

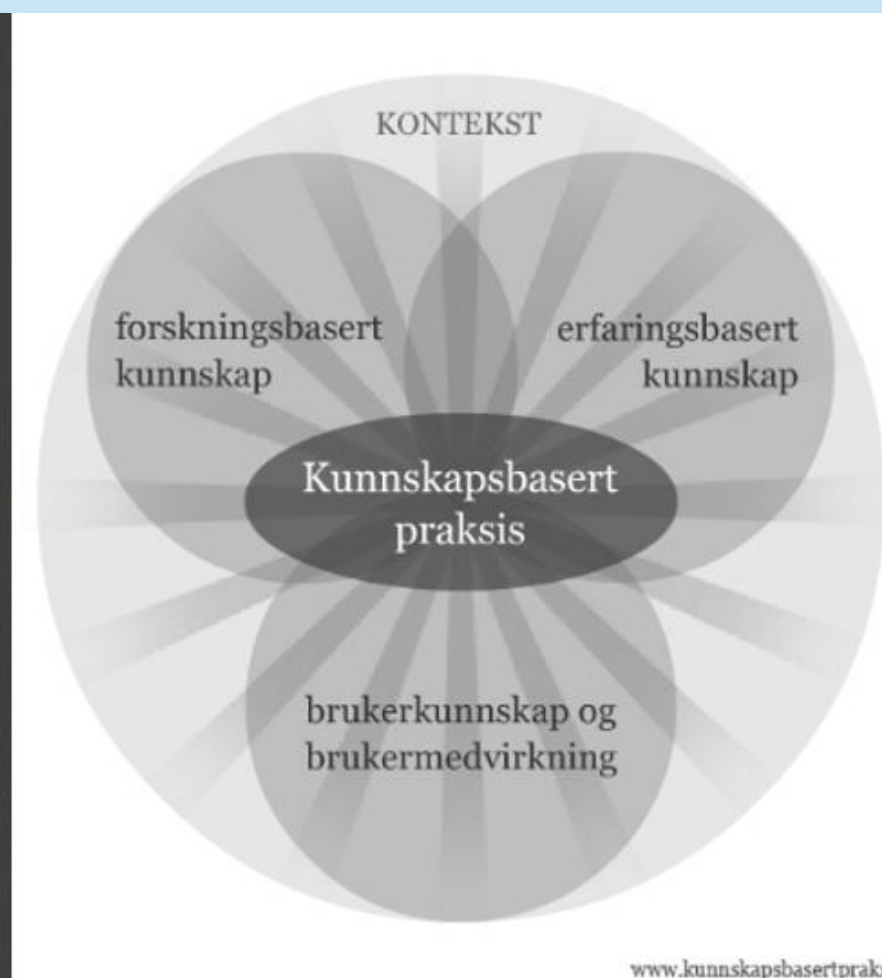
BRUK AV KUNSTIG INTELLIGENS I KUNNSKAPSBASERT PRAKSIS

Muligheter og utfordringer ved bruk av kunstig intelligens som støtte for helsepersonell i kunnskapsbasert praksis

20.10.2025

Nils Gunnar Landsverk

- Kunnskapsbasert praksis- definisjon og prosess
- Kunstig intelligens (KI) – definisjon og eksempler
- Bruk av kunstig intelligens i kunnskapsbasert praksis
 - KI i systematiske oversikter og retningslinjer
 - KI som hjelp til å finne og forstå forskning
 - KI som en digital sparringspartner/ assistent
- Tillitsverdige KI i helsetjenesten, forskning og utdanning



[www.kunnskapsbasertpraksis](http://www.kunnskapsbasertpraksis.no)

Kunnskapsbasert praksis (KBP) er å ta faglige avgjørelser basert på systematisk innhentet forskningsbasert kunnskap, erfaringsbasert kunnskap og pasientens ønsker og behov i en gitt situasjon (3)



(Kunnskapsbasertpraksis.no)

KBP-prosessen

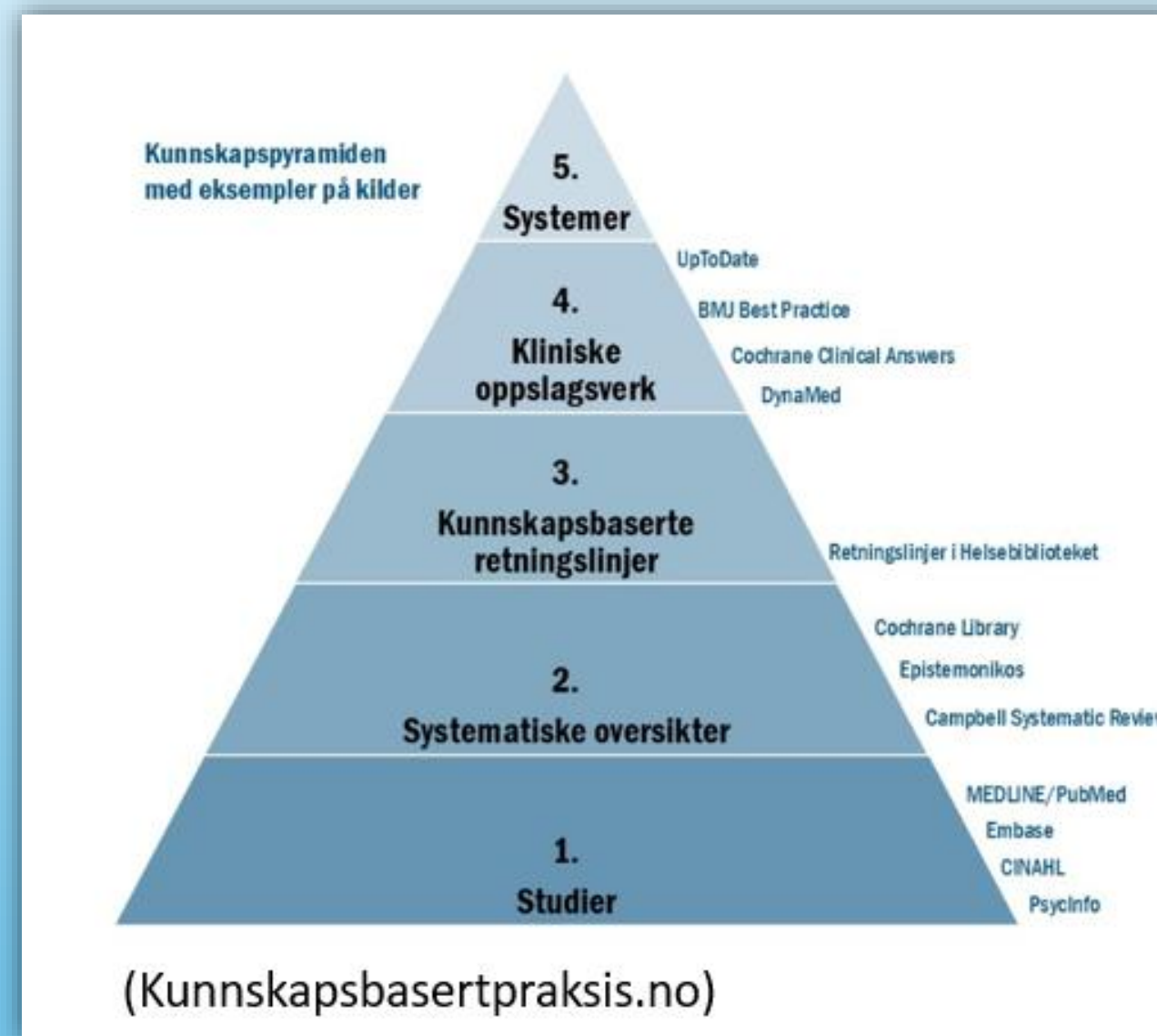
- 1) Reflektere over egen praksis
- 2) Formulere presise forskningsspørsmål

P	Population/problem	Hvilken populasjon eller hvilket problem dreier det seg om?
I	(Phenomenon of) Interest	Hvilken aktivitet, erfaring, opplevelse eller prosess dreier det seg om?
Co	Context	Hvilken kontekst eller setting dreier det seg om?

(Kunnskapsbasertpraksis.no)

KBP-prosessen

- 1) Reflektere over egen praksis
- 2) Formulere presise forskningsspørsmål
- 3) Finne forskningsbasert kunnskap



KBP-prosessen

- 1) Reflektere over egen praksis
- 2) Formulere presise forskningsspørsmål
- 3) Finne forskningsbasert kunnskap
- 4) Kritisk vurdere kunnskapen

Sjekkliste for vurdering av en oversiktsartikkel

Hvordan brukes sjekklisen?

Sjekklisen består av tre deler:

- A: Kan du stole på resultatene?
- B: Hva forteller resultatene?
- C: Kan resultatene være til hjelp i praksis?

I hver del finner du underspørsmål og tips som hjelper deg å svare. For hv skal du krysse av for «ja», «nei» eller «uklart». Valget «uklart» kan også oi



CASP Checklist:
For Randomised Controlled Trials (RCTs)

GRADE Handbook

Introduction to GRADE Handbook

Handbook for grading the quality of evidence and the strength of reco approach. Updated October 2013.

Editors: Holger Schünemann (schuneh@mcmaster.ca), Jan Brozek (brozek@mcmaster.ca), Guyatt (guyatt@mcmaster.ca), and Andrew Oxman (oxman@online.n

APPRAISAL OF GUIDELINES FOR RESEARCH & EVALUATION II

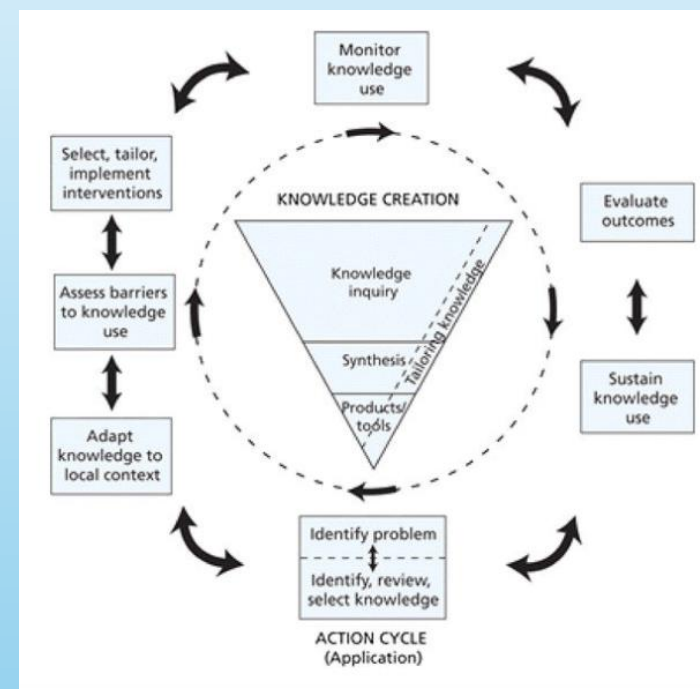


INSTRUMENT

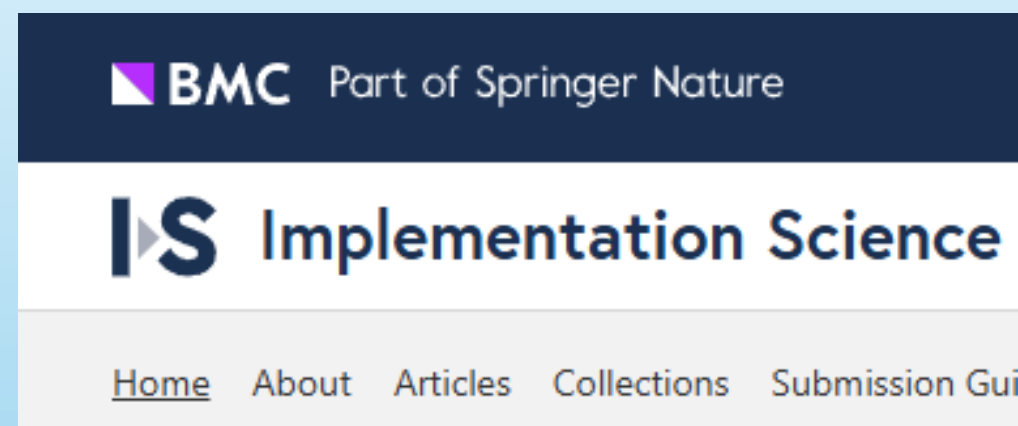
KBP-prosessen

- 1) Reflektere over egen praksis
- 2) Formulere presise forskningsspørsmål
- 3) Finne forskningsbasert kunnskap
- 4) Kritisk vurdere kunnskapen
- 5) Anvende den forskningsbaserte kunnskapen sammen med egen erfaring og pasientens behov

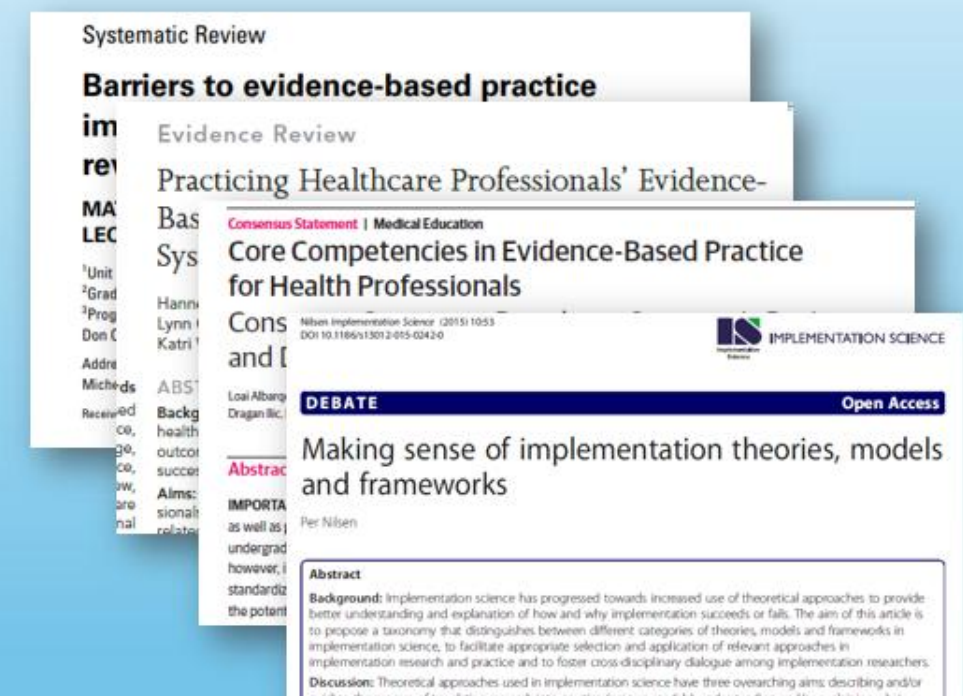
Implementeringsforskning!



Graham et al., 2006



Eccles et al., 2006



KBP-prosessen

- 1) Reflektere over egen praksis
- 2) Formulere presise forskningsspørsmål
- 3) Finne forskningsbasert kunnskap
- 4) Kritisk vurdere kunnskapen
- 5) Anvende den forskningsbaserte kunnskapen
sammen med egen erfaring og pasientens behov
- 6) Evaluere egen praksis

Å praktisere KBP-prosessen krever..

Tilstrekkelig KBP-kompetanse

- Kunnskap om KBP-konseptet og viktige begreper
- Grunnleggende KBP-ferdigheter
- Mestringstro tilknyttet egne KBP-evner
- Troa på viktigheten og nytten av KBP

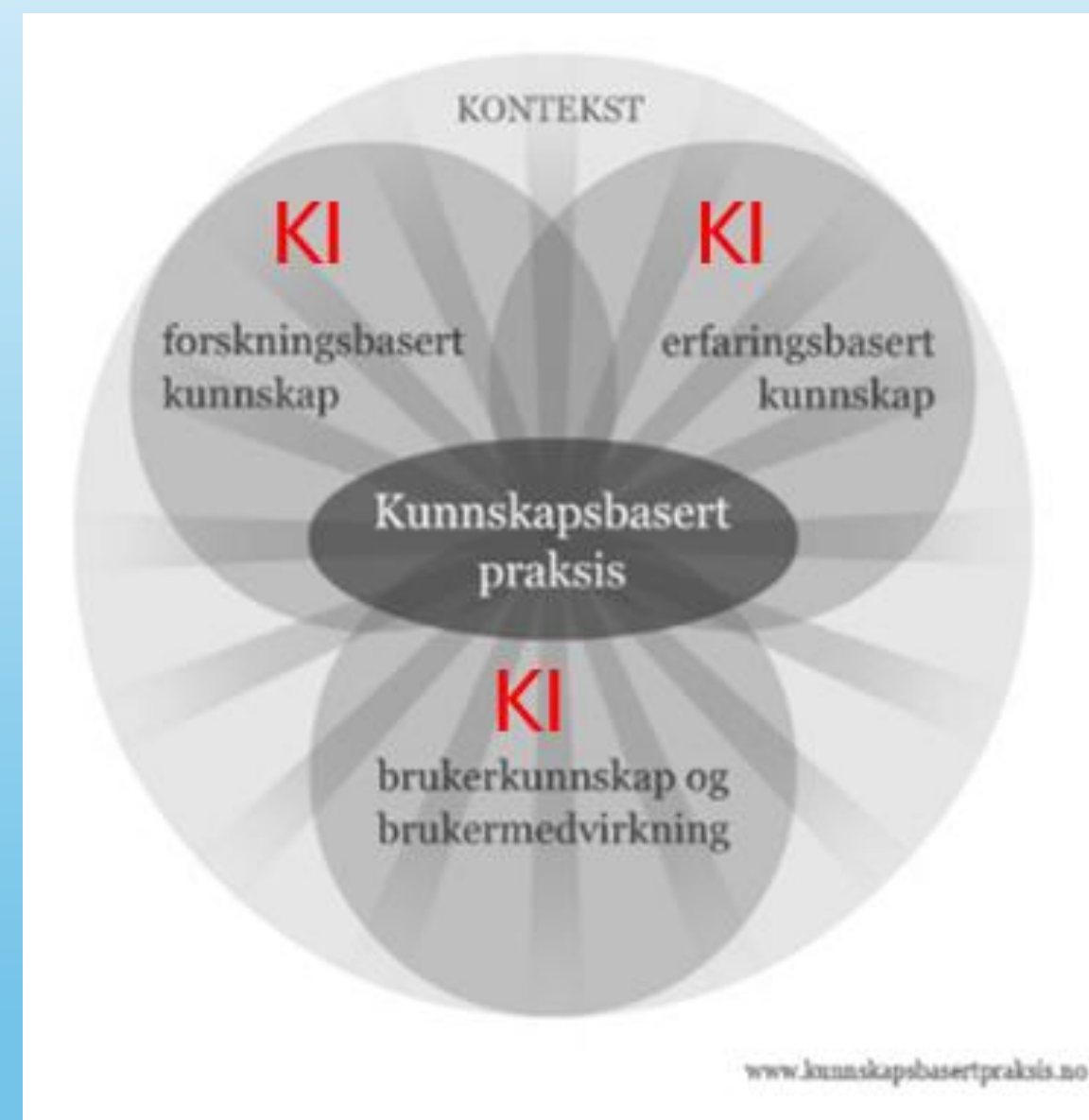
Men blir ofte hindret av...

- Mangel på tid og ressurser
- Mangel på lederstøtte
- Begrenset tilgang
- Arbeidsmiljø, tradisjon, vaner



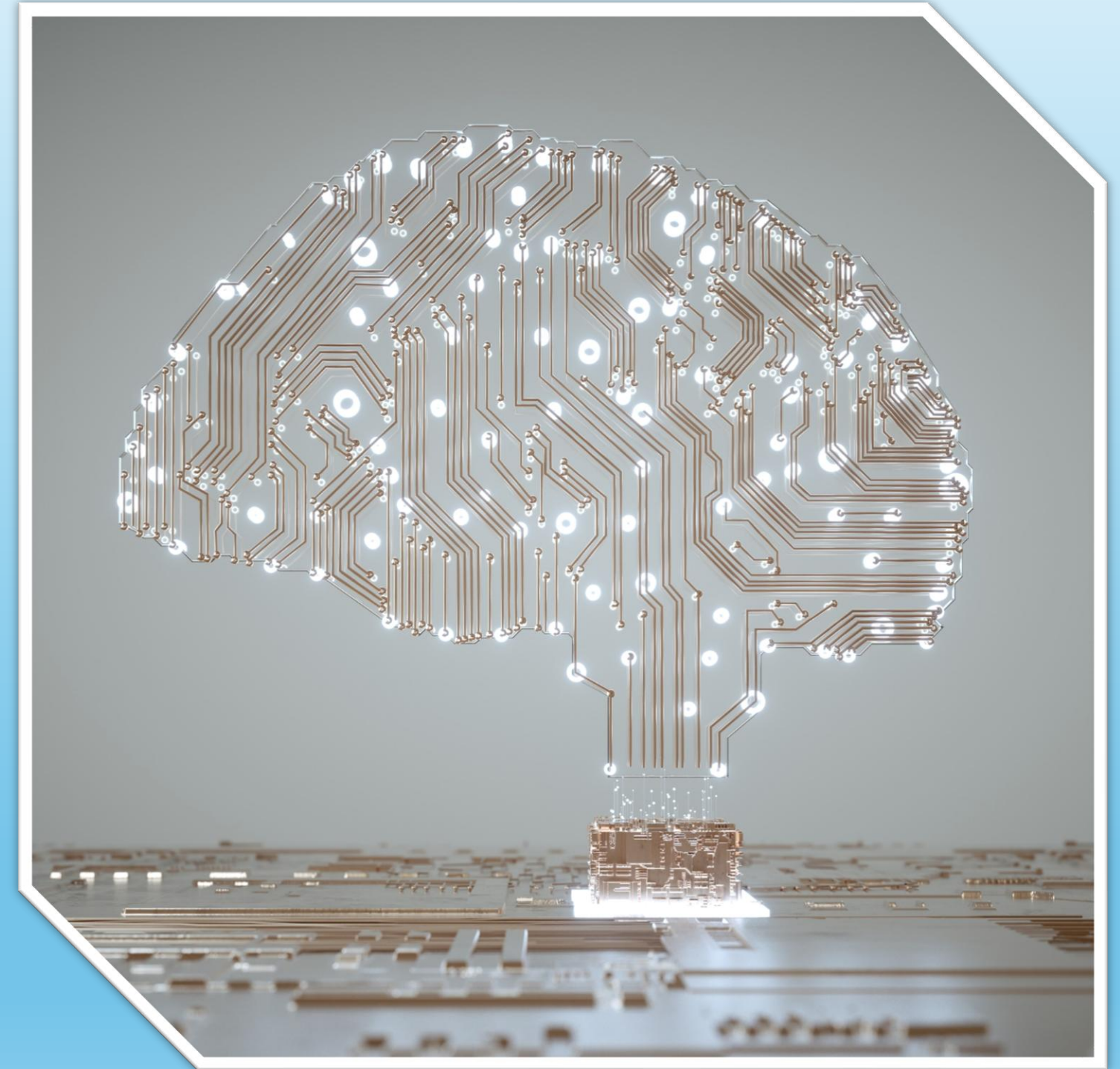
KI + KBP = sant?

«Kunstig intelligens (KI) kan bidra til raskere og mer presis diagnostikk, bedre beslutningsstøtte til personell, forenklet logistikk, automatisering av administrative oppgaver, og å sette innbyggerne bedre i stand til å følge opp sin egen helse. KI vil kunne utgjøre et betydelig bidrag til en bærekraftig utvikling av vår felles helsetjeneste» (1)

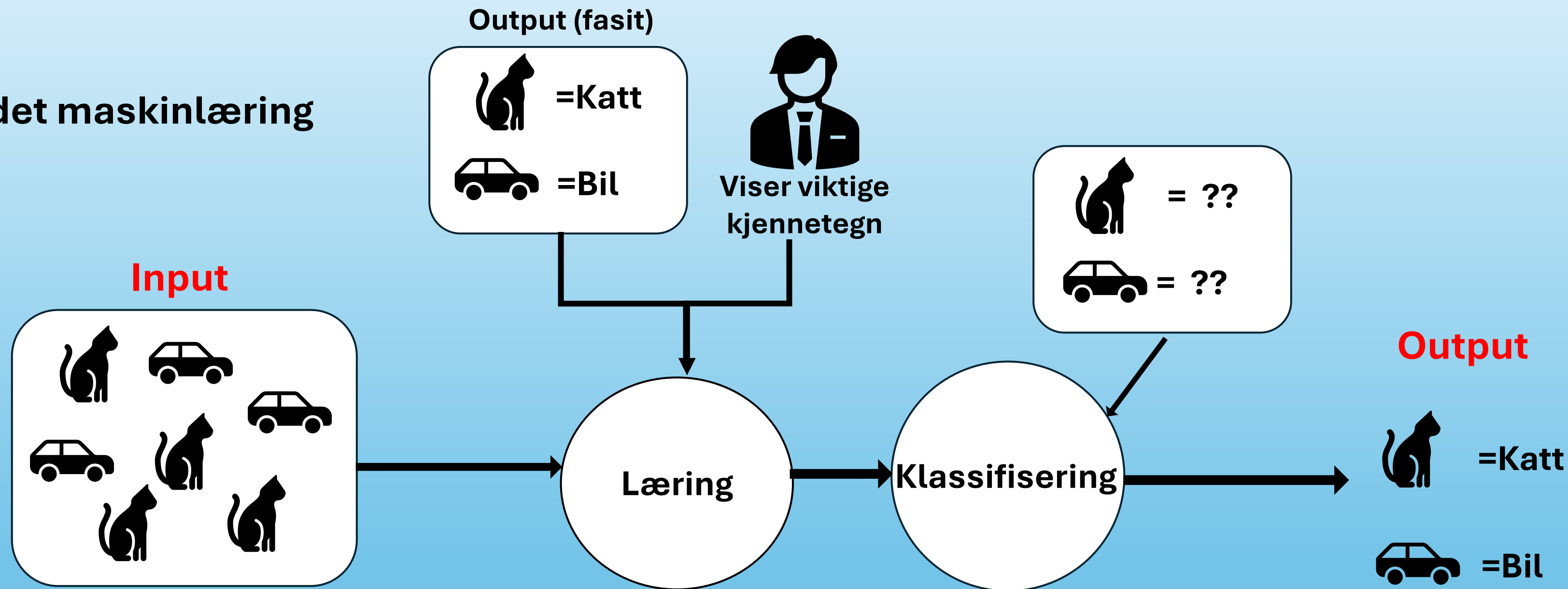


(Kunnskapsbasertpraksis.no), redigert for denne forelesningen

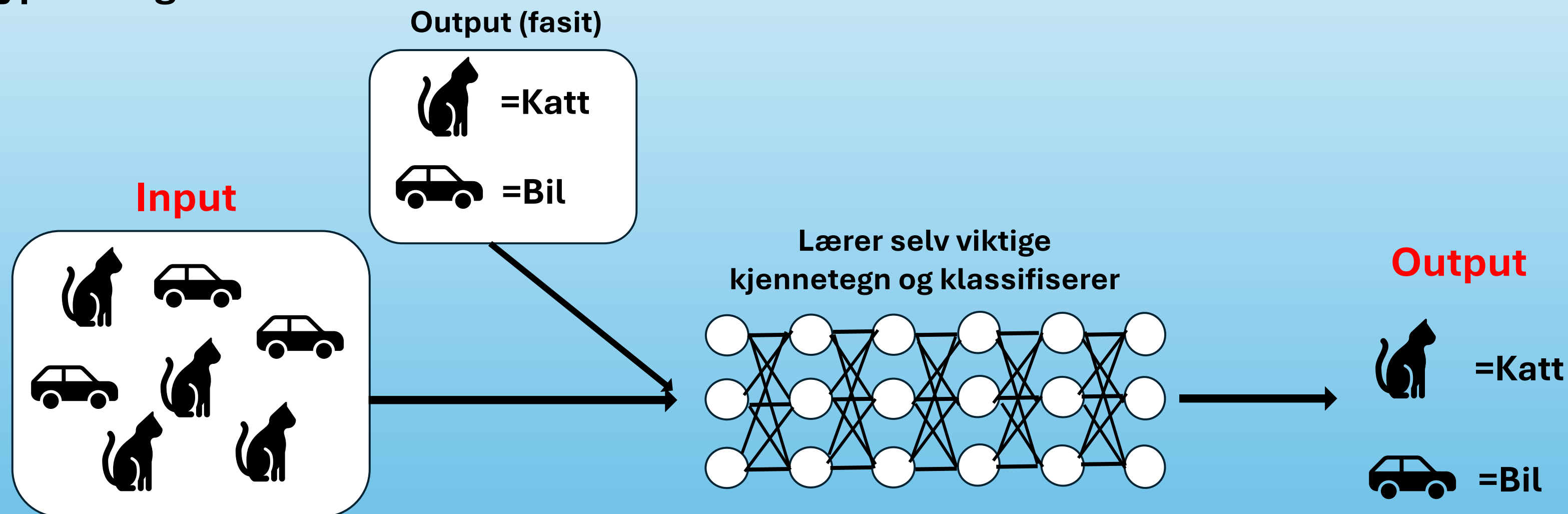
Teknologi som gjør det mulig for systemer å utføre oppgaver som vanligvis krever menneskelig intelligens, som læring, problemløsning og beslutningstaking, enten selvstendig eller med minimal menneskelig hjelp (1)



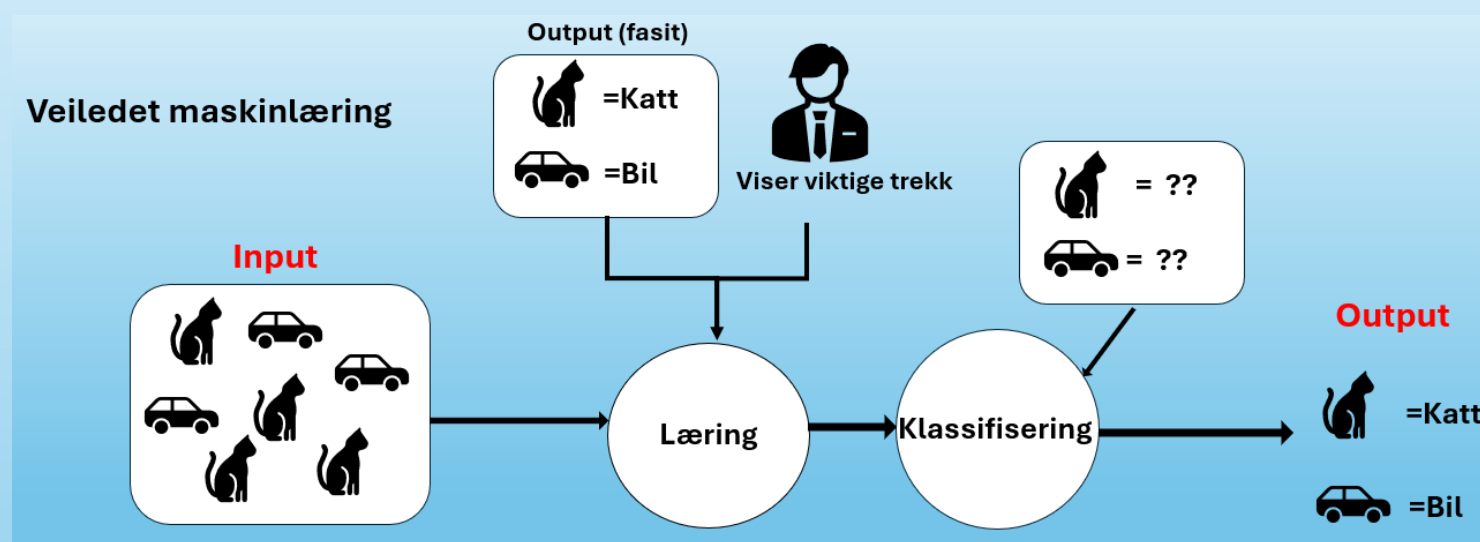
Veiledet maskinlæring



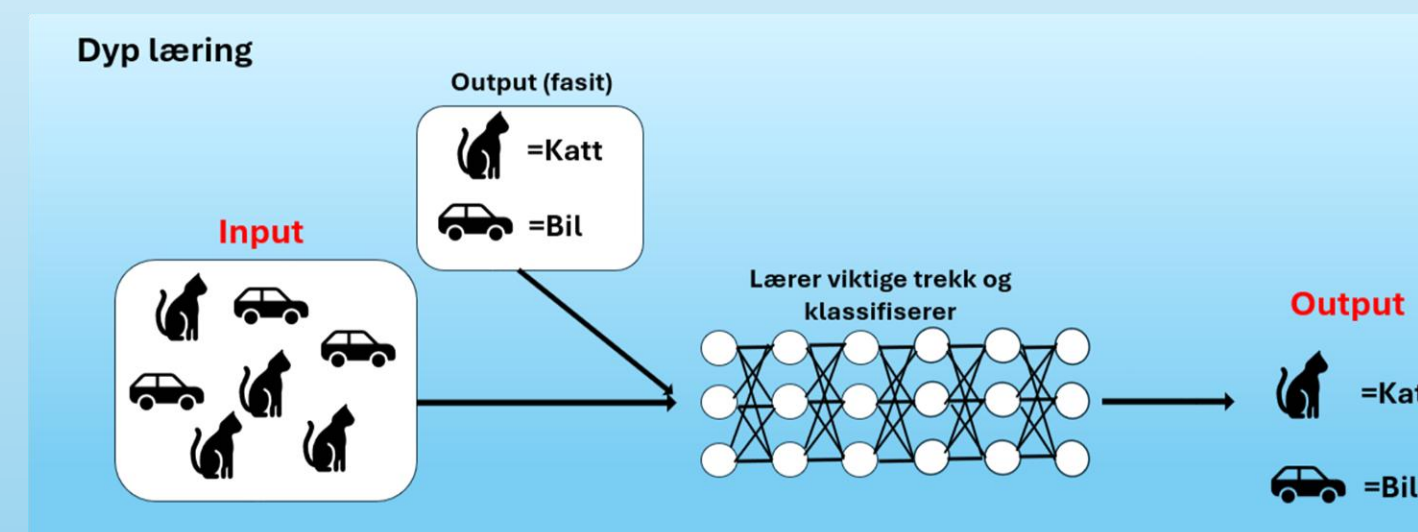
Dyp læring



Maskinlæring og dyp læring



- Krever strukturert data
- Relativt “lette” algoritmer
- Mulig å forstå logikken og tolke resultatene



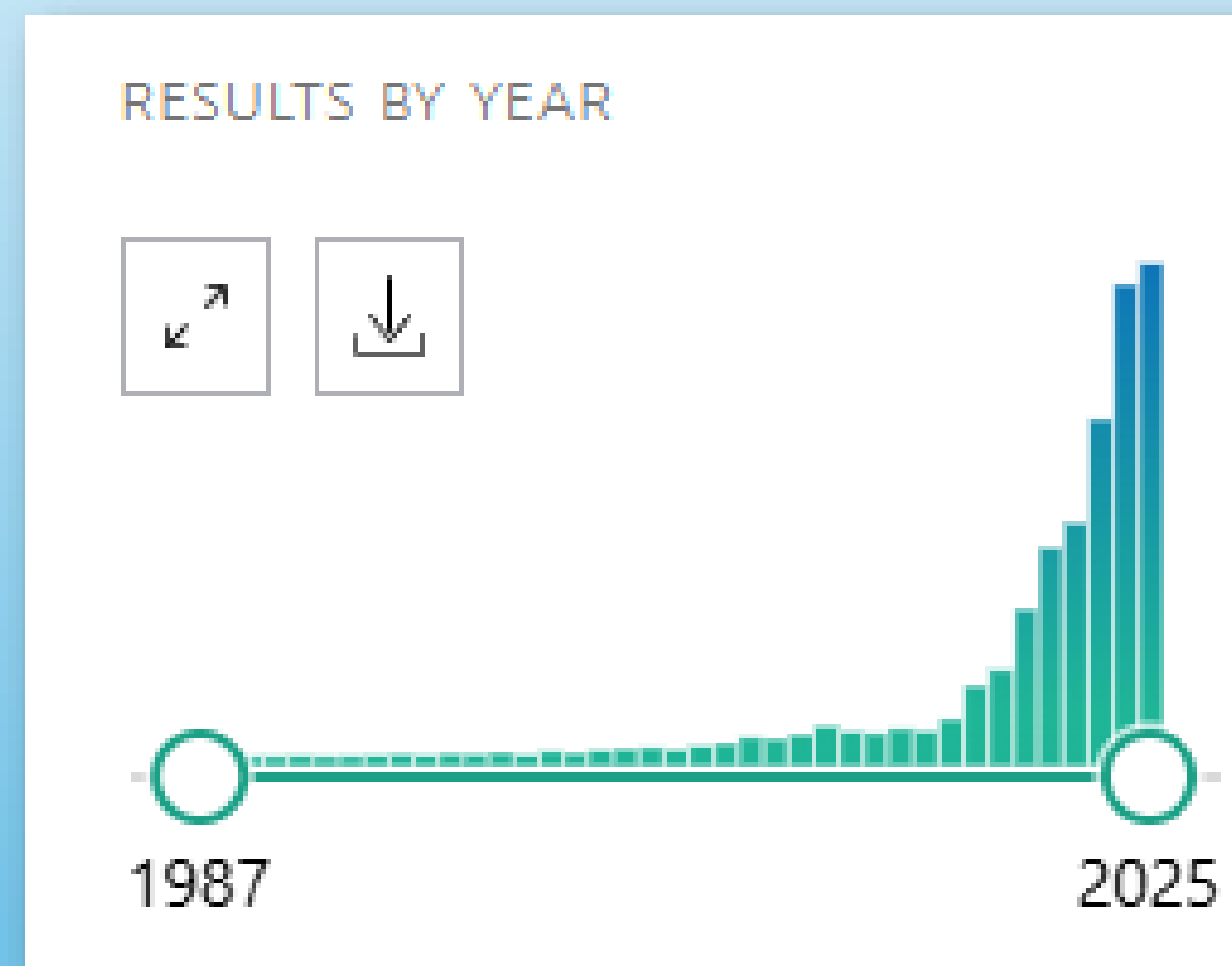
- Kan håndtere **u**strukturert data
- Kan gi mer sofistikerte svar
- Komplisert å forstå
- Datagrunnlag og logikken kan være ukjent (Black box)

KI som en del av KBP er et populært tema

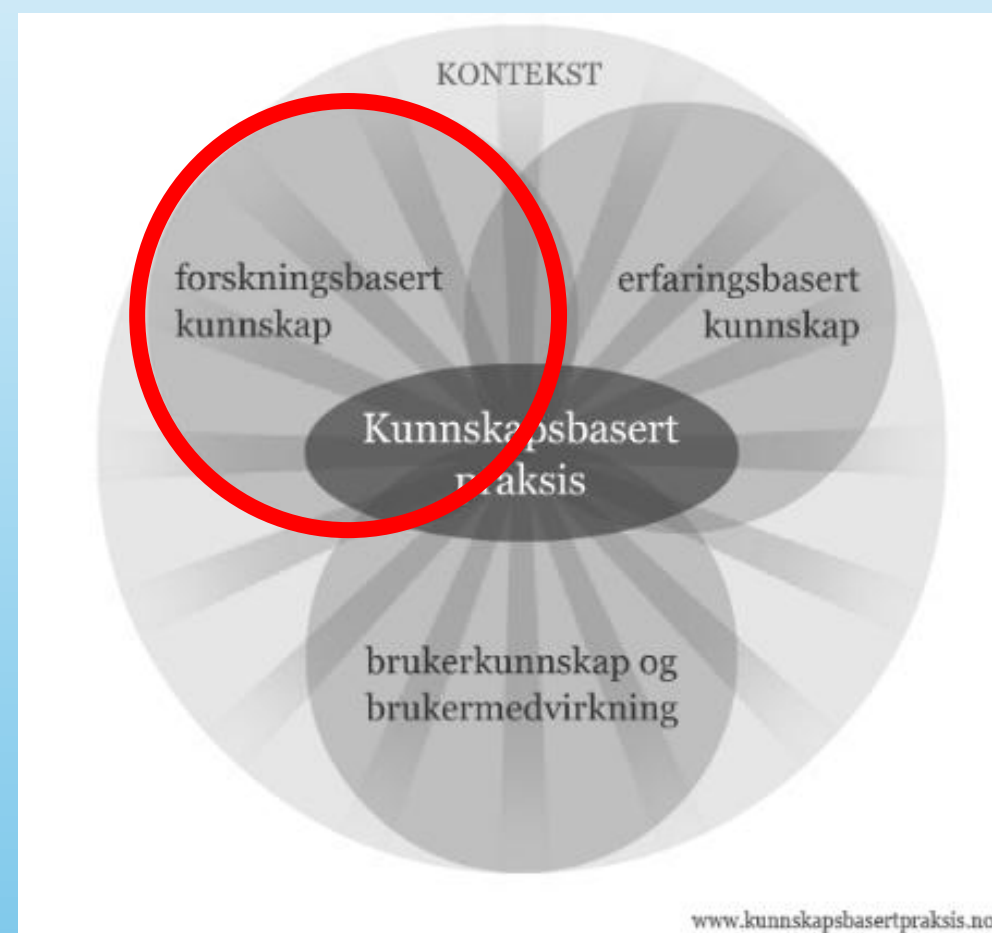
“Towards evidence-based practice 2.0: leveraging artificial intelligence in healthcare” (1)

“Artificial Intelligence in Evidence-Based Medicine: On the Verge of Breakthrough” (2)

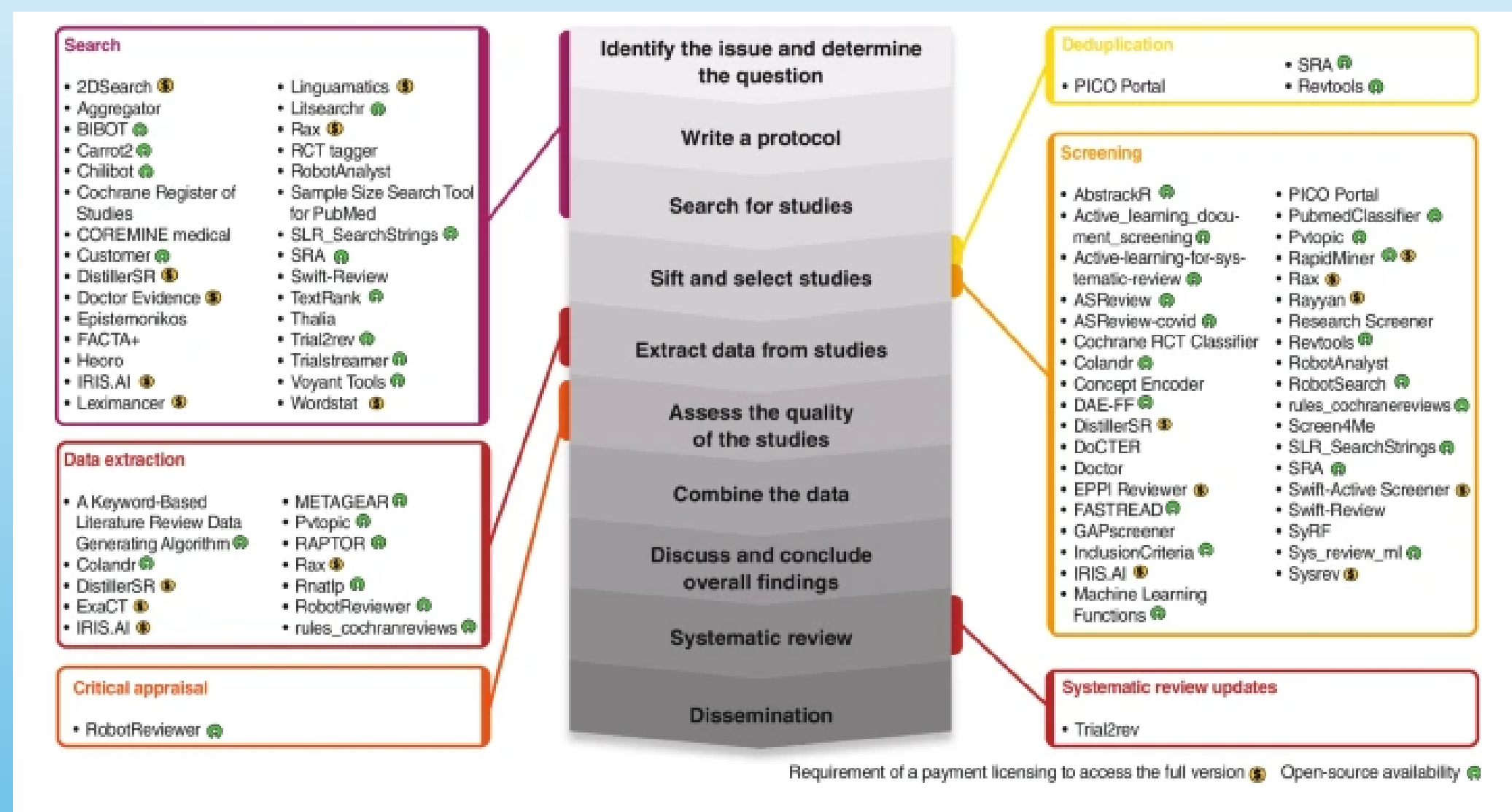
“New evidence-based practice: Artificial intelligence as a barrier breaker” (3)



Bruk av KI i systematiske oversikter



(Kunnskapsbasertpraksis.no)



(Cierco Jimenez et al., 2022)

1. Blaizot A, Veettil SK, Saidoung P, Moreno-Garcia CF, Wiratunga N, Aceves-Martins M, et al. Using artificial intelligence methods for systematic review in health sciences: A systematic review. Research synthesis methods. 2022;13(3):353-62.
2. Ge L, Agrawal R, Singer M, Kannapiran P, De Castro Molina JA, Teow KL, et al. Leveraging artificial intelligence to enhance systematic reviews in health research: advanced tools and challenges. Systematic Reviews. 2024;13(1):269.
3. Cierco Jimenez R, Lee T, Rosillo N, Cordova R, Cree IA, Gonzalez A, et al. Machine learning computational tools to assist the performance of systematic reviews: A mapping review. BMC Med Res Methodol. 2022;22(1):322.

KI til screening av artikler



KI til screening av artikler

Eksempel på KI-screening: (1)

- Forskerne startet screening av artikler (**input**)
- Markerte artikler som relevant/irrelevant (**fasit**)
- Verktøyet lærer av forskeren (**viktige kjennetegn**)
- Verktøyet rangerer artiklene basert på relevans (**klassifisering**)
- Forskerne vurderte så høyest rangert artikkel én etter én
- Forskerne stoppet etter 100 irrelevante **på rad**

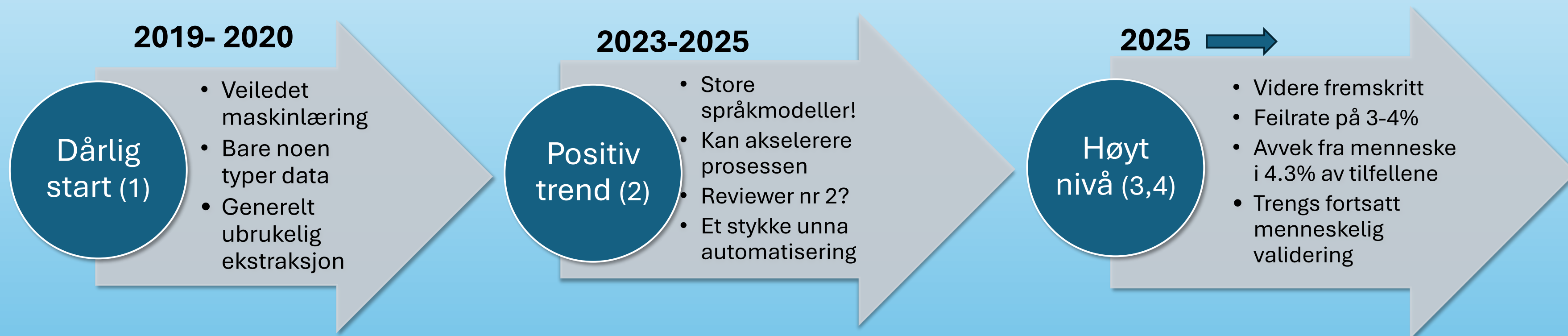
Resultat:

Screening stoppet etter
1063 av 4695 artikler (23%)

Mye tid spart!

Modellen prisgitt
forskeren

KI til dataekstraksjon



1. Blaizot A, Veettil SK, Saidoung P, Moreno-Garcia CF, Wiratunga N, Aceves-Martins M, et al. Using artificial intelligence methods for systematic review in health sciences: A systematic review. Research synthesis methods. 2022;13(3):353-62.

2. Schmidt L, Finnerty Mutlu AN, Elmore R, Olorisade BK, Thomas J, Higgins JPT. Data extraction methods for systematic review (semi)automation: Update of a living systematic review. F1000Res [Internet]. 2025; 10:[401 p.]. Available from: <https://f1000research.com/articles/10-401/v3>.

3. Bianchi J, Hirt J, Vogt M, Vetsch J. Data Extractions Using a Large Language Model (Elicit) and Human Reviewers in Randomized Controlled Trials: A Systematic Comparison. [Chichester] :2025.

4. Helms Andersen T, Marcussen TM, Termannsen AD, Lawaetz TWH, Nørgaard O. Using Artificial Intelligence Tools as Second Reviewers for Data Extraction in Systematic Reviews: A Performance Comparison of Two AI Tools Against Human Reviewers. Cochrane Evid Synth Methods. 2025;3(4):e70036.

Bruk av KI i systematiske oversikter

Ansvarlig bruk av KI i systematiske oversikter:

Gruppens mål:

- Definere og støtte ansvarlig bruk av KI
- Sikre at KI ikke går på bekostning av forskningsintegritet



KI i kunnskapsbaserte retningslinjer

› [Ann Intern Med.](#) 2025 Mar;178(3):408-415. doi: 10.7326/ANNALS-24-02338. Epub 2025 Jan 28.

Guidelines International Network: Principles for Use of Artificial Intelligence in the Health Guideline Enterprise

Bernardo Sousa-Pinto¹, Manuel Marques-Cruz¹, Ignacio Neumann², Yuan Chi³, Artur J Nowak⁴, Marge Reinap⁵, Mariette Awad⁶, Monika Nothacker⁷, Milana Trucl⁸, Jan Brozek⁹, Pablo Alonso-Coello¹⁰, Wojtek Wiercioch¹¹, Amir Qaseem¹², Elie A Akl¹³, Holger J Schünemann¹⁴; 2024 Board of Trustees of the Guidelines International Network

Collaborators, Affiliations + expand

PMID: 39869912 DOI: 10.7326/ANNALS-24-02338

Free article

MAGIC Evidence
Ecosystem
Foundation

Strategy 2024 - 2029



Bridging the gap between evidence and patient care

KI som hjelp til å formulere spørsmål og lage søkestrategi:

Resultat (ChatGPT, Bing og Bard):

- Alle programmene lagde spørsmål og PICO-elementer
- Noen retningsbestemte ord i O-elementet (eks: increased)
- Noen litt for ordrike PICO elementer
- Litt varierende nøyaktighet og omfang på søkestrenger

Konklusjon:

- Alle programmene viser stort potensial
- Kan være en støtte i utførelse av aktuelle KBP-trinn
- Bedre kvalitet og mindre bruk av tid på søk

Received: 31 January 2024 | Revised: 24 October 2024 | Accepted: 7 November 2024
DOI: 10.1111/jnu.13036

CLINICAL SCHOLARSHIP

JOURNAL OF NURSING
SCHOLARSHIP

PICOT questions and search strategies formulation: A novel approach using artificial intelligence automation

Lucija Gosak PhD candidate, MSc, RN¹ | Gregor Štiglic PhD, BSc^{1,2,3} |
Lisiane Pruinelli PhD, MSc, RN⁴ | Dominika Vrbnjak PhD, MSc, RN¹

¹Faculty of Health Sciences, University of Maribor, Maribor, Slovenia

²Faculty of Electrical Engineering and Computer Science, University of Maribor, Maribor, Slovenia

³Usher Institute, University of Edinburgh, Edinburgh, UK

⁴College of Nursing and College of Medicine, University of Florida, Gainesville, Florida, USA

Correspondence

Abstract

Aim: The aim of this study was to evaluate and compare artificial intelligence (AI)-based large language models (LLMs) (ChatGPT-3.5, Bing, and Bard) with human-based formulations in generating relevant clinical queries, using comprehensive methodological evaluations.

Methods: To interact with the major LLMs ChatGPT-3.5, Bing Chat, and Google Bard, scripts and prompts were designed to formulate PICOT (population, intervention, comparison, outcome, time) clinical questions and search strategies. Quality of the LLMs

KI som hjelp til å forstå forskning:

Mange klinikere:

- Har begrenset kunnskap om forskningsterminologi
- Og engelsk språk kan være en barriere

NotebookLM:

- Oppsummerer på Norsk
- Foreslår presentasjoner
- Lager podcast!

Correspondence | Published: 02 May 2025

Generative AI in clinical practice: novel qualitative evidence of risk and responsible use of Google's NotebookLM

[Max Reuter](#), [Maura Philippone](#), [Bond Benton](#) & [Laura Dilley](#) 

[Eye](#) 39, 1650–1652 (2025) | [Cite this article](#)

337 Accesses | 2 Citations | 10 Altmeteric | [Metrics](#)

The advent of generative artificial intelligence, especially large language models (LLMs), presents opportunities for innovation in research, clinical practice, and education. Recently, Dihan et al. [1] lauded LLM tool NotebookLM's potential, including for generating AI-voiced

Resultater Reuter et al., 2025:

- Kan hallusinere
- Faktasjekker ikke kildene
- 2+2= 5?

Bruk av KI for å hjelpe den enkelte kliniker



(Kunnskapsbasertpraksis.no)

Kan KI bli fremtidens digitale sparringspartner for helsepersonell?

Kan man forvente pålitelige svar fra KI på medisinske/ helsefaglige spørsmål?



Bruk av KI for å hjelpe den enkelte kliniker

KI som en digital sparringspartner:

Analysis

Meet generative AI... your new shared decision-making assistant

Glyn Elwyn ,¹ Padhraig Ryan,² Daniel Blumkin,³ William B Weeks⁴

10.1136/bmjebm-2023-112651

Introduction
Large language models (LLMs), a new family of tools to generate natural language, are likely to radically change how some aspects of healthcare are delivered.¹ It is difficult to comprehend the pace at which change is already occurring. For example, ChatGPT, released by OpenAI in November 2022, represented a leap in the capability of artificial intelligence (AI).² This was followed by the release of the Gemini family of models from Google and Claude-3 from Anthropic, as well as numerous other models with impressive capabilities. The aim of this analysis article is to illustrate how LLMs can already be used to improve healthcare based on their power to find and convey information. At the same time, we want to alert users to the pitfalls of relying on this form of AI: vigilance will

medical knowledge behind paywalls does not inform its training.
Concerns have been voiced, of course, at many levels.⁷⁻⁸ LLM responses can have errors. There are regular reports of 'hallucinations', where LLM responses are not factual.⁹ Nevertheless, these applications have been released. Millions have used LLMs and are learning how to use them with increasing levels of expertise. This has led to the sudden emergence of a new field known as prompt engineering, the strategic design and use of prompts to elicit desired responses from LLMs. Indeed, the concept of 'prompt engineers' is one of the fastest-growing areas of job development, highlighting the need for a new skill: how best to harness the power of LLMs.¹⁰
To illustrate how LLMs can be used to support

BMJ EBM: first published as 10.1136/bmjebm-2023-112651 on 12 June 2024
Protected by copyright, including for uses n

¹The Dartmouth Institute for Health Policy and Clinical Practice, Dartmouth College, Hanover, New Hampshire, USA
²Eron Corporation, Dublin, Ireland
³Dartmouth College, Hanover, New Hampshire, USA
⁴Microsoft Corp, Redmond, Washington, USA
Correspondence to: Professor Glyn Elwyn; glynelwyn@gmail.com

Resultat (1):

- Predikere sannsynlige spørsmål
- Komme med forslag til svar (muntlig og skriftlig)
- Gjøre kliniker bedre forberedt

Tidsskriftet

ARTIKLER FAGOMRÅDER UTGAVER PODKAST PUBLISERING LEGEJOBBER SØK

KORT RAPPORT

Kunstig intelligens og legers svar på helsespørsmål

TEKNOLOGI

ENGLISH

SAMMENDRAG

HOVEDFUNN

ARTIKKEL

INNLEDNING

MATERIALE OG METODE

Tiril Egset Mork, Håkon Garnes Mjøs*, Harald Giskegjerde Nilsen, Sindre Kjelsrud, Alexander Selvikvåg Lundervold, Arvid Lundervold**, Ib Jammer Om forfatterne*

Bakgrunn
Flere studier har undersøkt hvordan store språkmodeller besvarer helsespørsmål. I en studie fra 2023 ble svar på engelskspråklige helserelaterte spørsmål generert av språkmodellen GPT-3.5, oppfattet som mer empatiske og kunnskapsrike enn svar fra leger. Vi ønsket å anvende den nvere språkmodellen GPT-4 nå norsk for å

Publisert: 10. februar 2025
Utgave 2, 11. februar 2025
Tidsskr Nor Legeforen 2025 Vol. 145.
doi: 10.4045/tidsskr.24.0402
Mottatt 28.7.2024, første revisjon innsendt 16.10.2024, godkjent 14.12.2024.

Resultat (2):

- Mer empatiske svar
- Mer informative svar
- Kan KI gjøre oss bedre på kommunikasjon?

Kan vi forvente pålitelige svar på medisinske/helsefaglige spørsmål fra KI i 2025?

Kan KI hjelpe oss med å få til gode, kunnskapsbaserte, og pasientsentrerte beslutninger?

- Blir mer og mer nøyaktig!
- For mye fokus på nøyaktighet? (1)
- For lite fokus på sikkerhet, rettferdighet, kompleksitet? (1)

Mine tanker er:

- Stoler ikke helt på deg, men...
- Du kan være hjelpsom å diskutere med!

Sikt-KI chat sin egen oppsummering

*«I 2025 kan klinikeren bruke store språkmodeller som støtteverktøy for å finne informasjon og utforske medisinske spørsmål, men modellene er fortsatt **ikke pålitelige nok til å stå alene** i klinisk praksis. Menneskelig vurdering, verifisering og kontekstualisering av svarene er avgjørende for å unngå feil og sikre pasientsikkerhet. All informasjon jeg gir bør **alltid dobbeltsjekkes med kvalifisert helsepersonell** før den brukes i praksis»*

“Prioritizing ethical considerations is essential to harness AI's potential while preserving trust and integrity in healthcare and research practices. Acknowledging and addressing challenges such as ethical concerns, bias and the potential for overreliance is crucial” (1)

*“ **Tillitsverdige KI** er gyldig og pålitelig, trygg, sikker og motstandsdyktig, ansvarlig og gjennomsigtig, forklarbar og tolkbar, har godt personvern og er rettferdig med kontroll på skadelige skjevheter” (2)*



Egenskaper med KI som utfordre tillitsverdige KI:

Systematiske feil (Bias):

- når KI er trent på data med skjevheter
- Kan føre til feil og diskriminering

Hallusinasjon:

- Store språkmodeller har ikke begrep på «sannhet»
- Kan generere rent oppspinn

Sort-boks-problemet:

- KI som benytter dyp læring kan være umulig å forstå
- Datagrunnlag og logikk kan være ukjent
- Utfordre personvern og datasikkerhet

Tillitsverdig KI i den norske helsetjenesten:

Tillitsverdig KI krever:

- Tydelige rammer for å sikre lovlig, etisk og robust bruk av KI (1)

Rapporten omhandler:

- Grunnleggende regelverk og forordninger (2)
- Vurderinger som må tas før KI tas i bruk

Eksempel:

- CE-merking av KI-systemer som skal:

«Brukes på mennesker i den hensikt å, bidra til diagnostisering, forebygging, overvåking, prediksjon, prognostisering, behandling eller lindring av sykdom» (1)



Tillitsverdige KI i helseutdanningene:

- Fremtidens helsestudenter trenger KI-kompetanse for å sikre etisk, ansvarlig og kritisk bruk
- Mye bra gjøres i dag! (eksempel fra OsloMet)
 - Tilgang til trygge KI-verktøy
 - Gode resurser for undervisere for ansvarlig bruk av KI
- **Mer fokus på KI:**
 - Egne KI-emner i helseutdanningene
 - Nytt felt, men økende kunnskapsgrunnlag (1-5)

Kunstig intelligens (KI) ved OsloMet

På denne siden finner du blant annet informasjon om hvorfor det kan være lurt å bruke KI, hvilke standpunkt OsloMet tar til bruk av KI og hvilke KI-verktøy du har tilgang til som ansatt ved OsloMet.

[Bruke Chat GPT](#)[Policy for bruk av KI](#)[KI i undervisning](#)[KI i forskning](#)[Bruk av KI i innholdsproduksjon på nettsider](#)[Personvern hensyn ved KI i forskning](#)

FYB3100 Digital kompetanse og innovasjon i helse

Emneplan

Engelsk emnenavn

Digital Competence and Innovation in health

Studieprogram

Bachelorstudium i fysioterapi

Omfang

5 stp.

Studieår

2025/2026

Programplan

[Bachelorstudium i fysioterapi \(2025 HØST\)](#)

Emnehistorikk

2025 / 2026



Bruk av kunstig intelligens i kunnskapsbasert praksis – Muligheter og utfordringer ved bruk av kunstig intelligens som støtte for helsepersonell i kunnskapsbasert praksis

