



Statistisk analyse af perkolatdata

Fase 2

DepoNet

Dato: 5. december 2025

Rev.nr.	Dato	Beskrivelse	Udarbejdet af	Kontrolleret af	Godkendt af
1.0	05/12/2025	Notat for fase 2 af projektet Statistisk analyse af perkolatdata	CEKO	ECHR	CEKO

Indhold

1.	Introduktion	4
2.	Sammenfattende konklusioner	4
3.	Fremgangsmåde	5
3.1.	Databehandling og dataanalyse	5
3.2.	Modellering	5
3.3.	Evaluering	5
4.	Modellernes datagrundlag	6
5.	Modellering og resultater	7
5.1.	Analyse 1: Unsupervised ML	7
5.2.	Analyse 2 og 3: Supervised ML	8
6.	Afrunding	8
7.	Hvad så nu?	9
8.	Bilag	10

1. Introduktion

Dette notat er udarbejdet af NIRAS og samler resultaterne fra fase 2 af projektet Statistisk Analyse af Perkolatdata. Fase 1 havde til formål at kortlægge tilgængelige data, opbygge en samlet database og gennemføre en indledende dataanalyse. Fase 2, som dette notat omhandler, bygger videre på grundlaget fra fase 1 med fokus på at udvikle og afprøve avancerede statistiske metoder. Formålet med fase 2 har været at undersøge, i hvilket omfang Machine Learning (ML) og kunstig intelligens (AI) kan danne grundlag for modellering af stofkoncentrationerne over tid.

Succeskriterierne for projektet har været:

- At opstille og afprøve ML/AI modeller, der kan beskrive udviklingen i stofkoncentrationer.
- At vurdere modellernes anvendelighed som grundlag for at estimere stofkoncentrationer over tid.
- At udarbejde en samlet vurdering af potentialet ved AI.

2. Sammenfattende konklusioner

NIRAS vurderer, at der er et potentiale i at anvende AI på perkolatdata, men også klare begrænsninger, som bør tages i betragtning.

Under modelleringen er testet både en Random Forest (RF) model og et neuralt netværk (NN). Random Forest modellen viser bedre resultater end det neurale netværk, da den i højere grad end netværket formår at fange variationen i data. Begge modeller er dog begrænsede og falder over tid tilbage på en glat og generel kurve. Det understreger, at kvaliteten af AI forudsigelserne afhænger direkte af datagrundlaget: En AI model bliver aldrig bedre end de data, den er trænet på. Det etablerede datagrundlag fra fase 1 dækker flere forskellige enheder og geografiske lokationer, men omfatter primært korte og ufuldstændige tidsserier. Modellerne har derfor aldrig set et fuldt udviklingsforløb og har begrænset mulighed for at forudsige langt ud i fremtiden. Dertil kommer, at de kemiske parametre ikke er jævnt repræsenteret i data. Kravene til analyser har ændret sig over tid, og centrale stoffer som PFAS indgår derfor ikke i datagrundlaget. Det betyder, at modellerne reelt kun kan anvendes på de parametre, der historisk set er blevet målt, og ikke nødvendigvis på de mest kritiske for eksempelvis efterbehandlingstiden. Dette er en væsentlig pointe i forhold til anvendeligheden af modellerne. Hvis målet er en langtidsforudsigelse af stofkoncentrationer, vil AI ikke give et mere pålideligt resultat end mere simple metoder som for eksempel en eksponentiel fremskrivning. Til kortere tidshorisonter på for eksempel et år kan AI derimod bidrage med et mere nuanceret estimat, fordi modellerne kan supplere den eksponentielle trend med den kortsigtede variation, som ses i data.

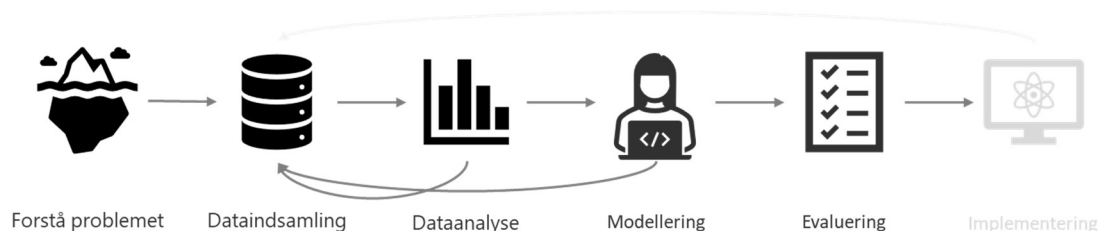
Desuden viser analysen, at AI kan anvendes til at identificere meningsfulde mønstre og grupper i data, der afspejler affaldstyper og andre relevante karakteristika. Dette indikerer, at data rummer strukturer, som kan udnyttes i fremtidige analyser. Eksempelvis kan nye eller datamæssigt svage enheder potentielt låne styrke fra lignende enheder, hvilket kan give et mere kvalificeret bud på koncentrationerne over tid. Derudover er der potentiale for at udvikle metoder til automatisk anomali detektion, som kan fungere som en indikator for uventede udviklingsforløb eller dataproblemer.

Samlet set vurderes det, at AI har potentiale som et supplerende beslutningsstøtteværktøj, der kan give et bedre overblik over variation, mønstre og datakvalitet på tværs af enheder. Der er dog klare begrænsninger ved brugen af AI på perkolatdata, og til langtidsprognoser er mere simple metoder et mere robust

valg, da udvikling og vedligehold af AI modeller ikke vil opveje den begrænsede forbedring i performance, der over tid alligevel flader ud.

3. Fremgangsmåde

Fase 2 har taget udgangspunkt i den datamodel og de analyser, der blev bygget i fase 1, men med et udvidet fokus på at udvikle og afprøve AI modeller til estimering af stofkoncentrationen i perkolat over tid. Arbejdet har været opdelt i tre hovedaktiviteter: *dataanalyse*, *modellering* og *evaluering*. Den overordnede tilgang fremgår af figuren nedenfor, hvoraf de første trin blev gennemført i fase 1.



3.1. Databehandling og dataanalyse

Det etablerede datagrundlag fra fase 1 blev gennemgået, kvalitetssikret og omstruktureret, så de kunne anvendes til modellering. Dette omfattede bl.a. ensretning af parameternavne, håndtering af forskelle i måleenheder, identificering og behandling af ekstreme værdier samt udvælgelse af pålidelige målepunkter. Derudover blev der tilføjet supplerende metadata, herunder tidspunkt for nedlukning og affaldstype, så data kunne kobles på tværs af enheder.

3.2. Modellering

Der blev udviklet modeller til at beskrive sammenhængen mellem tiden siden nedlukning og de målte stofkoncentrationer. Modelleringen blev gennemført trinvis med henblik på at klarlægge datamønstre, modelpotentiale og metodiske begrænsninger. Som første led blev der udført en klyngeanalyse efter en unsupervised tilgang for at vurdere, om tidsserierne udviser naturlige grupperinger baseret på mønstre i data uden forudgående kendskab til de underliggende forhold. Analysen giver en indledende vurdering af, om datagrundlaget rummer strukturer, der kan udnyttes i videre modellering. Efterfølgende blev der opstillet en supervised model baseret på Random Forest med det formål at undersøge, om modellen kan lære regler, som gør den i stand til at beskrive tidsseriernes udvikling på tværs af enheder og parametre. Dernæst blev der udviklet et neuralt netværk for at vurdere, om mere komplekse sammenhænge kan identificeres. Netværket kombinerer information på tværs af alle features og kan dermed opfange dybere mønstre, som potentielt kan forbedre forudsigelserne for enheder med varierende eller komplekse forløb.

For at muliggøre fremskrivninger blev der desuden implementeret en rekursiv forudsigelsesmetode inspireret af sprogmodeller. Dette indebærer, at modellens egne forudsigelser anvendes som input til efterfølgende forudsigelser.

3.3. Evaluering

Modellernes præstation vurderes ud fra tre opstillede scenarier, der repræsenterer forskellige datamængder og grader af kendskab til enhedernes historik. De opstillede scenarier bruges til at undersøge og vurdere, hvordan modellerne præsterer under realistiske forhold, samt i hvilken grad de er i stand til at generalisere til nye enheder. Scenarierne er opsummeret i Tabel 1.

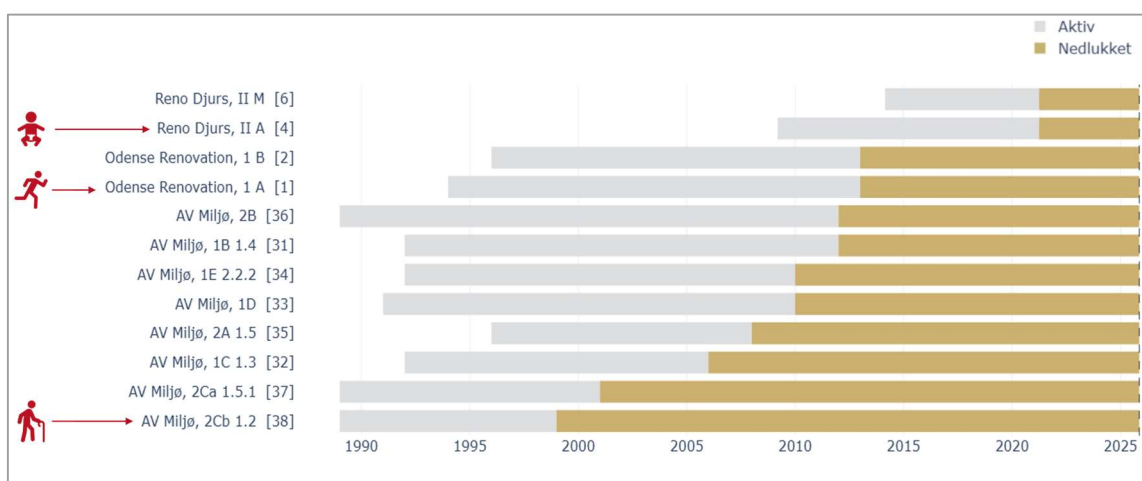
Scenarie	Beskrivelse
1. Den velkendte enhed	En velkendt enhed med lang tidsserie. Dette scenarie repræsenterer den bedste datatilstand for det etablerede datagrundlag og dermed det bedste datagrundlag for modellerne. Enheden har mange observationer over tid, og der findes flere lignende enheder i datagrundlaget. Det giver modellerne det stærkeste grundlag for at lære både den kortsigtede og langsigtede udvikling.
2. Den bekendte enhed	En delvist kendt enhed med begrænset historik. Dette scenarie repræsenterer en nyere enhed, med lavere repræsentation i datagrundlaget. Der findes målinger fra enheden, men tidsserien er kort. Samtidig er enhedens lokation eller karakteristika ikke stærkt repræsenteret i det samlede datagrundlag, hvilket giver modellerne færre sammenligningspunkter.
3. Den ukendte enhed	En ukendt enhed, som modellerne aldrig har set før. Dette scenarie simulerer en enhed, der står overfor nedlukning. Dette betyder, at modellen aldrig har set enheden før og dermed må basere sine forudsigelser på mønstre lært fra andre enheder.

Tabel 1: Scenarier som vurderingsgrundlag for modellernes præstation

På tværs af de tre scenarier vurderes modellerne med udgangspunkt i parameteren chrom. Set fra et data science perspektiv er dette en velegnet parameter, da den har en lang historik med mange målinger og optræder på tværs af mange enheder.

4. Modellernes datagrundlag

Efter behandling af det etablerede datagrundlag er 12 enheder identificeret og udvalgt til modellering og evaluering. Alle 12 enheder er afsluttede deponeringsenheder i efterbehandling. I figuren herunder er enhederne angivet med en illustration af deres aktive og nedlukkede perioder.



Figur 1: Modelgrundlaget med angivelse af scenarierne.

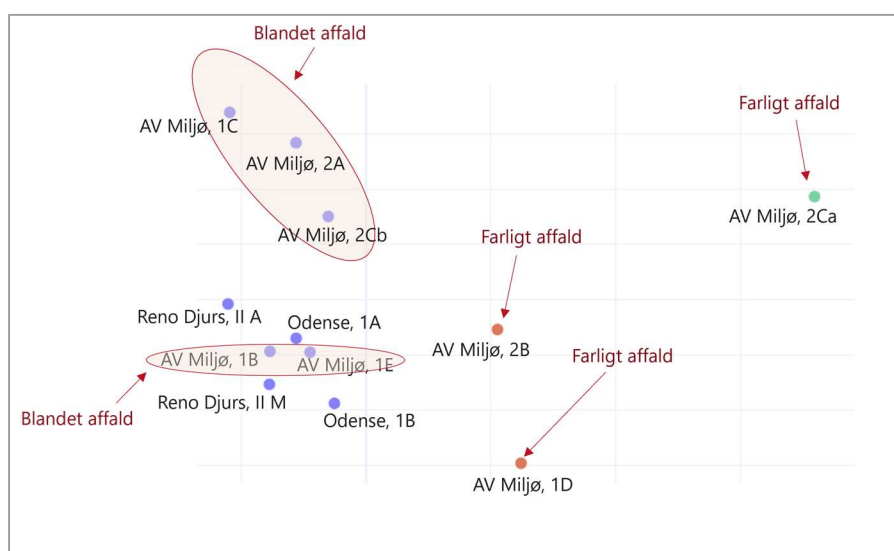
Oppefra ses scenarie 3 (Reno Djurs II A), scenarie 2 (Odense Renovation 1A) og scenarie 1 (AV Miljø 2Cb 1.2).

AV Miljø udgør den største del af datagrundlaget, både i antal enheder og i længden af de nedlukkede perioder, og en enhed herfra er derfor udvalgt som scenarie 1. Odense Renovation har et mere begrænset datagrundlag med to enheder og korte tidsserier for den nedlukkede periode, som dog er længere end hos Reno Djurs. En enhed fra Odense Renovation er derfor valgt som scenarie 2. Reno Djurs har ligeledes to enheder, men med væsentligt kortere nedlukningsperioder, og én af disse indgår derfor som scenarie 3. De to enheder fra Reno Djurs indgår ikke i træningsgrundlaget, da de anvendes til at afprøve modellernes evne til at håndtere enheder, der ikke tidligere er observeret. De øvrige ti enheder indgår i træningen sammen med alle relevante tidsserier. Modellerne er dermed trænet på et bredt og varieret datagrundlag, som afspejler forskelle i både enhedskarakteristika og parametre. Dette giver mulighed for at udnytte mønstre på tværs af enheder og parametre og er en væsentlig styrke ved anvendelsen af AI baserede metoder.

5. Modellering og resultater

5.1. Analyse 1: Unsupervised ML

I denne analyse har modellen fået et samlet datagrundlag bestående af alle de nedlukkede enheder. For hver enhed er der beregnet en række karakteristika baseret på deres tidsserier, så hver enhed indgår som ét datapunkt med tilhørende beskrivende karakteristika. Disse karakteristika omfatter blandt andet middelværdi, minimum, maksimum og varians for chrom, chlorid, nikkel, zink og øvrige relevante parametre. Resultatet af analysen fremgår af Figur 2. Analysen viser, at enhederne grupperer sig på måder, der i høj grad afspejler affaldstype og andre centrale forhold. Dette indikerer, at der er potentiale i at anvende AI til at identificere mønstre, som kan udnyttes i en beslutningsproces.



Figur 2: Resultat af unsupervised ML (cluster analyse)

Figuren viser tre identificerede grupper: blå, rød og grøn. Ved gennemgang af de tilhørende datapunkter ses det, at AV Miljø's enheder har haft særlig betydning for grupperingen. I den blå gruppe ligger AV Miljø's enheder med blandet affald, mens den røde og den grønne gruppe omfatter enheder med farligt affald. En nærmere analyse viser, at opdelingen mellem rød og grøn formentlig skyldes væsentlige forskelle i enhedernes volumen. Hverken affaldstype eller volumen er informationer, som modellen har fået som input. Grupperingen er således alene baseret på tidsseriernes forløb, hvilket understreger, at modellen

formår at identificere meningsfulde strukturer i data. Samlet set indikerer analysen, at der findes systematiske mønstre i datagrundlaget, som kan udnyttes af mere avancerede modeller.

5.2. Analyse 2 og 3: Supervised ML

I denne analyse har modellerne fået et samlet datagrundlag bestående af tidsserierne fra AV Miljø og Odense Renovation i Figur 3. Det betyder, at modellerne under træningen ser udviklingen i stofkoncentrationerne på tværs af både enheder og kemiske parametre. Formålet er at give modellerne mulighed for at identificere sammenhænge og mønstre på tværs af datagrundlaget. Resultaterne er opsummeret i Tabel 2, mens detaljerede resultater fremgår af bilagene.

Scenarie	RF-model	NN-Model
1. Den velkendte enhed	😊	😞
2. Den bekendte enhed	😞	😞
3. Den ukendte enhed	😞	😞

Tabel 2: Model resultater vurderet på hvert scenarie

Analysen viser, at AI generelt har udfordringer med at forudsige stofkoncentrationerne. Gennem flere iterationer blev det tydeligt, at modellerne havde svært ved at fange den overordnede trend, hvilket blandt andet resulterede i forudsigelser med stigende forløb eller negative koncentrationer for enkelte enheder og parametre. For at imødegå dette blev modellerne givet yderligere kontekst i form af en antagelse om, at udviklingen skulle være faldende og ikke måtte blive negativ. Dette svarer til en eksponentiel trend, som modellerne herefter skulle forudsige afvigelse omkring. Modellernes primære opgave bestod derfor i at modellere variation og udsving omkring denne trend, så de fremtidige forløb kunne blive mere robuste og afspejle den dynamik, der ses i træningsdata.

RF-modellen viste potentiale for den velkendte enhed, hvor den i nogen grad formåede at opfange de udsving, som kunne bruges til at forudsige fremtidige værdier. Ydeevnen faldt dog markant for enheder, som modellen kun havde begrænset kendskab til, og forløbene blev her mere glatte og gennemsnitlige.

Denne tendens sås allerede for NN-modellen ved den velkendte enhed. Neurale netværk kræver generelt store og varierede datamængder for at opnå høj præcision. Det kan sammenlignes med sprogmodeller, der bliver stærke, fordi de er trænet på meget omfattende tekstgrundlag bestående af hele sætninger, artikler og bøger. Derimod har den træned NN-model i denne analyse aldrig set komplette og lange tidsserier, og datagrundlaget er samlet set begrænset. Dette betyder, at modellen ikke har haft tilstrækkelig information til at lære den variation, der kendetegner udviklingen i stofkoncentrationerne.

6. Afrunding

Projektets resultater indikerer, at AI kan bidrage med indsigter i perkolatdata, særligt ved at samle information på tværs af celler og enheder og identificere mønstre, som ellers kan være vanskelige at opdage. Modellerne indikerer, at der er potentiale i at anvende AI til at beskrive den kortsigtede variation omkring den eksponentielle trend. Når der foreligger tilstrækkelig historik, kan modellerne give et mere detaljeret billede af de udsving, der præger udviklingen på kort sigt.

Resultaterne understreger dog, at deponier grundlæggende fungerer som en black box, hvor de underliggende processer ikke kan observeres eller kontrolleres direkte. Dette betyder, at modellerne udelukkende kan basere sig på de målte tidsserier og ikke har adgang til de mekanismer, der driver udviklingen. I praksis medfører det, at mere simple metoder som f.eks. en eksponentiel trend fortsat er det mest robuste grundlag for langtidsprognoser.

7. Hvad så nu?

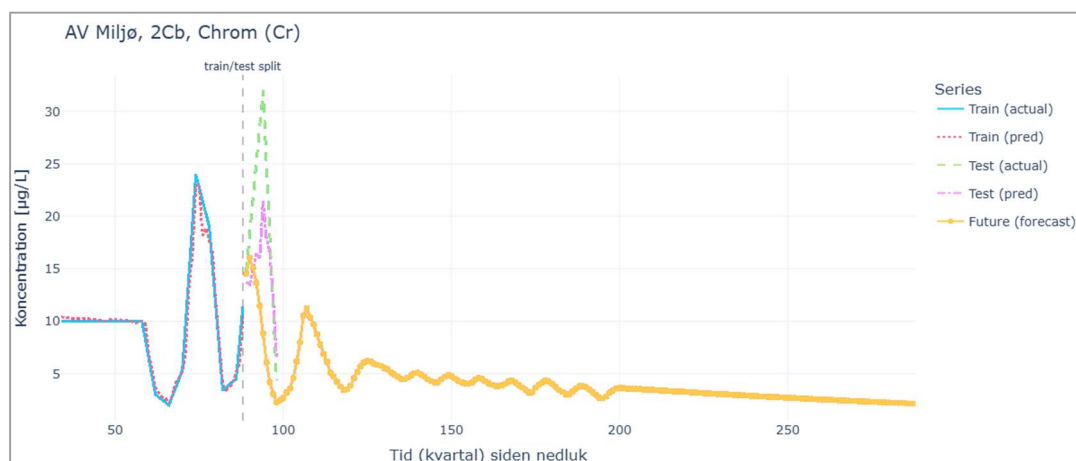
Gennem projektets faser har NIRAS fået indblik i de danske deponiers data og udfordringer. NIRAS har erfaret, at DepoNet har et klart ønske om at belyse og håndtere disse udfordringer samt en stor åbenhed over for anvendelsen af avancerede metoder. Projektet bekræfter, at AI rummer betydelige muligheder, men også at udbyttet i høj grad afhænger af datakvalitet og en ensartet tilgang til dataindsamling og datadisciplin. En systematisk og løbende indsamling af data på tværs af deponierne er derfor et centralt skridt fremad.

I projektets første fase etablerede NIRAS en database, som samler de tilgængelige data. Denne database kan med fordel videreudvikles, hvis DepoNet ønsker et fortsat fokus på at udnytte mulighederne i deres data. Med et mere robust og ensartet datagrundlag, struktureret i en fælles database, kan der udvikles et værktøj, f.eks. i Power BI, til løbende overvågning, kvalitetskontrol og tidlig identifikation af afvigelser. Under databehandlingen erfarede NIRAS, at der var betydelige afvigelser i data, som sandsynligvis stammede fra fejlbehæftede målinger. Sådanne afvigelser ville formentlig blive identificeret allerede ved modtagelse af analyseresultaterne, hvis man havde adgang til et sådant værktøj, og dermed ville det understøtte driften. Et fortsat fokus på datamængde og datakvalitet vil desuden kunne styrke de eksisterende, mere simple metoder, herunder den eksponentielle fremskrivning. Det betyder, at der også vil kunne hentes direkte værdi for de analysetilgange, der allerede anvendes i dag.

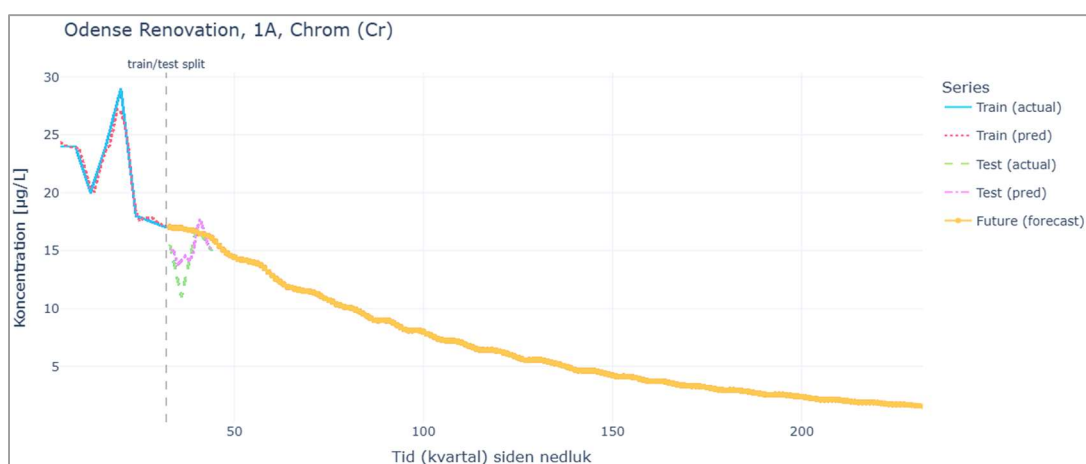
Med et stærkere datagrundlag på plads kan DepoNet begynde at afsøge konkrete anvendelsesscenarier, hvor AI kan skabe merværdi - eksempelvis ved klassifikation af nye enheder, automatisk detektion af afvigelser eller en dybere forståelse af variationen i de eksisterende tidsserier. Disse anvendelser vil kunne styrke en mere datadrevet tilgang til både overvågning og risikovurdering.

Samlet set kan AI tages i betragtning som et beslutningsunderstøttende supplement til de etablerede metoder. Teknologien kan styrke datakvaliteten og skabe bedre indsigt på tværs af enheder, forudsat at datagrundlaget fortsat udvikles og standardiseres. AI vil ikke kunne erstatte de eksisterende statistiske metoder og faglige vurderinger, men kan potentielt udvide og forbedre beslutningsgrundlaget.

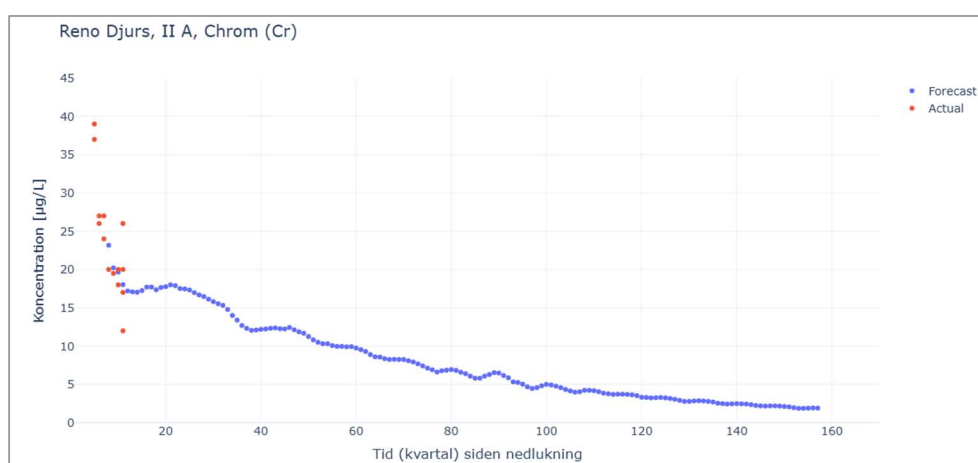
8. Bilag



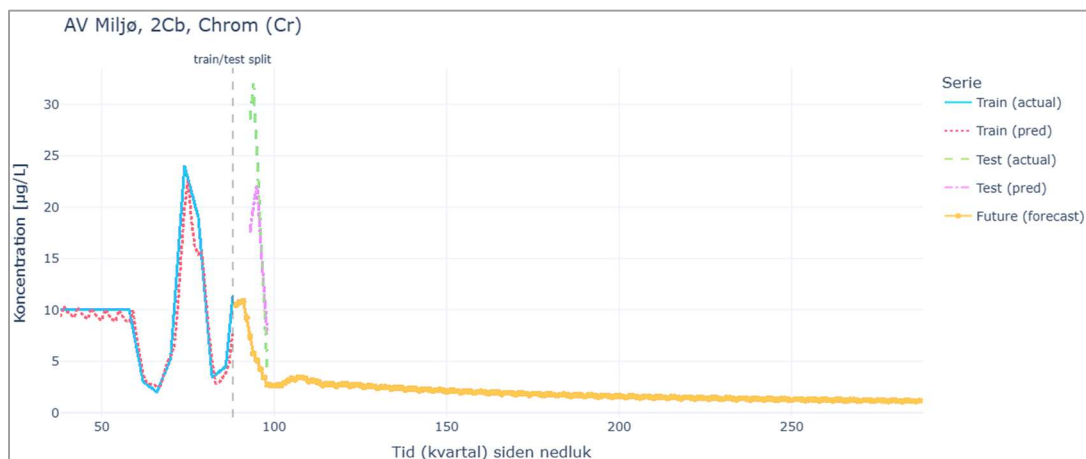
Figur 3: Resultat af RF-model på scenarie 1 (velkendt enhed)



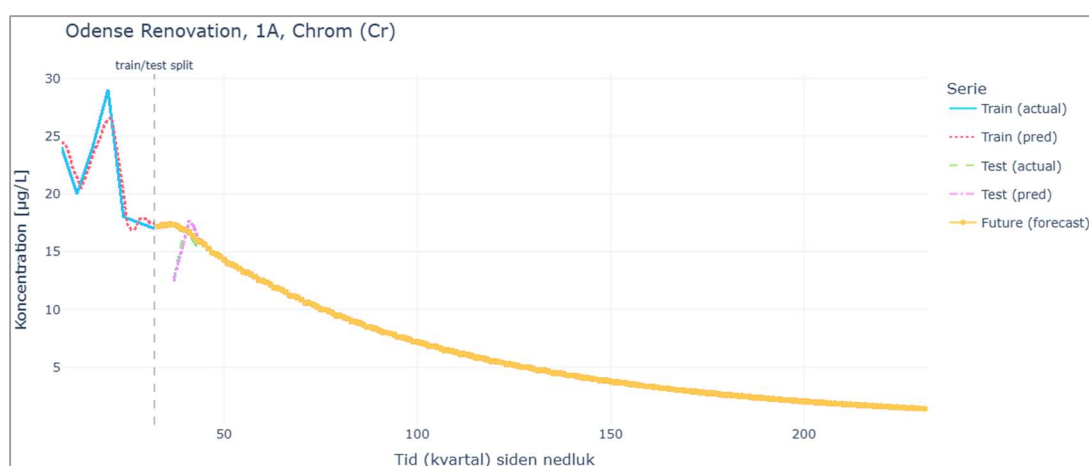
Figur 4: Resultat af RF-model på scenarie 2 (bekendt enhed)



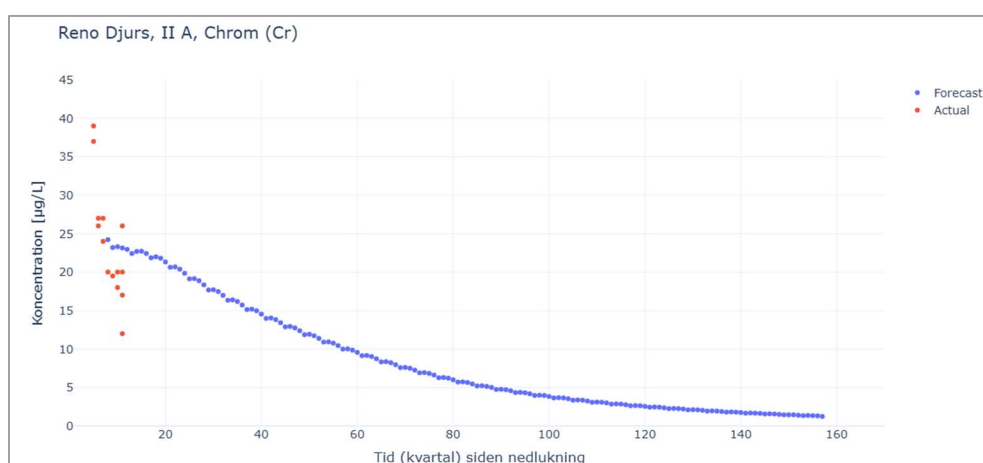
Figur 5: Resultat af RF-model på scenarie 3 (ukendt enhed)



Figur 6: Resultat af NN-model på scenarie 1 (velkendt enhed)



Figur 7: Resultat af NN-model på scenarie 2 (bekendt enhed)



Figur 8: Resultat af NN-model på scenarie 3 (ukendt enhed)