

Next Move

Et uafhængigt råd, et måltal der beskytter mennesker — og en opfølgning der tester rådet

AI-beslutningsstøtte · menneskelig beskyttelse · låste forudsigelser · dokumenteret opfølgning

Hvad Next Move er

Next Move er et **uafhængigt beslutningsprodukt** til organisationer, der skal vælge, om og hvordan et AI-initiativ bør gennemføres. Produktet omsætter en idé eller bekymring til et klart råd: gå i gang i det små, lav det om først, lad være, mål først eller gennemfør trinvis.

Rådet står ikke alene. Hver vurdering kobler den ønskede gevinst til et **beskyttende modmål**, sætter et loft over AI'ens autonomi, navngiver den forventede fejltilstand og låser en forudsigelse, der senere sammenlignes med organisationens egne resultater.

Portfolio-versionen viser den virkelige brugeroplevelse med **fiktive initiativer og tal**. Egne initiativer kan ikke indsendes, og der er ingen lagring eller produktionsforbindelse. Den offentlige produktidentitet hedder Cadence; Next Move er portfolioens produktnavn.



Landingssiden gør beslutningen konkret: beskriv udfordringen, få et uafhængigt råd, og følg det senere op.

Hovedkomponenter

Fra AI-idé til efterprøvet resultat. Hver komponent er beskrevet ved hvad den **gør**, med et konkret eksempel.

Disposition — Vurderingen ender i et klart handlingsudsagn — ikke et dashboard, der overlader fortolkningen til modtageren. *Eksempel: Svindel-scoring får dispositionen 'Lav det om først', fordi automatisk afvisning vil gøre modelbias til en kundeskadende beslutning.*

Uafhængig begrundelse — Rådet forklares med gevinst, menneskelig påvirkning, data- og styringsrisici samt den navngivne fejltilstand. *Eksempel: Et skævt træningssæt kan koncentrere flag på én demografi og skabe regulatoriske klager.*

Beskyttende måltal — Effektivitetsmålet får et modmål med baseline, aktuel værdi, rød streg og navngiven ejer. *Eksempel: Falsk-anklage-raten må ikke overstige 4 %, selv om initiativet reducerer svindellækage.*

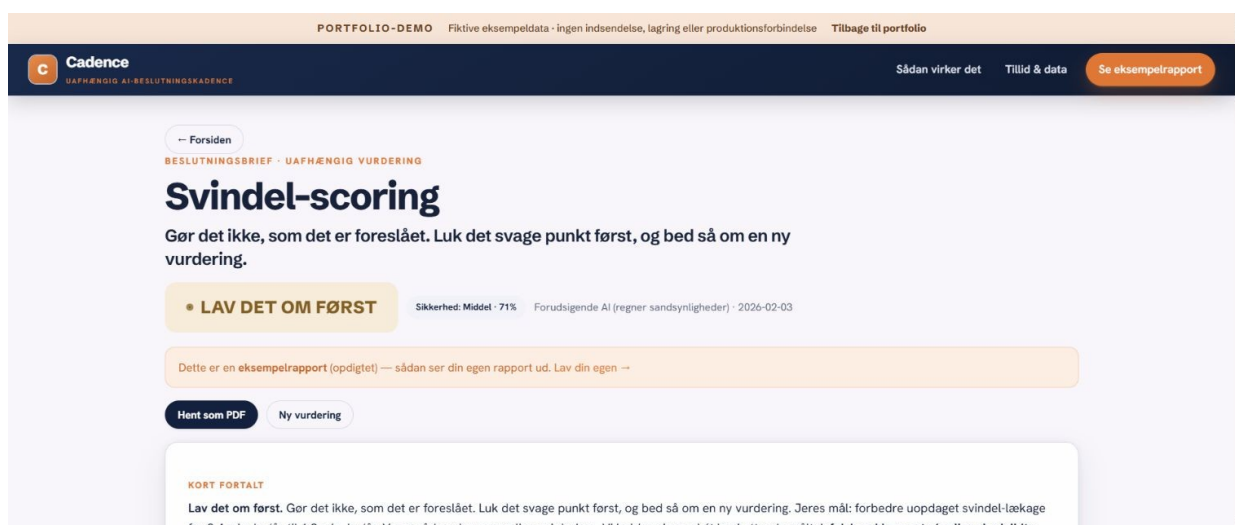
Autonomiloft — Produktet angiver præcis, hvad AI'en må gøre selv, og hvem der træffer den bindende beslutning. *Eksempel: Svindelmodellen må kun score; en menneskelig undersøger afgør hver sag.*

Readiness-gate — Roller, beslutningsmyndighed, oplæring og ret til at overstyre kontrolleres før ibrugtagning. *Eksempel: Et initiativ kan ikke få grønt lys, hvis medarbejderen ikke kan tilsidesætte modellens anbefaling.*

Låst forudsigelse — Forventet effekt, konfidens, målekilde, alternative forklaringer og opfølgingsdato fryses før udfaldet kendes. *Eksempel: Lækagen forventes reduceret fra 2,4 til 1,6 mio. kr., mens falsk-anklage-raten holdes under 4 %.*

Menneskepagt — Berørte medarbejdere får et beskyttende måltal, en tryk dissent-kanal, ret til override og synlig repræsentation. *Eksempel: Undersøgere kan bestride enhver score, og bestridte scores bliver input til forbedring.*

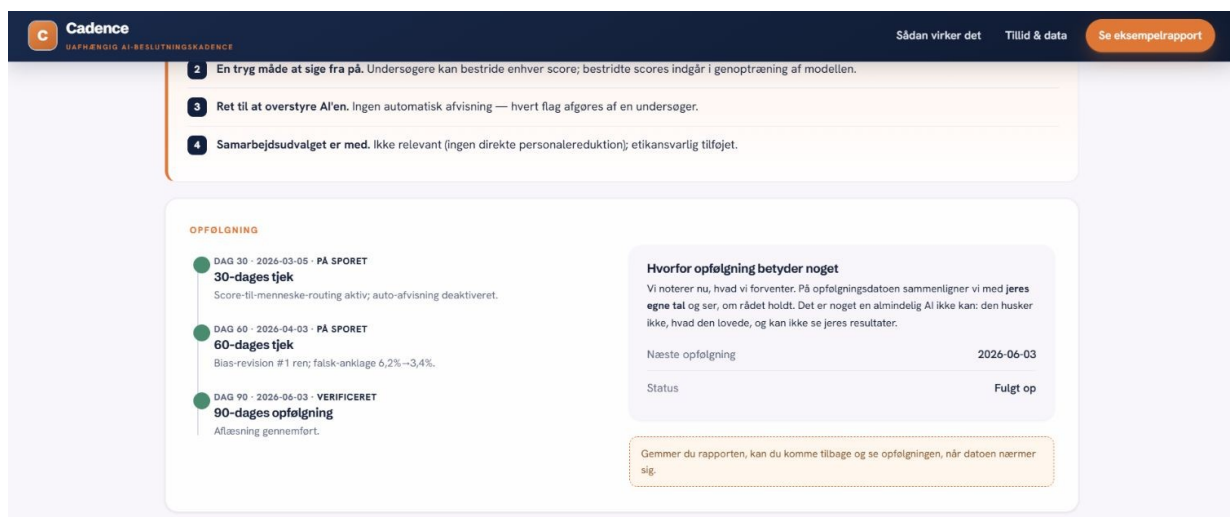
EU AI Act-bilag — Relevante styringspunkter kobles til vurderingen som input til organisationens dokumentation — uden at foregive certificering. *Eksempel: Eksempelrapporten peger på datastyring, menneskeligt tilsyn, nøjagtighed, bias og overvågning efter ibrugtagning.*



Eksempelrapporten åbner med én disposition, synlig sikkerhed og en kort, handlingsorienteret begrundelse.

Arbejdsgangen

- **1. Beskriv beslutningen.** Organisationen angiver initiativ, mål, berørte roller, AI-type, pres og kendte risici i almindeligt sprog.
- **2. Diagnosticér.** Produktet adskiller den påståede gevinst fra konsekvenserne for mennesker, ansvar, data og drift.
- **3. Sæt grænser.** Disposition, autonomiloft, readiness-krav og beskyttende måltal gør rådet operationelt.
- **4. Lås forventningen.** Forudsigelsen fryses med konfidens, målekilde, alternative forklaringer og dato, før resultatet kan pyntes.
- **5. Følg op.** Dag 30, 60 og 90 sammenligner de faktiske resultater med den oprindelige forudsigelse og registrerer, om rådet ramte.



Opfølgningsvisningen viser tre kontrolpunkter og forklarer, hvorfor et råd først bliver værdifuldt, når det kan testes mod virkeligheden.

Hvorfor produktet er anderledes

- **Ansvarlig dømmekraft frem for en mening.** Produktet tager stilling og viser betingelserne for, at beslutningen kan forsvares.
- **Mennesker er en del af målemodellen.** Produktivitet kan aldrig stå alene; skade, fairness, arbejdspress og ret til at sige fra instrumenteres samtidig.
- **Den kan troværdigt sige nej.** Uafhængighed af AI-leverandøren gør 'lad være' til et reelt muligt resultat.
- **Den husker, hvad den lovede.** En almindelig samtale-AI kan formulere et råd, men låser ikke en falsificerbar forventning og møder ikke senere resultatet.

Hvor jeg ser det gå hen

Produktets næste værdi ligger ikke i flere skærm billeder, men i mere dokumenteret læring:

- **Flere virkelige opfølgingsforløb.** Et outcome-ledger kan vise, hvilke dispositioner, grænser og modmål der faktisk beskytter mennesker og værdi.
- **Porteføljevisning for ledelsen.** Sammenlign initiativer efter disposition, readiness, autonome grænser og forestående opfølgninger.
- **Stærkere evidensgrundlag.** Den kuraterede casebase og failure library kan knyttes tættere til vurderingen med synlig proveniens.
- **Managed decision cadence.** Produktet kan fungere som en tilbagevendende, uafhængig kontrol frem for en engangsrapport.

Afslutning

Next Move gør AI-governance til en **gentagelig beslutningskadence** med et klart råd, en grænse for autonomi, et måltal der beskytter mennesker og en efterfølgende kontrol af virkeligheden. Produktets stærkeste idé er, at ansvarlig AI ikke bevises ved hensigter eller dashboards, men ved at skrive forventningen ned på forhånd og stå til regnskab for resultatet bagefter.