etiske diskusjoner ofte andre typer GenAI som har blitt brukt til å syntetisere data (f.eks. bilder) for forskning og praktiske formål, som løser noen etiske problemer og avslører andre.
- Vi gjorde en omfangsgjennomgang av de etiske diskusjonene om GenAI i helsevesenet for å analysere hull i forskningen grundig.
- For å redusere hullene har vi utviklet en sjekkliste for omfattende vurdering og evaluering av etiske diskusjoner i GenAI-forskning.
- Sjekklisten kan integreres i fagfelle vurderings- og publiseringssystemer for å forbedre GenAI-forskning og kan være nyttig for etikkrelaterte avsløringer for GenAI-drevne produkter og helseapplikasjoner av slike produkter og utover.

Panel 1: Definitions of the nine ethical principles for generative artificial intelligence (GenAI) in a health-care context

Accountability

- The explicit clarifications of to whom and to what extent responsibility or legal liability, or both fall
- The mandated and moral duty to establish regulatory mechanisms to prevent potential adverse effects on patients from the use of GenAI

Autonomy

- The preservation and fostering of patients' dignity, rights for self-determination, and capacity to make informed decisions
- Provision of understandable information to enable patients to use according to their values

Equity

- The use of GenAI to promote equity according to some notion of fairness (equality of opportunity, outcomes, etc) in health or health resources across diverse groups of patient populations, and to actively prevent or remedy systemic, unfavourable outcomes in specific patient populations
- The equitable access to AI or GenAI technology

Integrity (in medical education and quality of clinical research)

- The commitment to intellectual honesty and personal responsibility to abide by responsible research conduct, including data integrity, to establish accountability and prevent harm
- The rightful acknowledgment of contributions to and ownership of intellectual work when GenAI is used in clinical research

Non-maleficence

- The prevention of harm and potential risks to patients associated with GenAI use in health care, such as incorrect or misleading outputs (hallucinations)

Privacy

- The protection of patients' information from illegitimate access, and of confidentiality of personal sensitive information

Security

- The protection of health data integrity and safety, through careful assessments of vulnerabilities in data systems and the prevention of data breaches, cyberattacks, or other threats

Transparency

- The full disclosure and thorough documentation of information regarding GenAI development, including its dataset and evaluation of performance
- The ability to access and understand the processes underlying models' outputs, especially pertaining to black-box models, in so far as this is possible

Trust

- The confidence of users in GenAI and its developers, and expectations that the model is competent in performing its prespecified tasks and behaves in ways that serves patients and the medical community
- Evidence of performance and its limitations
- The willingness to accept and integrate GenAI tools to assist delivery of care or research
- Trustworthy GenAI processes and exhibition of a range of ethically reliable properties, including performance robustness, fairness, and security, among others

Panel 2: Ethics terms associated with the nine ethical principles of generative artificial intelligence in a health-care context

Accountability

Governance, responsibility, accountability, legal, and liability

Autonomy

Respect, human autonomy, human oversight, informed decisions, informed consent, and valid consent

Equity

Fair, fairness, bias, disparity, discrimination, justice, equity, equality, and inequity

Integrity (in medical education and quality of clinical research)

Integrity, plagiarism, data integrity, copyright, ownership, and intellectual property

Non-maleficence

Safe, safety, harm, harmful, misinformation, truthful, risk, and benevolence

Privacy

Privacy, private, confidential, and confidentiality

Security

Security and cybersecurity

Transparency

Transparent, transparency, explainable, explainability, explicability, interpretable, interpretability, non-interpretability, black box, opacity, and white box

Trust

Trust, trustworthy, and trustworthiness



Spørsmål?

Gunnar.hartvigsen@helgelandssykehuset.no