



Fremtidens internrevisjon, risk & Compliance

Sammen med dere i dag

Agenda

- 01 Fremtidens internrevisjon – fra kontrollfunksjon til strategisk verdiskaper
- 02 Fremtidens risk- og compliancefunksjon
- 03 Hvordan sikre at AI-modeller er pålitelige?



Kenneth Hansen
Head of Internal Audit



Tommy Knutsen
Senior Manager,
Internal audit &
Evaluation



Gisle Nerbråten
Head of
Regulatory
Transformation

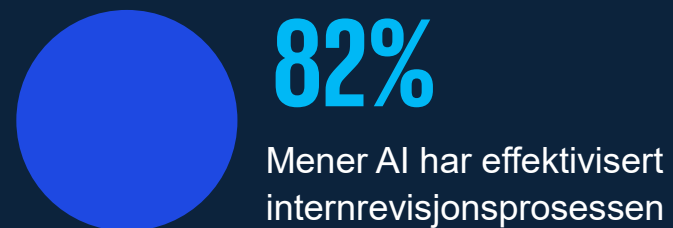


**Thor Aage
Dragsten**
Head of Risk
Analytics

Verden og internrevisjon

- Dynamisk og komplekst risikobilde
- Økte forventninger fra nøkkelinteressenter
- Nye IIA-standarder

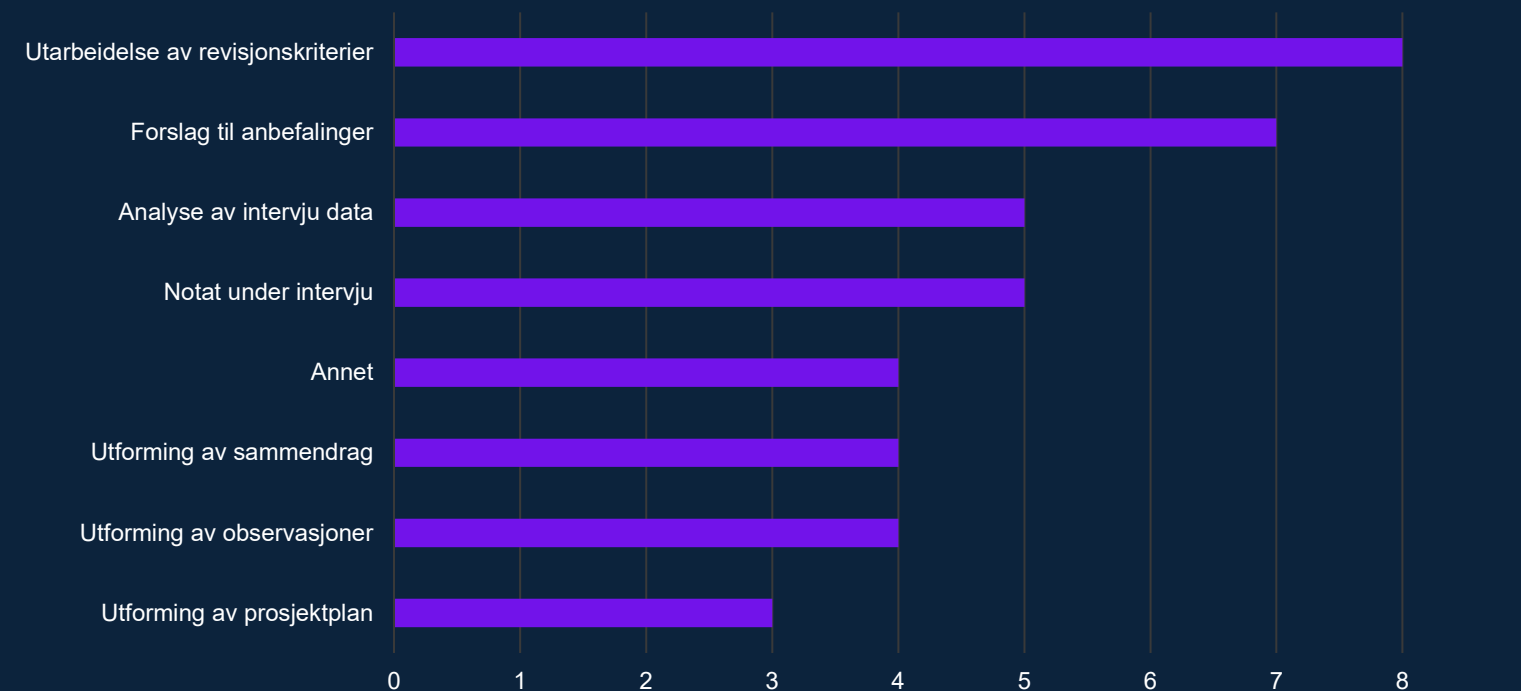
Nåsituasjon: Bruk av AI i internrevisjonen IR



Verktøy og teknikker vi bruker

CoPilot, KaiChat, AdvisoryGPT, Perplexity, Celonis, Relativity, Process Mining

På disse måtene bruker vi AI i internrevisjonsprosjekter



Største styrker

- Mer effektiv gjennomføring
- Identifisere helhetlige mønstre
- Bedre kvalitet

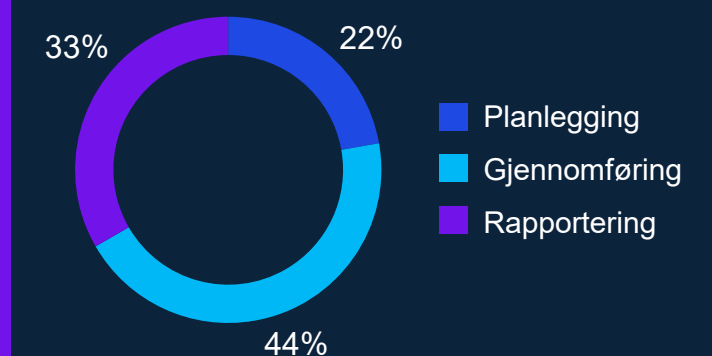
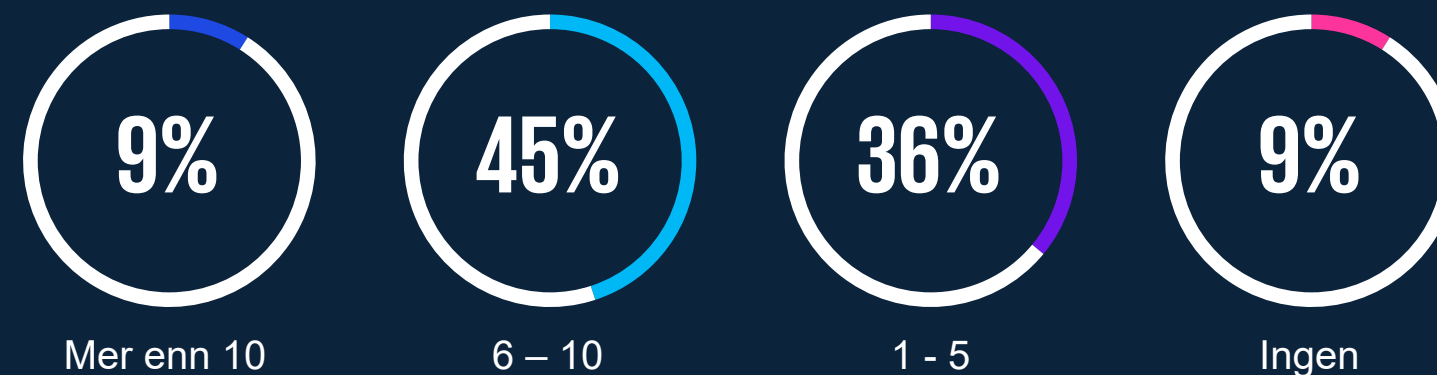
Største svakheter

- Manglende AI-innsikt i revisjon
- Kvalitet på tilgjengelige verktøy
- Kompetanse på AI

26%

Deltakere estimerer i gjennomsnitt at bruk av AI har effektivisert internrevisjonsprosjekter med 26%

Antall prosjekter der vi har benyttet AI



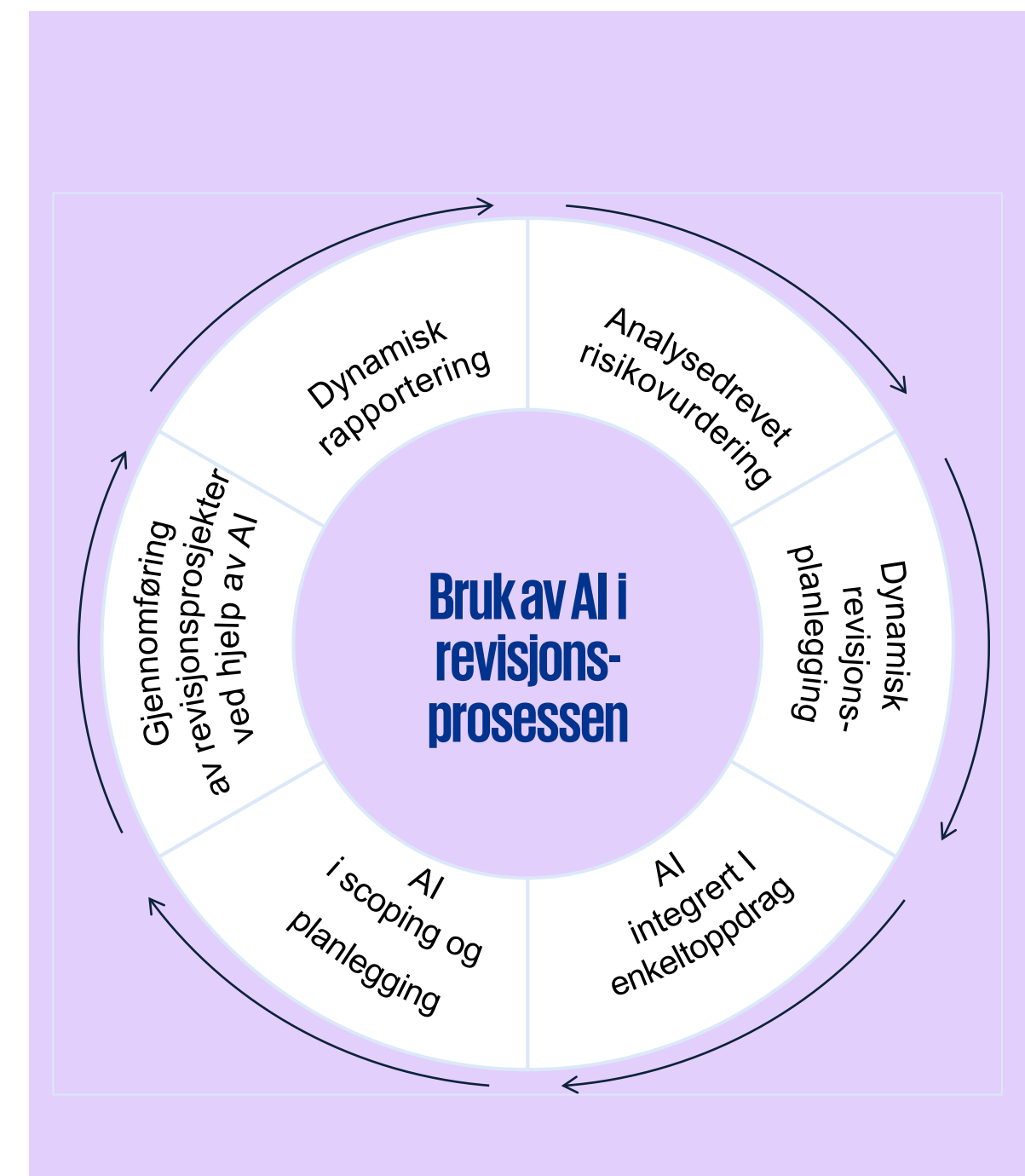
Flest ser størst økning i effektivitet i gjennomføringen av internrevisjonsprosjekter

Planleggingsprosess

AI benyttes i ulike deler av planleggings- og rapporteringsprosessen.

- Risikovurdering og risikofasilitering
- Monitorering av regulatoriske trender
- I enkeltrevisjoner

Eksempl



Revisjonskriterier

AI er effektiv til å foreslå revisjonskriterier.

Eksempel

«Det kule selskapet»

«Opprettholde en positiv selskapskultur med fokus på medarbeidertrivsel.

Fremme innovasjon, bl.a. gjennom å oppmuntre ansatte til å tenke utenfor boksen.

Tilby fleksible løsninger for ansatte (hjemmekontor, fleksible arbeidstider mv.).

Ta samfunnsansvar og bidra til en bærekraftig utvikling.»

MEN!

Revisjonskriterier må vurderes og kvalitetssikres av et menneske. AI kan benyttes til å lage utkast.

Vær åpen overfor revisjonsobjektet om hvordan AI benyttes.

Svakere på spesialiserte rettsområder.

Dokumentanalyse

Ved hjelp av AI sammenfattes store datamengder.

Eksempel

Gjennomgang av 100 tilsyn fra Finanstilsynet

Gjør automatisk uttrekk av ønskede tilsyn hos Finanstilsynet ved å laste ned hundrevis av tilsyn på få minutter. Deretter benytter vi en AI API til å gjennomgå og analysere tilsynsrapportene ved å trekke ut hovedfunn gjennom ønskede «prompts».

Kombinerte Python, HTML, ChatGPT og Microsoft 365 Copilot.

MEN!

Terskel for akseptabel kvalitet opp mot behov for effektivitet må etableres.

Menneskelig kvalitetskontroll/stikkprøver må gjennomføres.

Begrenses til offentlige tilgjengelig data for å sikre ansvarlig bruk.

Intervjuer

AI benyttes i ulike faser av en intervju- eller workshopprosess.

Eksempl

Forslag til intervjuguide: Bruk AI til å foreslå intervjuguide, f.eks. basert på funn i dokumentanalyse.

Ta referat fra samtalene: Bruk Microsoft 365 Copilot til å transkribere intervjuet. Kan gjøres både i Teams og fysisk.

Oppsummere hovedpunkter: Sentrale hovedpunkter kan trekkes ut for å påbegynne analysen av intervjudata.

MEN!

Samtykke må innhentes.

Sensitiviteten i informasjonen må vurderes.

Oppsummering

AI bidrar til å transformere IR **fra kontrollfunksjon til strategisk verdiskaper.**

AI er relevant i **alle** deler av revisjonsprosessen.

Ansvarlig bruk og **internrevisors integritet må vurderes kontinuerlig.**

AI skal bidra til **økt kvalitet**, ikke bare effektivitet.

Utfordringer for risk- og compliancefunksjonene

Omfang og utvikling i det regulatoriske rammeverket for finansforetak er økende. Mange andrelinjefunksjoner opplever utfordringer knyttet til kapasitet, metodikk for testing og for lite tid til rådgivning og vurderinger av nye regelverk. Manglende kapasitet knyttes gjerne til at mye tid brukes på rapportering, og at arbeidsprosesser er manuelle. AI kan bidra til å effektivisere, standardisere og gi bedre innsikt i etterlevelsesrisikoen som foretaket står overfor

RAPPORTERING

- 2. linje opplever at for mye av tiden benyttes til styre- og ledelsesrapportering.
- Rapporteringen er tidkrevende og representerer ofte oppdateringer heller enn nye vurderinger



KONTROLL

- 2. linje opplever ofte lite av tiden benyttes til å gjennomføre kontroller, og at for få kontroller gjennomføres basert på omfanget av kontrollområder



OVERVÅKNING

- Flere opplever at kontrollene som gjennomføres ikke gir nok trygghet på at brudd identifiseres, på grunn av manglende kapasitet

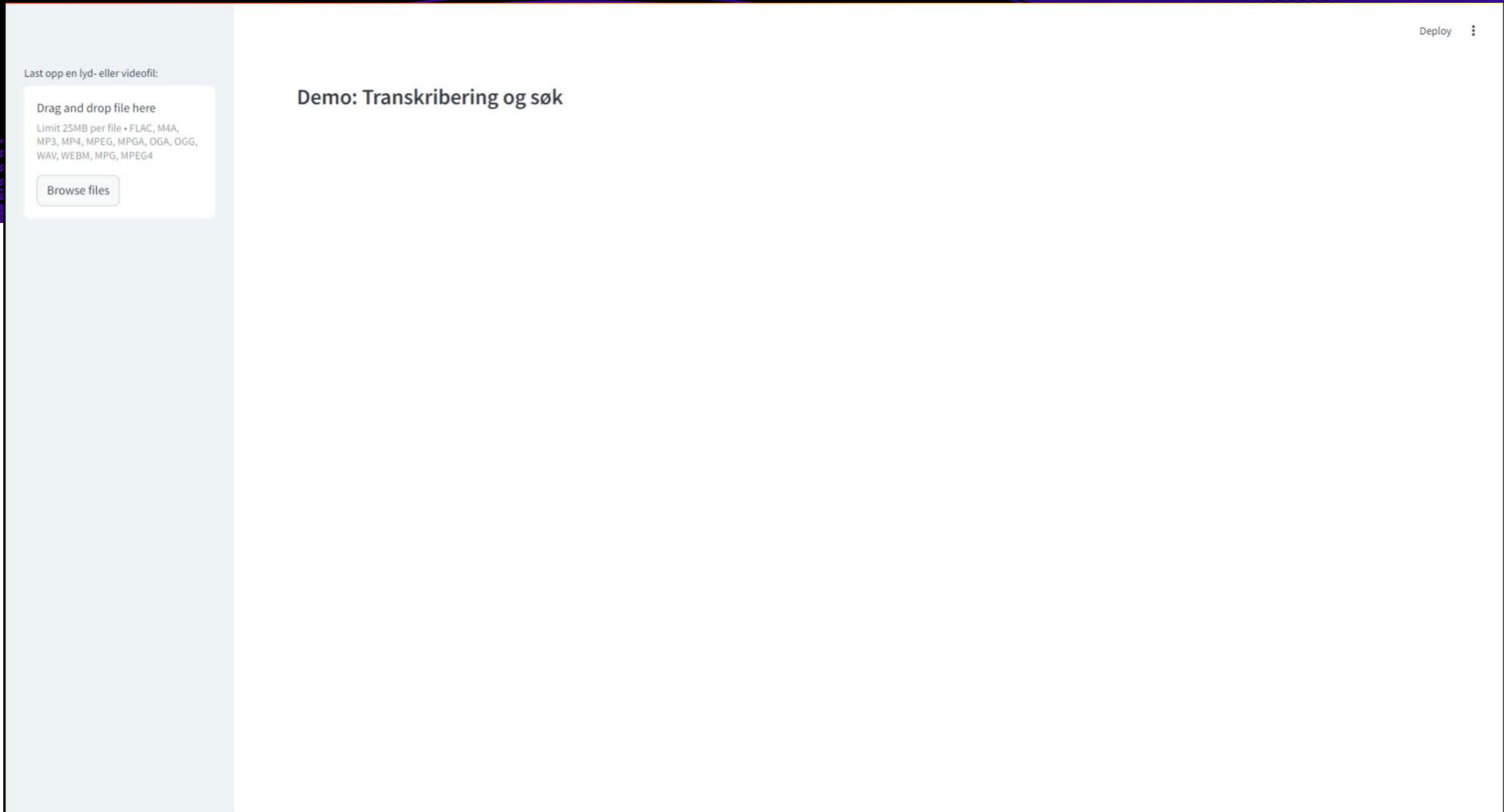


RÅDGIVNING

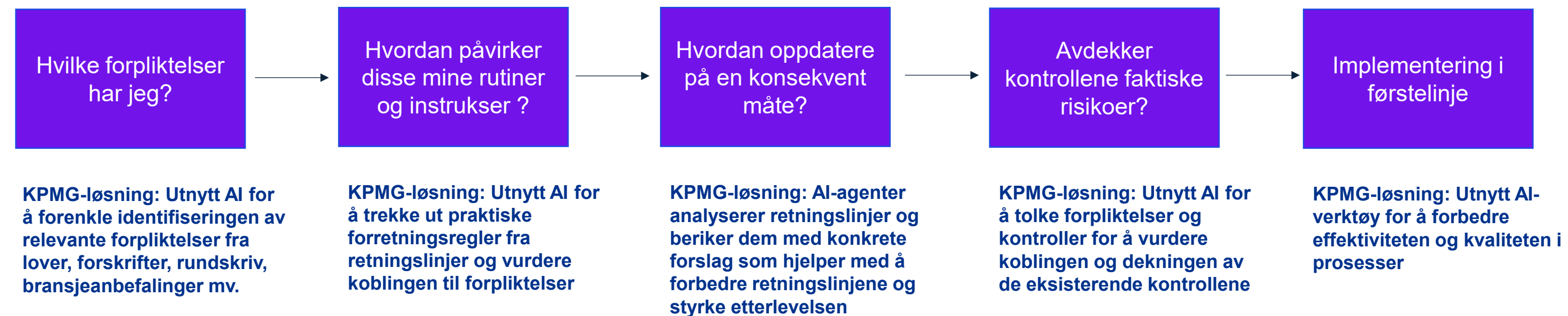
- Compliance-avdelinger opplever ofte at de har for lite kapasitet til å utøve sin rådgivende rolle for å bistå organisasjonen med å forstå hvordan nye regelverk vil påvirke virksomheten forretningsmessig



Effektivisering og kvalitetsforbedring ved bruk av RegTech – AI i compliance: Demo



Et E-risikog compliance



Vår Compliance Tracking AI-løsning, muliggjort av Kym

Ved å utnytte de store språkegenskapene til KPMGs generative AI-plattform, Kym, setter Compliance AI sammen og analyserer data for å identifisere sammenhenger og gap. Ytterligere perspektiver kan oppnås ved å bruke Kym til å vurdere nøkkelkomponenter i risikostyrings- og compliance-kunnskapsbasen mot KPMGs beste praksis-rammeverk og oversikt over kontroller av høy kvalitet.

Input, data og kilder:

Regulatoriske kilder

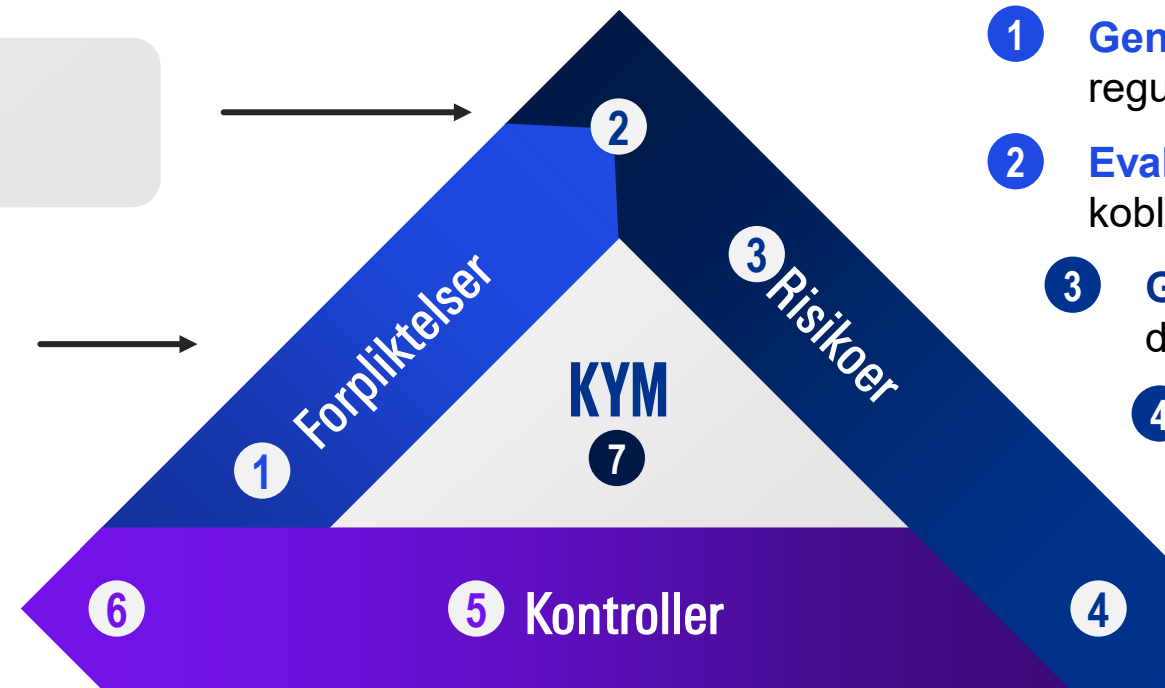
F.eks. lovgivning, forskrifter, guidelines, etc.

GRC / RMS

Eksisterende oversikt over forpliktelser, risikoer og kontroller.

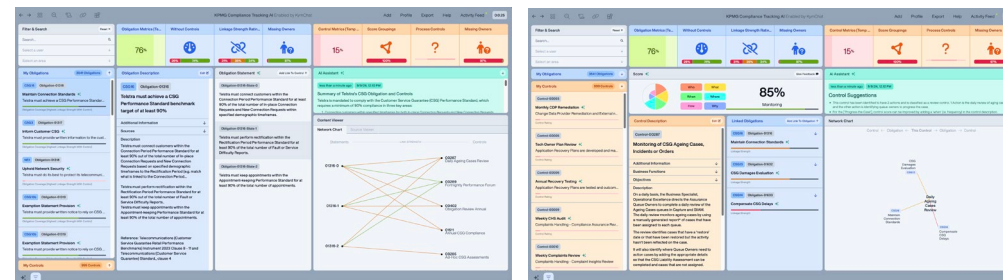
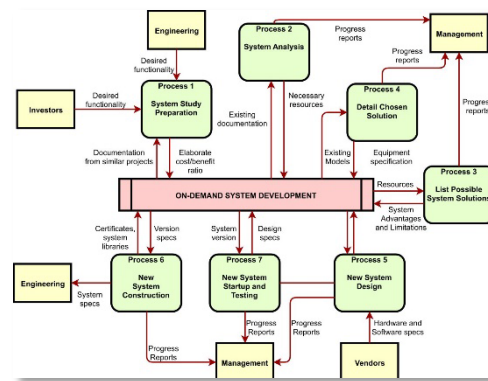
Operasjonelle dokumenter

Prosesskart, policyer, prosedyrer, etc.



COMPLIANCE AI:

- 1 **Genererer en oversikt over forpliktelser**, inkludert sporbare referanser til relevante regulatoriske krav.
- 2 **Evaluerer styrken av eksisterende koblinger mellom forpliktelser og risiko**. Nye koblinger identifiseres og anbefales ved bruk av det brede GRC/datagrunnlaget.
- 3 **Genererer et risikoregister** ved hjelp av regulatoriske krav og operasjonell dokumentasjon.
- 4 **Evaluerer styrken av eksisterende koblinger mellom risiko og kontroll**. Nye koblinger identifiseres og anbefales ved bruk av det brede GRC/datagrunnlaget.
- 5 **Vurderer kontrollbeskrivelser** i GRC/RMS-data mot KPMGs beste praksis-rammeverk og eksempler på kontroller. Nye kontroller anbefales for å adressere forpliktelser og risiko.
- 6 **Evaluerer styrken av eksisterende koblinger mellom kontroll og forpliktelser**. Nye koblinger identifiseres og anbefales ved bruk av det brede GRC/datagrunnlaget.
- 7 **Gjør det mulig å stille enkle spørsmål på engelsk til AI-assistenten** for å kontrollere deres kunnskap om forpliktelser og undersøke detaljene i kontroll- og risikovurderinger.



Hvordan AI kan skape verdi



Risikostyring

Compliance og kontroll

Før AI...

Risikostyringsprosesser er sterkt avhengig av manuell datainnsamling og -analyser.

Avhengighet av statiske risikomodeller. Dette resulterer i forsinket respons på fremvoksende risikoer.

Risikostyringstjenester levert i siloer, med minimal integrasjon på tvers av virksomheten.

Styringsprosesser overveiende manuelle, med fokus på compliance og kontroll, og begrenset smidighet i håndtering av nye risikoer.

Avhengighet av historiske data og personlig vurdering i styringsbeslutninger, mindre treffsikker prediktiv innsikt.

Ineffektive styringsmekanismer, med forsinket respons på endrede regulatoriske og risikomessige forhold.

Nå med AI...

➤ **Dynamisk risikomodellering og sanntidsrisikoanalyse muliggjort av AI, som gir umiddelbare innsikter i fremvoksende risikoer.**

➤ **AI-drevet automatisering effektiviserer prosessene for compliance-overvåking og risikovurdering, og reduserer manuelt arbeid og feilrater betydelig.**

➤ **Desentralisering av risikostyringsaktiviteter tilrettelagt av AI, som muliggjør mer smidige og responsive risikopraksiser på tvers av virksomheten.**

➤ **Implementering av AI-drevne styringsrammeverk, som forbedrer synlighet av risiko og compliance gjennom automatisert overvåking og rapportering.**

➤ **Innføring av dynamiske risikostyringspraksiser, som utnytter AI for sanntidsrisikovurdering og støtte til beslutningstaking.**

➤ **Forsterket risikostyring med AI-aktivert prediktiv analyse, som sikrer en proaktiv og smidig respons på regulatoriske endringer og fremvoksende risikoer.**

Verdiskaping*

Forbedret produktivitet

- 35 - 50% mer kapasitet
- 44% automatisering av driftsoppgaver



Økt lønnsomhet

- 30 - 40% lavere kostnader for å betjene viktige risikoreducerende funksjoner



Redusert risiko

- 25% reduksjon i non-compliance gjennom forbedret overvåking
- 35% forbedring i å identifisere trusler



*Kilde: Potensiell verdi estimert basert på ekstrapolering av KPMG AI & Technology-undersøkelser og tverrfaglige analyser av virkningen av AI de neste 24–36 månedene



PåliteAI

Sikre rettferdighet og forklarbarhet

AI kan brukes i alle forretningsområder

Kundeanskaffelse

- Individualisering av kundetilbud inkludert kryssalg,
- Selvbetjente løsninger og forbedring av kundeopplevelsen
- Identifisering av markedspotensial



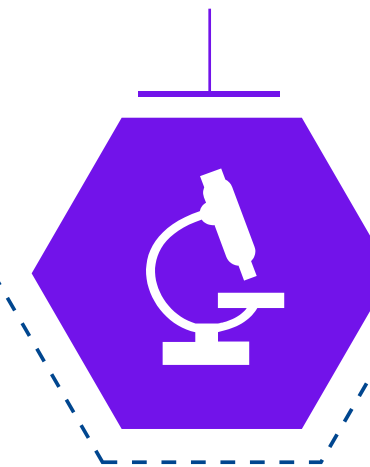
Databehandling

- Forbedring og automatisering, testing av datakvalitet
- Innholds- og prosessoptimalisering av intern og regulatorisk rapportering



Risikostyring

- Forbedring av beregningsmetodikk
- Forbedre datakvaliteten
- Intradag-kapasitet
- Prosessoptimalisering, validering



Onboarding

- Forbedring av kundeinformasjon inkl. KYC
- Forbedring av forretningsbeslutninger (kredittscoring, prising)



Etterlevelse og svindel

- Identifikasjon av hvitvasking av penger (AML)
- Oppdag konto-/kredittkortsvindel
- Støtte for cybersikkerhet
- Overvåking av forhandleraktiviteter



Ettersalgsprosesser

- Optimalisering av ettersalgsprosesser



Behov for forklarbarhet av AI-inntreffer i ulike situasjoner

Dataanalytiker / Modellutvikler



- Hvordan kan algoritmen forbedres?
- Hvordan kan algoritmen feilsøkes?

Risikokontrollør



- Oppfyller den nødvendige forretningskrav?
- Kan produktet frigis for bruk?

Kunde / Bruker



- Hvordan kom modellen frem til resultatet?
- Hva kan gjøres for å påvirke avgjørelsen?

Regulator/ tilsynsmyndighet



- Er modellen validert?
- Oppfyller modellen regulatoriske krav?

Validering



- Er modellen robust og stabil?
- Diskriminerer modellen utsatte målgrupper?
- Er resultatene akseptable?

EUs AI Act viser hvordan regulering kan påvirke bruken av AI

Formålet med AI Act: Beskytte enkeltpersoner og samfunnet mot risikoer som oppstår ved bruk av kunstig intelligens. Loven innebærer et stort behov for grundige tiltak knyttet til høyrisiko-tilnærminger, spesielt innen styring, validering og etterlevelse.



Definisjon av AI



et *maskinbasert* system
som er designet for å operere med *varierende nivåer av autonomi* og
som kan vise *tilpasningsevne* etter implementering, og
som, for eksplisitte eller implisitte mål, *utleder* fra inndata det mottar hvordan det skal
generere *utdata som prediksjoner, innhold, anbefalinger eller beslutninger*
som kan *påvirke* fysiske eller virtuelle miljøer

Høyrisiko AI-systemer i finansiell sektor (Annex II)

Tilgang til og bruk av essensielle private tjenester og offentlige tjenester og ytelser vurderes som høyrisiko:

Kreditt:

"AI-systemer beregnet for å vurdere kredittverdighet hos enkeltpersoner eller fastsette deres kredittscore, med unntak av AI-systemer som brukes for å oppdage økonomisk svindel."

Forsikring:

"AI-systemer beregnet for risikovurdering og prissetting i forhold til enkeltpersoner ved livs- og helseforsikring."

Hvorfor er de høyrisiko?

"...tjenester og ytelser som er nødvendige for at folk skal kunne delta fullt ut i samfunnet eller forbedre sin levestandard."

"AI-systemer som brukes til disse formålene kan føre til diskriminering mellom personer eller grupper."

Men andre høyrisiko AI-områder kan også være relevante...

Ikke ansett som høyrisiko hvis AI-systemet:

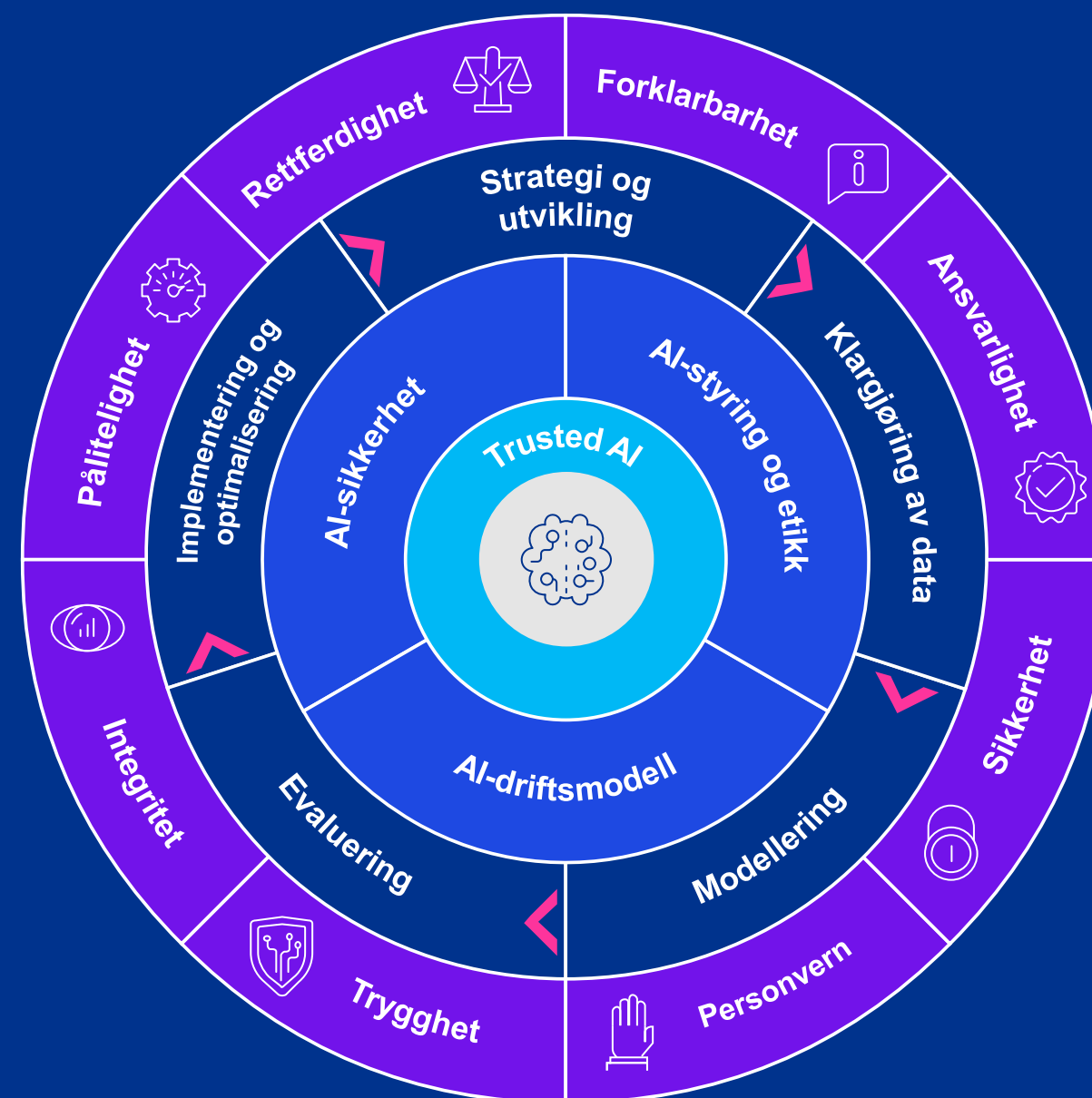
ikke utgjør en betydelig risiko for skade på helse, sikkerhet eller grunnleggende rettigheter til enkeltpersoner, inkludert ved ikke å påvirke utfallet av beslutningstaking vesentlig

Ingen betydelig risiko hvis AI-systemet er ment å:

- utføre en begrenset prosedyreoppgave,
- forbedre resultatet av en tidligere fullført menneskelig aktivitet,
- oppdage mønstre/avvik i beslutningstaking og ikke ment å erstatte eller påvirke tidligere fullført menneskelig vurdering, eller
- utføre en forberedende oppgave til en vurdering

Trusted AI

Vårt rammeverk Trusted AI bygger på åtte pilarer som hjelper selskaper imøtekomme regulatoriske krav og styre AI-relatert risiko.



1



Rettferdighet

Sikre at skjevheter i datasett blir redusert til virksomhetens akseptable nivå, slik at diskriminering og urettferdighet forhindres.

2



Forklarbarhet

Sikre at resultatene som AI produserer kan forklares, slik at brukerne kan opprettholde tillit, aktørene kan holdes ansvarlige, og løsningen forbedres.

3



Ansvarlighet

Sikre at mekanismer som fremmer ansvar er på plass, slik at negative konsekvenser med AI kan forutses, og rettslige hensyn ivaretas.

4



Sikkerhet

Beskytte mot uautorisert tilgang, lekkasjer og manipulasjon, slik at gjeldende prinsipper og regulatoriske krav til informasjonssikkerhet overholdes for AI.

5



Personvern

Benytte prinsippene for innebygget personvern, vurdere konsekvenser og andre krav til personvern slik at GDPR etterleves for AI.

6



Trygghet

Fastsette risikotoleranse og regulere tillatte bruksområder, slik at AI ikke får negativ innvirkning på mennesker, gjenstander eller miljøet.

7



Integritet

Sikre kvalitet i datahåndtering og berikelse gjennom hele livssyklusen, slik at AI produserer riktig beslutningsgrunnlag og overholder reguleringer.

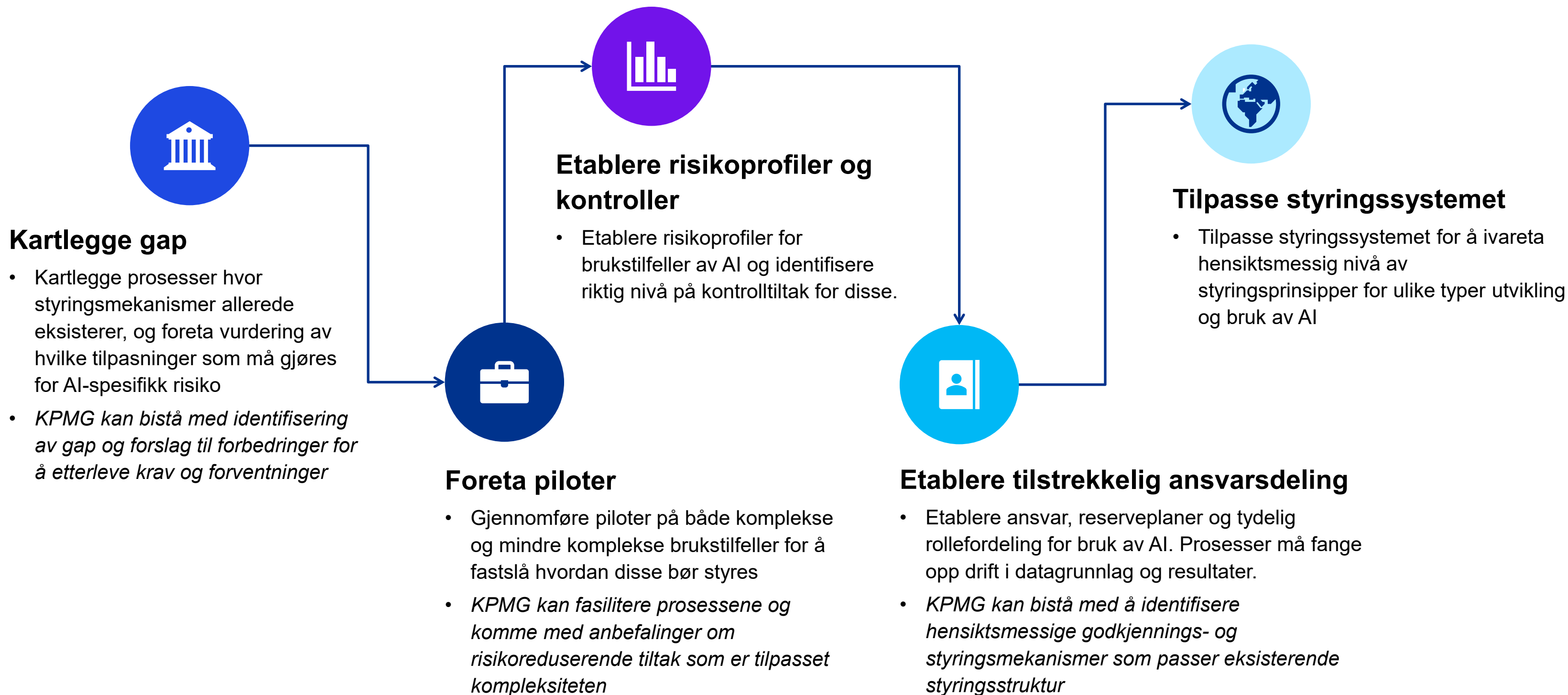
8



Pålitelighet

Sikre at innhold og prediksjoner fra AI er i tråd med tiltenkt formål og ytelse, slik at økonomisk- og sikkerhetsrisiko holdes innenfor terskelverdiene.

Hvordan sikre regulatorisk etterlevelse, samt styring og kontroll med AI-relatert risiko?



Skjevhet kan både introduseres og reduseres ved hvert trinn i implementeringen av en modell

Kilder til skjevhet

Data-innsamling



Data samles ofte inn av mennesker, som selv kan være påvirket av antagelser og fordommer. I ML-modeller kan vanligvis ikke alle data brukes. Å bestemme hvilke data som skal velges og neglisjeres er en kilde til skjevhet.

Gjennomsiktighet/ Forklarbarhet



Måten **modellen** blir utviklet på eller måten modellen blir trent på kan føre til urettferdige resultater.

Rettferdighet Etikk



Etter at modellen er trent, må resultatene **tolkes** for videre vurdering av rettferdighet og etikk.

Justering for mer rettferdighet

Forbehandling

Inndataene kan **endres direkte** på en måte som samsvarer med passende definisjon av rettferdighet, ved for eksempel å lage syntetiske data for underrepresenterte grupper.

Modellutvikling

Modellen bør oppfylle visse **rettferdighetskriterier**. Begrensninger som settes under design og utvikling, kan ta hensyn til både rettferdighet og nøyaktighet.

Etterbehandling

Resultatene kan **justeres direkte**, for eksempel ved å gjøre det lettere for enkelte utsatte målgrupper å få et «positivt modellresultat».

Behovet for forklarbarhet av AI-modeller kan vises ved «Husky vs Ulv»-modellen (1/2)



→ Bare én feil!

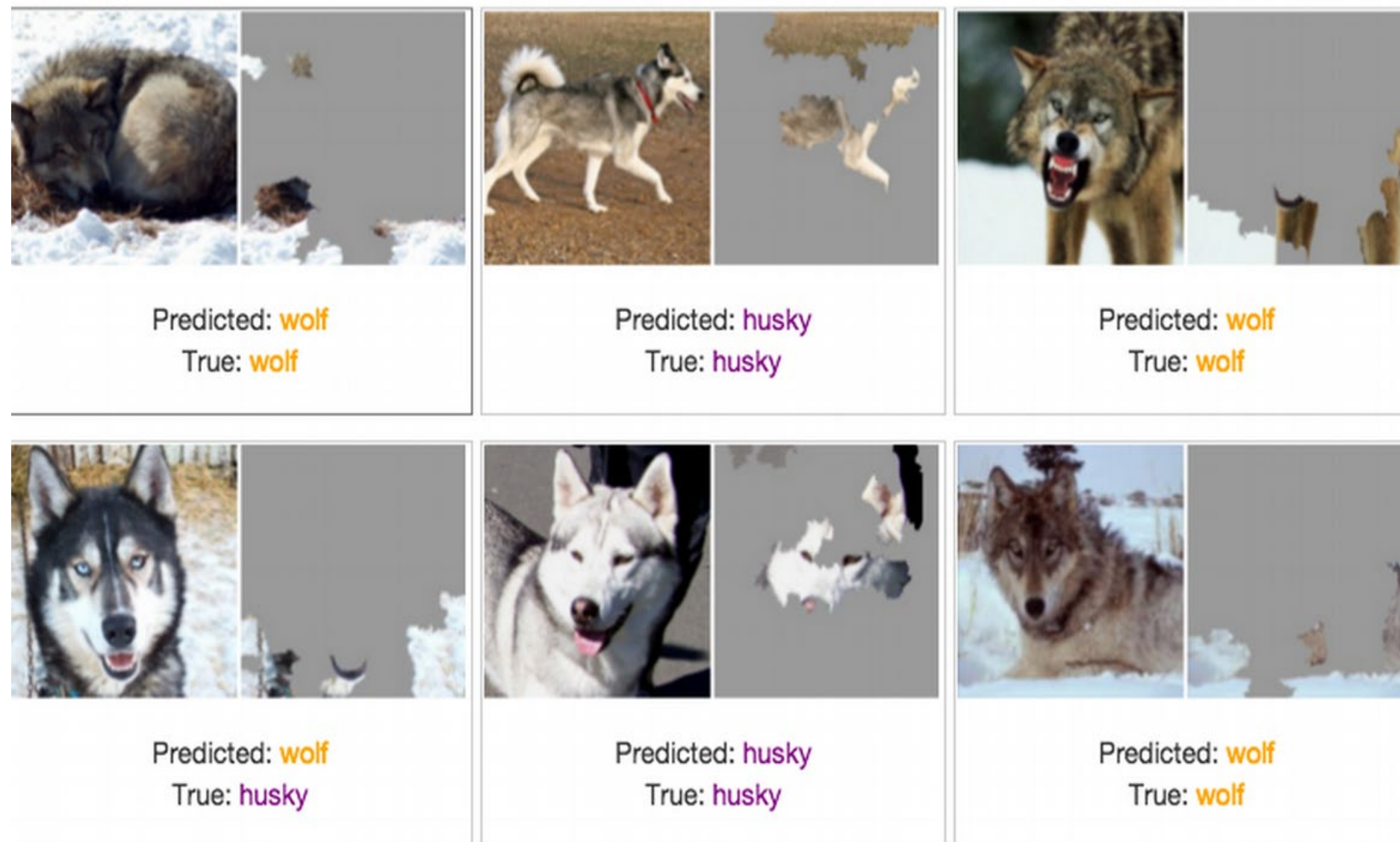
Klassifiseringen benyttet logistisk regresjon for å identifisere Husky vs Ulv.

Kan vi stole på denne modellen?

Hvorfor stoler vi på modellen?

Hvordan ser modellen forskjellen?

Behovet for forklarbarhet av AI-modeller kan vises ved «Husky vs Ulv»-modellen(2/2)



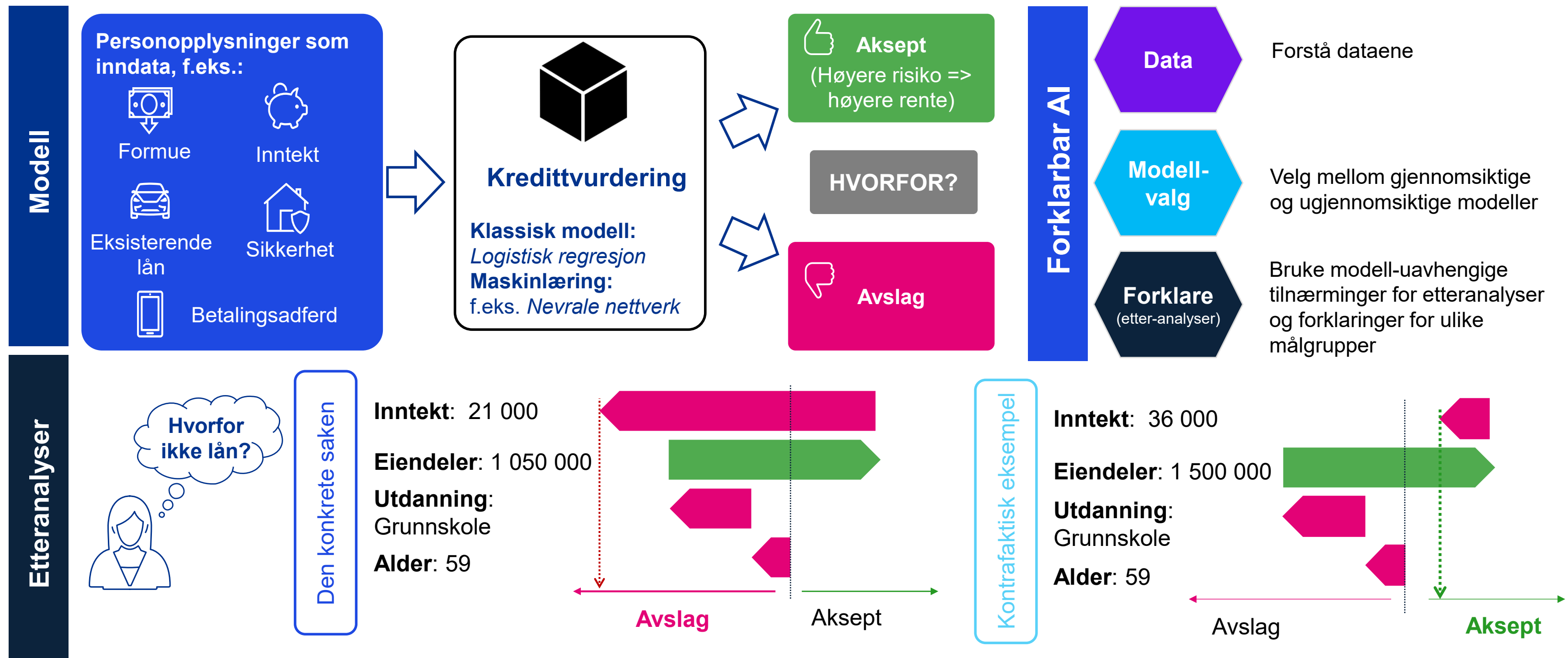
Vi kan benytte en tilnærming for å forklare algoritmen (her: LIME).

Visualiseringen viser pikslene som modellen benytter for å bestemme «Husky eller Ulv».

Maskinlæringsalgoritmen hadde i praksis lært seg å gjenkjenne snø!

Stoler du fortsatt på modellen?

Eksempel: Forklarbar AI for kredittmodell for personkunder



Oppsummering

AI åpner for mange nye **muligheter**, men også nye **utfordringer**.

Forklarbar og **rettferdig** bruk av AI krever innføring av nye prosesser og metoder.

Etablering av **rammeverk** for styring og kontroll bidrar til en **pålitelig bruk av AI**.



Kontakt oss



Kenneth Hansen
Head of Internal Audit

P: +47 907 37682

E: kenneth.hansen@kpmg.no

kpmg.no



kpmg.no/sosialemedier

The information contained herein is of a general nature and is not intended to address the circumstances of any particular individual or entity. Although we endeavour to provide accurate and timely information, there can be no guarantee that such information is accurate as of the date it is received or that it will continue to be accurate in the future. No one should act on such information without appropriate professional advice after a thorough examination of the particular situation.

©2025 Copyright owned by one or more of the KPMG International entities. KPMG International entities provide no services to clients. All rights reserved.

The KPMG name and logo are trademarks used under license by the independent member firms of the KPMG global organization.

Document Classification: KPMG Confidential