

Wanneer komt het boek terug?

De inzet van AI bij het voorspellen van het inlevermoment van uitgeleend bibliotheekmateriaal: bevindingen tijdens een pilot door bibliotheken, Probiblio en OCLC.

Auteurs: Ineke Goedhart, Paul Lucassen, Natalya Miltenburg,
Titia van der Werf, Rineke Zwanenburg



© 2023 OCLC.



Dit werk valt onder een Creative Commons Naamsvermelding 4.0 Internationaal-licentie (CC-BY 4.0).

Voor informatie over deze publicatie, neem contact op met: communications-nl@oclc.org

Bronvermelding:

Goedhart, Ineke; Lucassen, Paul; Miltenburg, Natalya; van der Werf, Titia; Zwanenburg, Rineke. 2023.

De inzet van AI bij het voorspellen van het inlevermoment van uitgeleend bibliotheekmateriaal: bevindingen tijdens een pilot door bibliotheken, Probiblio en OCLC. Leiden, OCLC.



Inhoudsopgave

Management Samenvatting	4
Inleiding	5
Pilot	7
Deelnemers	7
Onderzoeksvragen	7
Juridische aspecten	8
Juridische kaders	8
Privacy en bescherming persoonsgegevens	8
Wetgeving rondom AI	9
Verwerkingsovereenkomsten	10
Juridische toetsing	10
Verwerkingsgrondslag	10
Doelbeperking	11
Doelbinding en verenigbaar gebruik	11
Gegevensbeperking	11
Anonimiseren en pseudonimiseren	12
Geautomatiseerde besluitvorming	12
Toetsing aan het Europese AI-wetsvoorstel	13
Van pilot naar productiesysteem	13
Conclusie juridische aspecten	14
Ethische aspecten	15
Ethische kwesties	15
Ethische kaders	16
Ethische verkenning	18
Conclusie ethische aspecten	21
Technische aspecten	22
Relevante AI-begrippen	22
Machine Learning proces	23
Voorspellingsdoel en scope	23
Definitie van de dataset	24
Basisdataset	24
Uitbreiding met persoonsgegevens	25
Selecteren van het model	25
Eerste iteratie: Zonder persoonsgegevens	27
Tweede iteratie: Met persoonsgegevens, eerste verkenning	28
Derde iteratie: Met persoonsgegevens, tweede verkenning	29
Conclusie technisch aspecten	31
Bevindingen en aanbevelingen	32
Conclusie	32
Nieuwe inzichten	32
Aanbevelingen	33
Dankbetuigingen	34
Bijlage 1: Lijst met gebruikte afkortingen	35
Bijlage 2: Gebruikte data, data dictionary	36
Bijlage 3: Tellingen unieke combinaties en buckets voor CBS-data	38
Bijlage 4: Relevante functionaliteit en data beschikbaar in OCLC Wise	40

Management Samenvatting

Probiblio, OCLC en de openbare bibliotheken van Kennemerwaard, Bollenstreek en Krimpenerwaard, hebben onderzocht of en hoe kunstmatige intelligentie (artificial intelligence/AI) kan helpen met het nauwkeurig voorspellen van de retourdatum van uitgeleend bibliotheekmateriaal. Zij hebben in 2022 een AI-pilot uitgevoerd om de juridische, ethische en technische aspecten van deze casus te onderzoeken. De pilot fungeerde tevens als casus voor het bredere onderzoek naar de juridisch-ethische aspecten van AI in de cultuur-, media en erfgoedsector, in het kader van de werkgroep Cultuur en Media van de Nederlandse AI Coalitie (NLAIC).

Resultaat

De AI-experimenten op basis van bibliotheekmateriaal- en uitleentransactiegegevens, leverden een voorspellingsnauwkeurigheid op van +/-9 dagen ten opzichte van de nulmeting (+/- 14 dagen). Van belang is de bevinding dat de resultaten niet noemenswaardig beter werden door het toevoegen van (gepseudonimiseerde) persoonsgegevens van leners (+/- 7 dagen).

Proces

Tijdens de pilot vonden drie iteraties van AI-computermodelbouw en data-exploratie plaats, waarbij verschillende AI-methoden werden ingezet en de dataset uitgebreid. Bij elke iteratie werd opnieuw nagedacht over de gegevens, de methode en de juridisch-ethische consequenties.

Bevindingen juridische aspecten

Met het toevoegen van persoonsgegevens van de leners, tijdens de tweede en derde iteraties van de data-exploraties, lag de focus op het rechtmatig werken met de data. Privacy by design was het leidende principe en dit zorgde ervoor dat tijdens het traject, elke pilotdeelnemer – ongeacht zijn/haar rol binnen het team – doordrongen raakte van de AVG-beginselen, wat nodig was om naar de geest van de wet te kunnen handelen.

Bevindingen ethische aspecten

Tijdens de pilot was de verwerking van persoonsgegevens weliswaar rechtmatig, maar achteraf gezien, niet *noodzakelijk*. Dit was echter een bevinding op basis van de AI-experimenten, niet een conclusie op basis van ethische overwegingen. Een ethische verkenningssessie aan het einde van de pilot maakte duidelijk wat het gemis was van een ethicus aan boord van het team. Het dwong de pilotdeelnemers langer stil te blijven staan bij de probleemdefinitie, en aannames, keuzes en waarden beter te verhelderen. Zo werden de bibliotheekwaarden die in deze casus meespelen beter geëxpliciteerd: “de lener centraal”, “de betrouwbare bibliotheek” en “solidariteit”. Een interessante vraag hierbij was: hoe kunnen de uitleendienstenaansluitingen beter aansluiten bij deze waarden, die mee veranderen door de verplaatsing naar online?

Bevindingen technische aspecten

De gebruikte AI-methoden – lineaire regressie met mixed random effects en clustering – waren het meest geschikt voor de doeleinden van de pilot. Het bijeenbrengen van de uitleengegegevens van de verschillende deelnemende bibliotheken resulteerde in een dataset die voor het doel van de pilot onvoldoende consistent en compleet was. De data-exploraties duurden in totaal 6 maanden en in dit tijdsbestek werden te weinig werkbare patronen ontdekt. Er zijn volgens de pilotdeelnemers nog een aantal data-exploraties blijven liggen voor een eventueel vervolgonderzoek. Echter, tijdens de pilot is het inzicht gegroeid dat er waarschijnlijk beter resultaat geboekt kan worden door data-analyses te doen van processen of deelverzamelingen waar men meer grip op heeft.

Inleiding

De lenersvraag “Hoelang duurt het voordat ik mijn gereserveerde boek kan komen ophalen?” vormde de aanleiding van de pilot.



Figuur 1 - Illustratie van BiebPanel Infographic

Leners maken steeds meer gebruik van de reserveringsoptie voor boeken die niet direct leverbaar zijn. Het aantal reserveringen stijgt zowel in absolute zin als in relatie tot het totaal van de uitleningen. Het kan tot meer dan 6 weken duren voordat een gereserveerd boek beschikbaar is. Bijna 75% van de reserveringen wordt binnen 2 weken geleverd, 13% wordt tussen 2 à 6 weken geleverd en 12% wordt na 6 weken geleverd.¹

In de praktijk merken bibliotheekmedewerkers dat leners steeds meer behoefte hebben aan duidelijkheid over de levering van een gemaakte reservering. Wanneer kan een gereserveerd boek worden opgehaald?

Een BiebPanel-onderzoek uit 2017 naar het reserveren en verlengen van boeken bevestigt dat er onder de panelleden behoefte is aan een indicatie van de wachttijd.² Momenteel is het niet mogelijk hier een voorspelling over te doen. Een groot contrast met een webshop, waar klanten direct bij de aankoop de bezorgtijd gepresenteerd krijgen.

Waarom is het niet mogelijk inzicht te geven in de levertijd van gemaakte reserveringen? In het hele logistieke proces en de administratie rondom uitleenhandelingen is de retourdatum van uitgeleend materiaal de belangrijkste onbekende factor. De bibliotheek hanteert weliswaar een maximale uitleentermijn en soms ook een boeteregeling voor te laat inleveren, toch blijft het werkelijke inlevermoment onvoorspelbaar omdat deze afhankelijk is van het inlevergedrag van de lener. In hoeverre is dit gedrag voorspelbaar?

Kunstmatige intelligentie - ook Artificial Intelligence (AI) genoemd - wordt met name ingezet om voorspellingen te doen en kan mogelijk een uitkomst bieden om tot een nauwkeuriger verwachte inleverdatum te komen. Naast het beantwoorden van de lenersvraag kan het nauwkeurig voorspellen van de retourdatum van uitgeleend materiaal een (belangrijke) component zijn bij de inrichting van

¹ Probiblio Monitor Noord-Holland, Augustus 2022

(https://probiblio.nl/uploads/2022_bestanden/Monitor%20Collectie%20NH%202022%20Probiblio.pdf)

en Probiblio Monitor Zuid-Holland, November 2022

(https://probiblio.nl/uploads/2022_bestanden/Monitor%20Collectie%20Provincie%20Zuid-Holland%202022-Probiblio.pdf)

² Probiblio BiebPanel regulier onderzoek 2 2017 - rapport – maart 2018



uitleendiensten binnen het bibliotheekstelsel. Ook in andere deelgebieden (collectiemanagement, communicatie met leners) is zo'n voorspelling mogelijk relevant bij het verbeteren van de dienstverlening of optimaliseren van de collectie.

Het soort AI waar we op doelen heet machine learning (ML). ML richt zich op het bouwen van algoritmische modellen die patronen en relaties in gegevens kunnen identificeren. Een algoritmisch model wordt eerst getraind op een gegeven dataset en vervolgens gebruikt om bijvoorbeeld (nieuwe) data te classificeren of om voorspellingen te doen. In het technische deel van dit rapport gaan we dieper in op de gebruikte technologie en methode.

Het is nog verre van bewezen dat computermodellen geschikt zijn om menselijke gedragingen te voorspellen en het gebruik van AI op maatschappelijk vlak roept nog veel vragen op. De focus van het rapport ligt daarom op de juridisch-ethische aspecten van de pilot.

Hiermee beoogt het rapport ook bij te dragen aan het bredere onderzoek naar de juridisch-ethische aspecten van AI in de cultuur-, media en erfgoedsector, dat in het kader van de werkgroep Cultuur en Media van de Nederlandse AI Coalitie (NLAIC) uitgevoerd wordt³.

³ De use case van deze AI-pilot staat vermeld in het NLAIC WG-Cultuur en Media Position Paper 2022 (zie pagina 38). (https://nlaic.com/wp-content/uploads/2022/02/NL_NLAIC_position-paper-cultuur-en-media-def.pdf)

Pilot

Probiblio en OCLC hebben in 2022 samen met de bibliotheken Kennemerwaard, Bollenstreek en Krimpenerwaard een pilot uitgevoerd om te onderzoeken of en hoe AI kan helpen bij het voorspellen van het inlevermoment van uitgeleend materiaal. OCLC is de ontwikkelaar van Wise, een bibliotheek-automatiseringssysteem in gebruik bij de meeste openbare bibliotheken in Nederland⁴. Probiblio is een Provinciale ondersteuningsinstelling (POI) voor openbare bibliotheken in Noord- en Zuid-Holland en verzorgt het functioneel beheer van het Wise-systeem voor aangesloten bibliotheken. Niet eerder hebben deze partijen AI ingezet op data uit het Wise-systeem om een functionaliteit te ontwikkelen of verbeteren.

Deelnemers

Bibliotheken

De bibliotheekdeelnemers waren:

- Lenette Logtenberg – Programmaregisseur Ontwikkeling en Ontplooiing van de Bibliotheek Kennemerwaard;
- Lydia Verbart – Domeinspecialist Collectie van de Bibliotheek Bollenstreek;
- Marit van den Dool – ICT-coördinator van Bibliotheek Krimpenerwaard.

Zij vertegenwoordigden de pilotbibliotheken en fungeerden als opdrachtgever en klankbordgroep gedurende de hele pilot.

Probiblio

Probiblio was initiatiefnemer, projectleider, had een adviserende rol en leverde juridische ondersteuning. Betrokken waren:

- Rineke Zwanenburg – Privacy Officer;
- Ineke Goedhart – Netwerkadviser Digitale Infrastructuur.

OCLC

OCLC was de technologiepartner en leverde ondersteuning op het juridische vlak. OCLC ontwikkelde het algoritme op basis van de door de

bibliotheek bepaalde gegevens. Betrokken waren:

- Titia van der Werf – Adviseur Research;
- Natalya Miltenburg – Legal Council;
- Paul Lucassen – Architect Wise.

Onderzoeksvragen

Gedurende de pilot zijn de volgende onderzoeksvragen aan de orde geweest:

Hoofdonderzoeksvraag: In hoeverre is het mogelijk het inlevermoment van uitgeleend materiaal nauwkeurig te voorspellen met behulp van AI?

Deelvragen:

- Welke (historische) gegevens over uitleningen zijn hiervoor beschikbaar en relevant?
- Zijn voorspellingen op basis van niet-persoonsgebonden gegevens nauwkeurig genoeg?
- Welke persoonsgegevens zijn nodig voor (meer) betrouwbare en nauwkeurige voorspellingen?
- Aan welke juridisch-ethische voorwaarden moet de inzet van AI hierbij voldoen?

De directies van de drie pilotbibliotheken hadden als belangrijke voorwaarde aangegeven dat zij akkoord gingen om persoonsgegevens van hun lenersbestand voor de pilot in te zetten “mits het binnen de wettelijke regels past en het gebruik van deze gegevens in verhouding staat tot het resultaat”. Dit criterium had de vervlechting van de juridisch-ethische aspecten van datagebruik voor AI niet beter kunnen formuleren.

Het rapport is onderverdeeld in drie hoofdstukken die achtereenvolgens de juridische, ethische en technische aspecten van de pilot behandelen.

⁴ OCLC Wise is een bibliotheekautomatiseringssysteem dat bij (een meerderheid van Nederlandse) openbare bibliotheken gebruikt wordt voor ondersteuning van de bedrijfsprocessen: o.a. collectiebeheer, circulatie (uitleening) en lenersadministratie, zie Bijlage 4 voor een uitgebreidere beschrijving van de relevante functionaliteit.

Juridische aspecten

Voordat we, als pilotdeelnemers, aan de slag konden gaan met de ontwikkeling van AI, moesten we eerst bepalen welke gegevens we zouden gaan gebruiken en deze toetsen aan de vigerende juridische kaders. In dit hoofdstuk gaan we in op de juridische kaders en de juridische toetsing.

Welke gegevens zijn beschikbaar en relevant voor het voorspellen van het inlevermoment van uitgeleend materiaal? Te denken valt aan zowel titel- en exemplaargegevens van het materiaal, vestigingsgegevens van de deelnemende bibliotheken, alsook persoonsgegevens van de leners. Tot deze laatste categorie behoren gegevens zoals: leeftijd, geboortjaar, geslacht, postcode, leenhistorie en financiële historie (bijvoorbeeld opgelegde boetes, die inzicht geven of een lener een boek te laat heeft ingeleverd). Het is met name de verwerking van deze persoonsgegevens die toetsing behoeft. Ook de aard van de AI-verwerkingen op de dataset hebben we juridisch getoetst.

Juridische kaders

De op deze pilot van toepassing zijnde juridische kaders vallen uiteen in drie categorieën: privacy en bescherming persoonsgegevens, wetgeving rondom AI en verwerkingsovereenkomsten tussen de pilot-partners.

Privacy en bescherming persoonsgegevens

Privacy is een recht dat een individu heeft tegen inmenging in zijn privéleven, een recht om vrij te zijn van bespieding en beïnvloeding. Privacy gaat om de afscherming van de woning of leefruimte, familie- en gezinsleven, het eigen lichaam, maar ook het recht om vertrouwelijk te communiceren. Het is een universeel mensenrecht, een fundamentele vrijheid en een grondrecht⁵.

Meer specifiek omvat privacy ook bescherming van iemands persoonsgegevens. Een persoon heeft het recht om enige controle te hebben over hoe de eigen persoonlijke informatie wordt verzameld en gebruikt. Bescherming van persoonsgegevens wordt doorgaans geassocieerd met het juiste gebruik van persoonlijke gegevens oftewel Persoonlijk Identificeerbare Informatie, zoals namen, adressen, burgerservicenummers en creditcardnummers. Het begrip persoonsgegevens strekt zich echter ook uit tot andere waardevolle of vertrouwelijke gegevens, zoals financiële gegevens of gezondheidsinformatie. Van belang bij de verwerking van persoonsgegevens is de vraag in hoeverre deze direct dan wel in combinatie met andere gegevens tot een persoon herleidbaar zijn. Het enkele gegeven dat een bibliotheekbezoeker 35 jaar oud is, is in beginsel onvoldoende om een persoon te herleiden. Daarentegen kan een persoon op basis van informatie bestaande uit de combinatie van een geboortjaar en de postcode sneller herleidbaar zijn (zeker in dunnerbevolkte postcodegebieden). Het gebruik van (de combinatie van) gegevens leidt ertoe dat er sprake is van een verwerking van persoonsgegevens.

De bescherming van persoonsgegevens vormt een veel grotere uitdaging met de steeds verdergaande digitalisering van de samenleving, waar deze gegevens in machineleesbare vorm op grote schaal worden geregistreerd, bewaard, uitgewisseld en verwerkt. Zowel wetgevers als verschillende industrieën hebben initiatieven genomen om de regels rondom bescherming van persoonsgegevens vast te leggen om juist gebruik van persoonsgegevens te waarborgen. De belangrijkste regels voor de omgang met persoonsgegevens in Nederland zijn vastgelegd in de Algemene verordening gegevensbescherming (AVG). De AVG gaat over het rechtmatig omgaan met persoonsgegevens.

⁵ Zie: <https://nl.wikipedia.org/wiki/Privacy>

De belangrijkste basisbeginselen uit de AVG zijn als volgt samen te vatten⁶:

- **Transparantie**

De persoon van wie de gegevens verwerkt worden, weet dat zijn gegevens bewaard worden. Hij heeft hier toestemming voor gegeven en kent zijn rechten in dezen.

- **Doelbeperking**

Persoonsgegevens worden enkel en alleen voor het vooraf bepaalde en inzichtelijke doel gebruikt.

- **Gegevensbeperking**

Alleen de persoonsgegevens die noodzakelijk zijn voor de bestemming worden verzameld en bewaard.

- **Juistheid**

De persoonsgegevens moeten kloppend zijn.

- **Bewaarbeperking**

Persoonsgegevens mogen niet langer bewaard worden dan nodig voor het beoogde doel.

- **Integriteit en vertrouwelijkheid**

Alle persoonsgegevens moeten veilig bewaard worden, zodat ze niet toegankelijk zijn voor onbevoegden en niet verloren gaan of vernietigd kunnen worden.

Wetgeving rondom AI

Voor het gebruik van AI zijn diverse wetteksten van belang. Wetten en documenten waar we naar gekeken hebben, zijn:

De Algemene verordening gegevensbescherming (AVG)

Voor de pilot is met name Artikel 22 van toepassing, waarin onder andere staat dat een betrokkene zich moet kunnen onttrekken aan geautomatiseerde besluitvorming en profilering. Ook mag een dergelijke verwerking alleen met uitdrukkelijke toestemming van de betrokkene worden uitgevoerd.

Het Voorstel voor Europese Harmonisatie van Kunstmatige Intelligentie

In 2021 heeft de Europese Commissie een voorstel voor de wet Artificial Intelligence gepubliceerd. Het idee erachter is dat “Europeanen kunnen genieten van nieuwe technologieën, ontwikkeld en functionerend volgens de waarden, grondrechten en beginselen van de Unie” (uit Art.1.1 van het voorstel).

Het voorstel is ‘risicogestuurd’: het verwachte risico van de AI bepaalt hoe zwaar de regels zijn. Hoe hoger het risico, hoe zwaarder de ingrepen van de wetgever. Het minst gereguleerd zijn de systemen met een laag risico. Meer specifiek, kent de AI-verordening vier risicoklassen:

1. **Onaanvaardbaar risico** – dit type AI

vormt zulke grote risico’s voor de grondrechten en de samenleving dat ze onaanvaardbaar zijn zoals algoritmische sociaalkreditsystemen, die burgers een bepaalde score toekennen op basis van gedrag.

Deze AI-systemen zijn verboden binnen de Europese Unie.

2. **Hoog risico** – AI-systemen met hoog risico voor de gezondheid en veiligheid of de grondrechten van natuurlijke personen. De AI-verordening vereist concrete verplichtingen met betrekking tot de kwaliteit van gebruikte datasets, technische documentatie en administratie, transparantie en informatievoorziening aan gebruikers, menselijk toezicht en robuustheid, nauwkeurigheid en cybersecurity.

3. **Beperkt risico** – AI-systemen met concrete beheersbare risico’s.

De AI-verordening voorziet in transparantieplichtingen voor chatbots, deep fakes, biometrische categoriseringssystemen voor emotieherkenning

4. **Laag risico** – gebruik of toepassingsgebieden van AI met een laag risico is niet gereguleerd, wel worden vrijwillige initiatieven en gedragscodes in de AI-verordening aangemoedigd.

De risicoklasse bepaalt de mate van regulering onder de AI-verordening.

⁶ Dit zijn de in artikel 5 AVG opgenomen basisbeginselen. Zie ook:

<https://autoriteitpersoonsgegevens.nl/nl/over-privacy/wetten/algemene-verordening-gegevensbescherming-avg>

Verwerkingsovereenkomsten

Bibliotheken verzamelen persoonsgegevens die zij nodig hebben voor het uitvoeren van hun kerntaken, waaronder het verzorgen van uitleningen. Deze gegevens verkrijgen zij met toestemming van de lener. Toestemming is de grondslag waarop deze gegevens verwerkt mogen worden. Deze gegevens vormen tevens onderdeel van de verwerkingsovereenkomsten die bibliotheken sluiten met derde partijen, zoals Probiblio, voor het uitvoeren van hun kerntaken.

Voordat er met persoonsgegevens gewerkt kan worden, moeten de partijen die betrokken zijn bij de verwerking van deze gegevens, contractuele afspraken met elkaar maken over de wijze waarop zij invulling geven aan hun verwerkingsverantwoordelijkheid, conform de AVG. In dergelijke contracten zijn onder meer de aard en het doel van de verwerking, het soort persoonsgegevens, en de rollen van de verschillende partijen omschreven. De afspraken hebben betrekking op zaken zoals:

- het informeren van betrokkenen over de verwerking van hun persoonsgegevens, conform het transparantiebeginsel (meestal in de vorm van een privacyverklaring);
- het wissen/vernietigen van de persoonsgegevens, conform het beginsel van gegevensbeperking;
- de beveiliging van de gegevens, conform het beginsel integriteit en vertrouwelijkheid.

Er zijn bestaande verwerkingsovereenkomsten tussen de bibliotheken, Probiblio en OCLC over de doeleinden en middelen voor de verwerking van persoonsgegevens. Daarin zijn de verwerkingsrollen als volgt verdeeld: de bibliotheek is verwerkingsverantwoordelijke, Probiblio verwerker en OCLC subverwerker. Deze rolverdeling brengt met zich mee dat elk van de partijen eigen verantwoordelijkheden en verplichtingen heeft.

De bibliotheken hebben een rol van Verwerkingsverantwoordelijke, volgens de AVG “een natuurlijke persoon of rechtspersoon die, alleen of samen met anderen, het doel van en de middelen voor de verwerking van persoonsgegevens vaststelt”. Probiblio fungeert daarbij als Verwerker, oftewel “een rechtspersoon die ten behoeve van de verwerkingsverantwoordelijke persoonsgegevens verwerkt”.

Omdat Probiblio OCLC inschakelt voor het leveren van het Wise-systeem, is OCLC in deze contractrelatie een Subverwerker.

Juridische toetsing

Met een juridische toetsing vooraf wilden wij bepalen welke persoonsgegevens we voor de pilot konden gaan gebruiken en onder welke voorwaarden. Dit was nodig om binnen de juridische kaders voor datagebruik te werken en voor de pilotbibliotheken om toestemming te kunnen geven voor gebruik van deze gegevens in de pilot.

De belangrijkste vragen die daarbij aan de orde kwamen hadden betrekking op de verwerkingsgrondslag en doelstelling, de toetsing aan de AVG-basisbeginselen, de mogelijkheden en beperkingen van anonimiseren dan wel pseudonimiseren, en tenslotte, de vraag of er bij de pilot sprake was van geautomatiseerde besluitvorming.

Verwerkingsgrondslag

In de eerste gesprekken die we hielden kwam de wens naar voren om hoog in te zetten als het gaat om de hoeveelheid persoonsgegevens die in de pilot gebruikt zouden kunnen worden. We wilden gezien het korte tijdsbestek van de pilot geen nieuwe persoonsgegevens verzamelen. Wel wilden we bestaande lenersgegevens waar bibliotheken over beschikken zoveel mogelijk kunnen hergebruiken. We wilden daarmee de kans op een betere voorstelling vergroten en de grenzen van de juridische en ethische kaders onderzoeken.

Een belangrijke eerste keuze was dus om te werken met de bestaande persoonsgegevens die onderdeel uitmaken van de bestaande verwerkingsovereenkomsten tussen de pilotpartijen. Dit betekende tevens dat de verwerkingsgrondslag van deze gegevens in het kader van de pilot gebaseerd was op reeds eerder verleende toestemming van leners.

Een gevolg hiervan was dat de pilot rekening moest houden met de verwerkingsdoelstelling waarvoor toestemming gegeven is.

Doelbeperking

Het werd ons al snel duidelijk dat vanuit juridisch oogpunt⁷, de eisen aan de bescherming van persoonsgegevens niet anders zijn wanneer deze gegevens ingezet worden voor onderzoeksdoeleinden dan wanneer ze voor een operationele dienst worden gebruikt. In beide gevallen gaat het om “verwerking van persoonsgegevens” en moet er dus aan alle eisen uit de AVG en andere regelgeving worden voldaan.

Juridisch gezien is het niet mogelijk om onder de noemer “onderzoek” of “pilot” ongelimiteerd met persoonsgegevens te experimenteren en deze pas te toetsen aan de AVG als de gegevens in een operationele context worden gebruikt.

Doelbeperking is een belangrijk begrip in de AVG. Dit dwong ons tot nadenken over het doel van het gebruik van persoonsgegevens voor de pilot en of dit doel afwijkt van het doel waarvoor de bibliotheken deze gegevens verwerken, zoals vastgelegd in hun verwerkingsovereenkomsten. We kwamen al snel tot de conclusie dat de pilotdoelstelling gericht was op de verbetering van de informatievoorziening aan de lener met betrekking tot reserveringen, en dus in het verlengde lag van de verwerkingsdoelstelling voor uitleningen van bibliotheken.

Doelbinding en verenigbaar gebruik

We leerden dat er in de wetgeving ook uitzonderingen gemaakt zijn voor de verdere verwerking met het oog op archivering in het algemeen belang, wetenschappelijk of historisch onderzoek of statistische doeleinden. Deze verdere verwerkingsdoeleinden worden overeenkomstig de AVG

niet als onverenigbaar met de oorspronkelijke doeleinden beschouwd. Dit is waar de begrippen “doelbinding” en “verenigbaar gebruik” om de hoek komen kijken. Voor de pilot was het dus bovendien mogelijk om de bestaande persoonsgegevens te gebruiken voor statistische doeleinden.

Mag je dan zomaar alles onder de noemer doelbeperking of statistisch onderzoek? Nee, we kwamen erachter dat we alsnog te maken hadden met het AVG-beginsel van gegevensbeperking en de AVG-hoofregel van anonimisering.

Gegevensbeperking

De AVG gaat uit van het principe van privacy by design, wat inhoudt dat er al bij de ontwikkeling van producten en diensten aandacht moet zijn voor privacy – met andere woorden, dat er vanaf het begin nagedacht moet worden over de persoonsgegevens die nodig zijn voor het doel van de verwerking. Gegevensbeperking is een AVG-beginsel dat zeker bij AI-toepassingen strikt gehanteerd moet worden. AI is immers het meest effectief als het ingezet kan worden op grote hoeveelheden data, zoals klantgegevens, bij het doen van aankoopvoorspellingen bijvoorbeeld. Ongelimiteerde dataverzameling kan echter ook schade toebrengen aan de privacy van de betrokkenen. De AVG is juist bedoeld om hieraan paal en perk te stellen. Dataminimalisatie werd daarmee een belangrijk onderdeel van de pilot: wat zijn de grenzen van het gebruik van persoonsgegevens? Wat is echt nodig om tot een goed resultaat te komen, wat niet?

⁷ Om de juridische grenzen en randvoorwaarden in kaart te helpen brengen is advies gevraagd aan privacyspecialisten van een advocatenkantoor.

Anonimiseren en pseudonimiseren

Met het gebruik van anonieme data speelt de AVG geen rol. Bij het anonimiseren van persoonsgegevens, kan er geen link meer worden gemaakt met de persoon en is er dus geen sprake meer van persoonsgegevens. Bijvoorbeeld door het pseudo-ID – zoals het lenerspasnummer – te verwijderen én de gegevens niet of in zeer beperkte mate te combineren. Dan zijn de gegevens niet meer herleidbaar tot een persoon en kan worden gesproken over geanonimiseerde gegevens. Dat wordt anders met de verwerking van gepseudonimiseerde persoonsgegevens.



Figuur 2 - Illustratie uit Wikimedia Commons (Akshayhallur, Anonymous, CC BY-SA 3.0)

Gepseudonimiseerde persoonsgegevens kunnen in bepaalde gevallen herleid worden tot de betrokken persoon. Dit kan bijvoorbeeld via een pseudo-ID, zoals het lenerspasnummer. Het kan ook door een gepseudonimiseerde dataset te koppelen aan of te vergelijken met een andere dataset, waardoor de gegevens (indirect) naar een persoon herleid zouden kunnen worden. Dit betekent dat gepseudonimiseerde data als persoonsgegevens worden gezien en dat daarom de AVG hierop van toepassing is.

De keuze was om tijdens het pilottraject te beginnen met geanonimiseerde data en van daaruit te kijken of meer identificeerbare persoonsgegevens nodig waren, afhankelijk van de resultaten die hiermee werden bereikt.

Tijdens de pilot hebben we meerdere keren de dataset uitgebreid met mogelijk herleidbare persoonsgegevens zoals geslacht, geboortjaar, postcode en familierelatie (op hetzelfde adres). Daarbij hebben we ook pseudonimisering toegepast. Bij elke dataset uitbreiding en aanpassing in de verwerkingsmethode van de persoonsgegevens, hebben we onze aanpak opnieuw getoetst aan de huidige verwerkersovereenkomsten, de AVG en de classificatie uit de AI-conceptwet.

Geautomatiseerde besluitvorming

Geautomatiseerde besluitvorming leidt tot een besluit dat zonder menselijke tussenkomst op basis van een geautomatiseerde verwerking tot stand is gekomen. Het besluit kan bestaan uit een maatregel, die tot stand is gekomen op basis van persoonlijke informatie, en die de betrokkene in aanmerkelijke mate treft.⁸ Een voorbeeld van geautomatiseerde besluitvorming is de automatische weigering van een kredietaanvraag of de automatische weigering van een sollicitant.

Profilering is een vorm van geautomatiseerde besluitvorming en bestaat uit de geautomatiseerde verwerking van persoonsgegevens waarbij aan de hand van deze gegevens, bepaalde persoonlijke aspecten van een natuurlijke persoon worden geëvalueerd met het doel de gedragskenmerken van de betrokkene te kunnen analyseren en/of voorspellen.⁹ In de AVG worden als voorbeeld onder andere de volgende gedragskenmerken genoemd: beroepsprestaties, gezondheid, persoonlijke voorkeuren of interesses, betrouwbaarheid, gedrag. Met andere woorden, bij profilering worden gedragskenmerken van de betrokkene beoordeeld om vervolgens op basis van de beoordeling een voorspelling te kunnen doen.

⁸ Overweging 71 AVG.

⁹ Artikel 4 sub 4 AVG.

Het is dus mogelijk dat er profilering plaatsvindt als de persoonsgegevens niet alleen op groepsniveau worden gebruikt, maar ook op individueel niveau.

De toepassing van geautomatiseerde besluitvorming en profilering kan tot een zwaardere inbreuk op de rechten van de betrokken individu leiden. De impact is groter dan wanneer de persoonsgegevens slechts op groepsniveau worden gebruikt voor statistische doeleinden. Hoe groter de impact op de individu, hoe zwaarder zijn/haar belang weegt. De mate waarin geautomatiseerde besluitvorming wordt toegepast kan ertoe leiden dat de fundamentele rechten en vrijheden van de mens zwaarder wegen dan het gerechtvaardigd belang. Mocht de geautomatiseerde besluitvorming zelfs zover gaan dat het de betrokkene in aanmerkelijke mate treft (zoals de geautomatiseerde beëindiging van iemands bibliotheeklidmaatschap), dan is dit in principe verboden.

Is er bij deze pilot sprake van geautomatiseerde besluitvorming en/of profilering?

Om te kunnen spreken van geautomatiseerde besluitvorming moet aan de volgende voorwaarden worden voldaan:

1. de gegevens moeten bestaan uit persoonsgegevens; en
2. persoonlijke aspecten van de lener worden geëvalueerd; en
3. op basis van de evaluatie wordt een besluit genomen die de lener in aanmerkelijke mate treft; of
4. op basis van de evaluatie wordt een voorspelling gedaan van de gedragskenmerken van de lener.

Deze beslisboom hebben we tijdens de pilot gebruikt bij elke uitbreiding van de dataset en aanpassing van de data verwerkingsmethode. Geen enkele keer is er sprake geweest van geautomatiseerde besluitvorming of profilering. Wanneer er wel persoonsgegevens verwerkt werden, was er

geen sprake van het evalueren van persoonlijke aspecten van de lener (voorwaarde 2). De analyse van de gegevens was niet gericht op individuele personen. Bovendien nam het algoritme geen besluiten op basis van de evaluatie (voorwaarde 3) en deed het geen voorspelling over gedragskenmerken van individuele leners (voorwaarden 4).

Toetsing aan het Europese AI-wetsvoorstel

We hebben tenslotte gekeken naar de risicoklasse van de pilot volgens het Voorstel voor Europese Harmonisatie van Kunstmatige Intelligentie (de 'AI-verordening').

Uit de toepassing van de risicoclassificatie van de AI-verordening bleek dat deze AI-pilot in de laagste klasse viel (categorie 4: minimaal risico) en daarmee onder bestaande wetgeving kon worden uitgevoerd.

We hebben ervoor gekozen iedere aanpassing in gebruikte data telkens af te zetten tegen het AI-wetsvoorstel. Zo zijn we er zeker van dat de pilot nog steeds juist was geclassificeerd. Hiermee wilden we de pilot "future-proof maken", zodat de AI-oplossing naar verwachting ook voldoet aan de aankomende Europese wet- en regelgeving.

Van pilot naar productiesysteem

Deze pilot had als doel om de juridische, ethische en technische mogelijkheden van AI voor bibliotheekdata te onderzoeken. De toepassing moet opnieuw getoets worden aan de AVG, het AI-wetsvoorstel en de bestaande verwerkingsovereenkomsten mocht het in de toekomst in de praktijk worden gebracht. Dit zou eventuele aanpassingen teweeg kunnen brengen, zoals het aanpassen van de privacyverklaringen en de bibliotheekgebruikers een optie voor opt-out geven.



Conclusie juridische aspecten

Een van onze onderzoeksvragen was: *Aan welke juridisch-ethische voorwaarden moet de inzet van AI voldoen, bij het voorspellen van het inlevermoment van uitgeleend bibliotheekmateriaal?*

Op juridisch vlak is het allereerst van belang te bepalen welke data gebruikt gaan worden. Als persoonsgegevens in het spel zijn, dan is het noodzakelijk de verwerking van deze gegevens te toetsen aan de AVG voor privacybescherming en het rechtmatig omgaan met persoonsdata. *Privacy by design* is dan het leidende principe. Toetsing aan de AVG moet vanaf het begin gebeuren. Bij iteratief ontwerp van een AI-systeem, waarbij de dataset in vervolgitaties wordt uitgebreid met (meer) persoonsgegevens en/of de aard van de dataverwerking verandert, moet ook toetsing iteratief plaatsvinden. In hun rol van Verwerkingsverantwoordelijke geven de betrokken bibliotheken toestemming de persoonsgegevens die zij verzamelen van hun leners te hergebruiken voor de casus. Daarbij is doelbeperking een belangrijke voorwaarde. Omdat de casus gericht is op de verbetering van de informatievoorziening aan

de lener met betrekking tot reserveringen, past de doelstelling binnen de verwerkingsdoelstelling van de betrokken bibliotheken. Wanneer de verwerking van geanonimiseerde data onvoldoende resultaat oplevert en overgegaan wordt op pseudonimiseren en/of profilering, dan speelt het gerechtvaardigd belang een rol. Anders gezegd, dan moet het gebruik van persoonsgegevens in verhouding staan tot het doel en wegen de rechten van de lener zwaarder. Voor deze casus weegt het gerechtvaardigd belang niet zo zwaar: verbetering van de informatievoorziening aan de lener is een goed streven, maar rechtvaardigt niet een zware inbreuk op diens rechten. In hoeverre is de verwerking van persoonsgegevens hiervoor *noodzakelijk*? Deze belangenafweging wordt verder besproken in het hoofdstuk over de ethische aspecten van deze casus.

Voorts speelt de wetgeving rondom AI een rol voor het reguleren van AI-systemen op basis van het verwachte risico dat het systeem met zich mee brengt. De casus valt in de categorie minimaal risico. Dit betekent dat de betreffende AI-oplossing in principe niet onder een zwaarder regime van regulering en monitoring hoeft te vallen.

Ethische aspecten

Naast de juridische aspecten, spelen de ethische aspecten van het gebruik van persoonsgegevens en AI een rol. Hierbij staat de vraag centraal: ook al mag het (juridisch) en kan het (technisch), moet je het wel willen (ethisch)? In dit hoofdstuk gaan we in op de ethische kaders en de ethische verkenning rond de casus van de AI-pilot.

Men is al snel geneigd om te denken: „als het mag van de wet“ en we legaal handelen, dan zit het wel goed.

Maar, de casus riep toch kritische vragen op: Waarom zou je persoonsgegevens willen gebruiken – met alle risico's van dien – om te voorspellen wanneer een uitgeleend boek beschikbaar komt? Met andere woorden: heiligt het doel de middelen?

Bij de juridische toetsing kwamen telkens ook ethische vragen naar boven, waardoor we ons steeds bewuster werden van de vervlechting van de juridische en ethische aspecten. Neem bijvoorbeeld het anonimiseren van persoonsgegevens: wat zijn de precieze afwegingen wanneer persoonsgegevens wel of niet „herleidbaarheid“ zijn? Dat is een glijdende schaal van risicoloos naar risicovol. En hoe anoniem kun je het maken met zoveel online bronnen, waar ieders gegevens openbaar voor het oprapen liggen, zoals op Facebook, Twitter, Instagram of LinkedIn?

Hoeveel weten we al wel niet over onze leners met alle persoonsgegevens die bibliotheken in hun systemen verzamelen? En wat doen we met die kennis? Op basis van verlengingsgegevens weten we bijvoorbeeld dat leners van 28 jaar en jonger een boek gemiddeld langer dan 28 dagen in huis hebben. Wat betekent dat voor de mogelijkheid die informatie te gebruiken voor profilering? Stel dat bibliotheken mensen van onder de 28 jaar eerder een herinnering gaan sturen? Het feit dat iets mogelijk is betekent niet automatisch dat het ethisch gezien ook verantwoord is.

Wetten ontwikkelen zich meestal niet in hetzelfde tempo als technologie. Hoewel wettelijke kaders enige ethische aanknopingspunten bieden voor evaluatie van de belangen van betrokkenen kunnen deze soms de ethische normen niet volledig weergeven of zijn eenvoudigweg niet goed geschikt om bepaalde dilemma's aan te pakken. Daarom zijn ook ethische vraagstukken van groot belang bij het ontwikkelen van een AI-oplossing.

Ethische kwesties

Een van onze onderzoeksvragen was: Aan welke juridisch-ethische voorwaarden moet de inzet van AI voldoen bij het voorspellen van het inlevermoment van uitgeleend materiaal?

Aan het begin van de pilot hadden we een aantal algemene bedenkingen en vragen naar aanleiding van het soort gebruik van persoonsgegevens dat wij tijdens de pilot wilden gaan uitvoeren, zoals:

- Welke basiseigenschappen (van persoonsgegevens) staan we (als bibliotheken) toe voor AI-projecten in het algemeen (bijvoorbeeld: voor de inlevermomentpilot, of een recommender-project waarbij de lener op basis van uitleengegegevens geattendeerd wordt op een andere titel)?
- Is het acceptabel voor bibliotheken om persoonsgegevens te delen voor dit doel?
- Wil je als bibliotheek lenersprofielen opbouwen? En, wie mag daar dan bij? Alleen het algoritme of ook medewerkers/ andere systemen?
- Hoe zit het met de balans tussen de belangen van de leners van wie de gegevens worden gebruikt en die van (andere) leners die daardoor beter kunnen worden geïnformeerd?

De voorwaarde dat het gebruik van lenersgegevens in verhouding moet staan tot het resultaat, vereist ook een ethische afweging. Tijdens de iteratieve data-exploraties van de pilot zelf lag de focus vooral op het binnen de wettelijke regels

werken met de data en ervoor zorgen dat het gebruik van persoonsgegevens in verhouding staat tot het resultaat. Dat laatste zorgde ervoor dat we tijdens de derde iteratie besloten de pilot te beëindigen. De resultaten werden namelijk niet beter door het toevoegen van meer detail in de persoonsgegevens en het categoriseren van leners (zoals incidentele leners of grootgebruikers).

Toch wilden we de ethische aspecten - die tot dan toe een impliciete rol gespeeld hadden - beter in kaart brengen en expliciteren. Daartoe hebben we gezocht naar ethische kaders en richtlijnen, en een ethische verkenning uitgevoerd.

Ethische kaders

Ter oriëntering hebben wij een aantal richtlijnen en kaders onder de loep genomen. Allereerst, in het verlengde van onze juridische verkenning, hebben we gekeken naar de Europese richtlijnen:

Europese richtlijn voor verantwoord gebruik van Kunstmatige Intelligentie (2019) van de High-Level Expert Group on Artificial Intelligence van de Europese Commissie (EC). De richtlijn adviseert een viertal ethische beginselen in acht te nemen bij de ontwikkeling, installatie en het gebruik van AI-systemen:

1. respect voor menselijke autonomie;
2. preventie van schade;
3. rechtvaardigheid;
4. verantwoording.

Deze beginselen zijn gebaseerd op de fundamentele waarden die zijn vastgelegd in de EU-verdragen en het Handvest van de grondrechten van de EU, zoals: respect voor de menselijke waardigheid, voor democratie, justitie en de rechtsstaat.

Voortbouwend op deze beginselen heeft de Expert Group zeven vereisten opgesteld waaraan AI-systemen moeten voldoen om als betrouwbaar te kunnen worden bestempeld:

1. menselijke controle en menselijk toezicht;
2. technische robuustheid en veiligheid;
3. privacy en data governance;
4. transparantie;
5. diversiteit, non-discriminatie en rechtvaardigheid;
6. milieu- en maatschappelijk welzijn;
7. verantwoordingsplicht.

The Assessment List for Trustworthy Artificial Intelligence (ALTAI) for self-assessment (2020).

Dit is een checklist die ontwikkelaars en gebruikers van AI helpt om te voldoen aan de zeven eerder genoemde vereisten van de Europese richtlijn. ALTAI geeft handvatten voor het uitvoeren van een zelfevaluatie bij de implementatie van AI-systemen. ALTAI was niet een geschikte checklist voor onze pilot, aangezien implementatie en in productie nemen van een AI-systeem buiten onze scope was.

Vijf van de zeven vereisten van de Europese richtlijn over verantwoord gebruik van AI overlappen met bestaande wetgeving en benadrukken enerzijds de noodzaak om binnen wettelijke kaders en verplichtingen te werken en anderzijds de noodzaak van veilige, betrouwbare en robuuste technologie. Wij waren vooral op zoek naar referentiekaders voor een waardensysteem waarmee we een discussie konden voeren rond de ethische kwesties die onze AI-pilot opwierp.

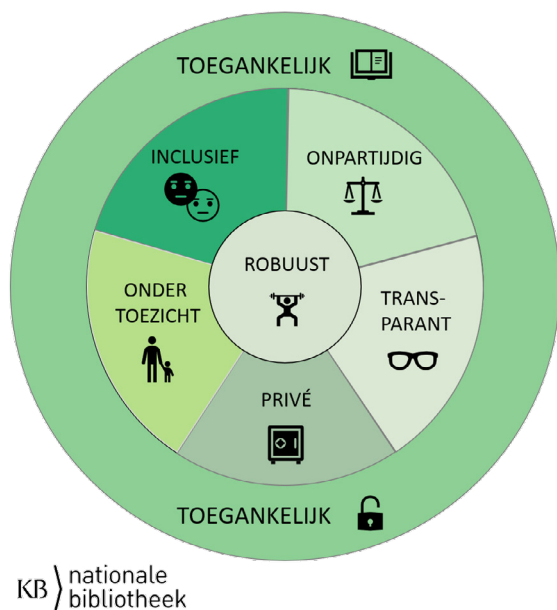
De openbare bibliotheek in Nederland is een openbare voorziening die een publieke taak voor het algemeen publiek vervult. We hebben daarom gekeken naar ethische richtlijnen voor de publieke sector in Nederland en openbare bibliotheken, in het bijzonder.

Toolbox voor Ethisch Verantwoorde Innovatie

van het ministerie van Binnenlandse Zaken (BZK). Deze toolbox wil overheidsorganisaties helpen om op een ethisch verantwoorde manier te innoveren, "met respect voor belangrijke publieke waarden en grondrechten". Aan de basis liggen zeven kernprincipes:

1. Zet publieke waarden centraal;
2. Betrek burgers en andere belanghebbenden;
3. Zorg dat je relevante wet- en regelgeving respecteert;
4. Let op de veiligheid van technologie;
5. Zorg voor biasvrije data, algoritmes en analysemethoden;
6. Wees transparant en leg verantwoording af;
7. Monitor, evalueer en stel bij.

De door de Koninklijke Bibliotheek (KB) geformuleerde “**AI en de bibliotheek: zeven principes**”. De visuele weergave in Figuur 3 geeft de samenhang tussen de zeven principes aan:



Figuur 3 - De KB zeven principes

- Om de cirkel heen bevinden zich de gebruikers van de bibliotheek. Het vergroten van de toegankelijkheid van informatie ligt op deze buitenste cirkel het dichtst bij de gebruikers. Hiermee worden de positieve mogelijkheden van AI benadrukt.
- Robuustheid bevindt zich in de binnenste cirkel – hiermee is een AI-systeem in de kern niet wezenlijk anders dan elk ander software-systeem: het moet aan dezelfde betrouwbaarheidseisen voldoen.
- De principes in de middelste ring wijzen naar vijf ethische kwesties die door de toepassing van AI specifiek aandacht behoeven: inclusief, onpartijdig, transparant, privé, onder toezicht.

Grondrechten (van de mens, van de EU) en fundamentele waarden (publieke waarden), en de waarden die openbare bibliotheken specifiek uitdragen vormen een reeks principes die potentieel allemaal van belang zijn: rechtvaardigheid, onpartijdigheid, solidariteit, respect voor menselijke autonomie, recht op informatie, enzovoorts. Maar hoe ga je hiermee aan de slag?

Er zijn richtlijnen en instrumenten die zich ook bezig houden met de methode rond ethiek en AI.

- Een van de zeven principes van de eerder genoemde Toolbox van het ministerie van BZK is om AI op een democratische wijze te ontwikkelen. Dat doe je door burgers en andere belanghebbenden te betrekken in het proces: “Wanneer we technologieën op een democratische wijze ontwikkelen en implementeren, kunnen we waarden, belangen, verwachtingen en zorgen beter in het ontwerp meenemen. Bovendien draagt participatie van gebruikers bij aan hun begrip van de technologie (en bijbehorende regelgeving).”
- **De Ethische Data Assistant (DEDA) van Utrecht Data School (2022)** is ontstaan uit de behoefte naar ethische kaders voor Big Data-onderzoek en daarna als instrument aangepast voor ethisch gebruik van data binnen overheidsorganisaties. DEDA is een methode die – in tegenstelling tot het aftikken van checkboxes – participatief is. Het stelt een reeks vragen die in teamverband, met mensen met verschillende functies en achtergronden, besproken kunnen worden om de ethische kwesties beter in beeld te krijgen en samen te bespreken en overwegen.
- Na afronding van onze pilot publiceerde de NLAIC de publicatie ‘**Ethiek en AI – Zeven methoden in theorie en praktijk**’. Hierin wordt een overzicht gegeven van een zevental mogelijke methoden die zich richten op de ethische toepassing van AI.

De verkenning van ethische kaders leidde tot de conclusie dat er geen kant en klare methode was die geschikt zou zijn voor onze pilot.

We hadden weliswaar een multidisciplinair team gevormd, maar we misten de expertise en ervaring op het ethische vlak. Daarom riepen we ondersteuning in van Filosofie in actie, een adviesorganisatie over het ethisch gebruik van data en technologie.

Ethische verkenning



Figuur 4 - Illustratie van Filosofie in actie

Tijdens de ethische verkenning in de vorm van een workshop hebben we nader gekeken naar de probleemdefinitie van de pilot, de belangen van de betrokkenen en de waarden die in het spel zijn. Hieruit kwamen een aantal constatering en aanbevelingen.

Allereerst het belang van een gezamenlijk, goed gedefinieerde probleemstelling. Wat proberen we

nu precies op te lossen? We wilden “meer duidelijkheid kunnen geven over wanneer een boek beschikbaar komt, omdat dit beter aansluit bij huidige verwachtingen van de lener.” Een goed bedoeld voornemen, maar wel eentje met een flink aantal aannames die we van tevoren niet voldoende beargumenteerd hadden en ook niet gevalideerd met de gebruikers.

Zo voorspelt het algoritme het inlevermoment en dat is iets anders dan wanneer het boek ook daadwerkelijk beschikbaar komt. De logistieke processen om het boek bij de lener te krijgen nemen ook tijd in beslag die eveneens moeilijk in te schatten zijn, omdat de processen vaak erg complex zijn.

Zelfs als ons algoritme nauwkeurig had kunnen voorspellen, blijft het een voorspelling. Het biedt nog geen zekerheid aan de lener; het is slechts een waarschijnlijkheid. In hoeverre zou dit dan een verbetering van de dienstverlening zijn?

Er speelden in de probleemdefinitie ook belangen, motieven en probleempercepties mee van de verschillende pilotdeelnemers die we niet geëxpliciteerd hadden. Dit werd duidelijk toen we een stakeholderanalyse hadden gedaan.

Wie zijn alle betrokken partijen en hoe zitten zij erin met betrekking tot het probleem dat we wilden aanpakken?

Stakeholdersanalyse Bibliotheek

Stakeholder	Recht/plichten	Belang	Wens
Bibliotheek	- Vrije informatievoorziening - Toegankelijkheid van informatie	- Klantenbinding - Leners/uitleencijfers - Bestaansrecht kunnen verantwoorden - Maatschappelijke rol blijven vervullen (ipv commerciële partij)	- Ontwikkelen/informereren burger - Goede dienstverlening - Bepalen termijn waarop een exemplaar nog niet geleverd wordt aan andere bibliotheek
Betaalde medewerkers		- Informatie over de lening (voorspelbaarheid) en logistiek proces	- Om de lener te kunnen informeren - Geen discussies met klanten / leuke gesprekken - Serieus genomen worden door klanten
Vrijwilligers in bieb		- Informatie/duidelijkheid over proces	- Leners informeren - Prettige sfeer

Figuur 5 - De bibliotheek - Stakeholderanalyse - Workshop Filosofie in actie - 2022

Zo bleek bijvoorbeeld dat bibliotheekmedewerkers de behoefte hebben om leners goed te informeren en duidelijkheid te verschaffen. Dit is een duidelijke drijfveer voor deze casus geweest.

Zij hebben steeds meer te maken met veranderende verwachtingen van leners, die een online reservering als een online bestelling zien. Door de verplaatsing naar online, hebben ze steeds minder contact met leners – en contact over reserveringen betekent al gauw het afhandelen van een klacht. Ze merken dat de vraag naar de beschikbaarheid

van boeken sinds de COVID-19-pandemie blijvend is toegenomen door de flinke stijging van het aantal reserveringen. Dit knelpunt zien ze niet snel opgelost.

Medewerkers hebben behoefte om de digitale processen uit te kunnen leggen en daar hoort een nauwkeurige voorspelling kunnen geven ook bij.

Stakeholdersanalyse Lener

Stakeholder	Recht/plichten	Belang	Wens
Lener (reserveert)	- Lidmaatschap	- Financieel belang	- Snelle/adequate levering - Lezen
Leners als geheel / collectief		- Betaalbaar	- Interessante en beschikbare collectie - Gastvrije bibliotheek-medewerkers / servicegericht
Lener (huidige lener)	- Plicht: tijdig inleveren - Recht: bepaalde tijd lenen		- Boek kunnen (uit)lezen - Niet opgejaagd worden - Prettige leeservaring (ongestoord)

Figuur 6 - De lener - Stakeholderanalyse - Workshop Filosofie in actie - 2022

De grote afwezige was de lener zelf. We werkten weliswaar vanuit de bevindingen van BiebPanel en de praktijkervaringen van bibliotheekmedewerkers met het afhandelen van reserveringen, maar een diepgaander gesprek met leners over hun behoeften en verwachtingen ontbrak.

Gaat het de lener wel om een voorspelling te krijgen - over wanneer een boek beschikbaar komt - of gaat het toch vooral om de termijn? Om het lange wachten voor een reservering? Mogelijk waren we op een aangepaste probleemdefinitie uitgekomen

als we direct vanaf het begin met een lenerspanel in gesprek waren gegaan.

Tijdens de workshop keken we naar de noodzaak voor het inzetten van AI en mogelijke alternatieve middelen om tot hetzelfde doel te komen. We dachten aan: de bibliotheek koopt een nieuw exemplaar van het gereserveerde boek zodat er meer in omloop zijn; alternatieve leenopties/ander materiaal-soort (bijvoorbeeld e-book) aanbieden; populaire boeken korter uitlenen of het verlengen niet toe te staan; een puntensysteem invoeren waarbij leners voorrang kunnen verdienen in de reserveringswachtrij; reserveringen afschaffen.

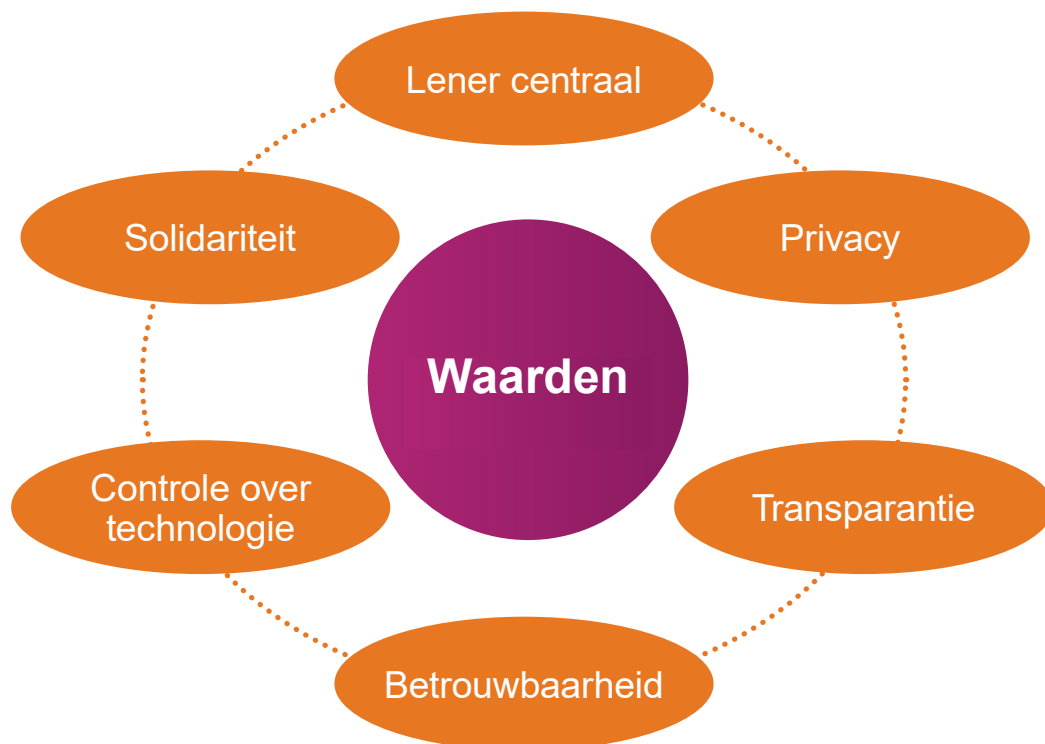
Bibliotheekmedewerkers zouden in de huidige situatie ook naar de opgebouwde leenhistorie van de huidige lener kunnen kijken en op basis daarvan het inlevertmoment inschatten. In hoeverre is dit een inbreuk op de privacy van de huidige lener? Sommige bibliotheken gebruiken de leenhistorie op individueel niveau om leners te attenderen op het feit dat ze een aangevraagde titel al eerder hebben geleend. Dat is ethisch gezien minder problematisch.

We vroegen ons af in hoeverre er niet meer resultaat en voordeel te behalen valt met het vereenvoudigen of efficiënter maken van de achterliggende logistieke processen dan met het voorspellen van het inlevertmoment? Een aantal van deze processen zijn weliswaar in kaart gebracht (bijvoorbeeld transport van boeken tussen bibliotheken), maar het geheel is te complex om er op korte termijn iets aan te kunnen doen. Dat zou een signaal kunnen zijn om er juist nog eens goed naar te kijken – al dan niet met behulp van AI. Welbeschouwd, valt er beter data-analyse te doen op processen waar men wel grip op heeft.

Tenslotte hebben we een waardenverkenning gedaan rondom de casus. De waarden ten aanzien van het algoritme van onze AI-pilot die het hoogste scoorden waren: betrouwbaarheid, privacy, transparantie en controle over technologie.

De 3 bibliotheekwaarden die door de deelnemende bibliotheken van cruciaal belang werden geacht voor deze casus, waren: de lener centraal, de betrouwbare bibliotheek en solidariteit.¹⁰

Vooraf het bespreken van deze laatste categorie gaf goed inzicht in hoe de bibliotheek worstelt met de waarde solidariteit en de vraag: Hoe gaan we om met de klassieke bibliotheekwaarde 'solidariteit' in tijden waarin de verwachtingen van de leners veranderen? Stroken de nieuwe vormen van solidariteit (bijvoorbeeld gelegenhedsvrijwilligers; "sharing is caring") nog wel met de soms omslachtige uitleenpraktijken in bibliotheken?



Figuur 7 - Waarden-analyse – Workshop Filosofie in actie - 2022

¹⁰ Ter vergelijking, artikel 4 van de Wet stelsel openbare bibliotheekvoorzieningen (Wsob) zegt: 'Een openbare bibliotheekvoorziening heeft een publieke taak die zij voor het algemene publiek vervult op basis van de waarden onafhankelijkheid, betrouwbaarheid, toegankelijkheid, pluriformiteit en authenticiteit'. (<https://wetten.overheid.nl/BWBR0035878/2022-07-01>)

Moet de verantwoordelijkheid meer bij de lener gelegd worden? Willen leners liever onderling met elkaar bemiddelen over het inlevermoment? En tegelijk hun leeservaring met elkaar delen? Kan solidariteit beter gefaciliteerd worden door digitale platforms dan bemiddeld door de bibliotheek? Een discussieonderwerp op zichzelf!

Samengevat heeft de ethische verkenning ons veel inzichten opgeleverd en ons doordrongen van het belang om ook ethische vragen al in de ontwerpfase van technologiegedreven projecten mee te nemen (ethical-by-design).

Conclusie ethische aspecten

Een belangrijk doel van de pilot was om aan de hand van de casus (het voorspellen van het inlevermoment van uitgeleend bibliotheekmateriaal) de grenzen te verkennen van de juridische en ethische kaders bij de inzet van AI-technologie en het gebruik van persoonsgegevens. De directies van de drie pilotbibliotheken hadden vanaf het begin wel een beperkende voorwaarde gesteld, namelijk dat het gebruik van persoonsgegevens uit hun lenersbestand in verhouding moest staan tot het resultaat. Dit vereiste naast een juridische, ook een ethische afweging tijdens de iteraties waarbij persoonsgegevens werden gebruikt. Bovendien kwamen tijdens de juridische toetsingen ook ethische vragen naar voren, waardoor we ons steeds bewuster werden van de vervlechting van de juridische en ethische aspecten.

Op ethisch vlak zijn er betrekkelijk weinig bruikbare kaders, richtlijnen, methoden of best practices. De richtlijnen die er zijn hanteren een aantal basisbeginselen of principes die gestoeld zijn op grondrechten, rechten van de mens en fundamentele waarden of publieke waarden. Er zit veel overlap met de principes die gehanteerd worden in bestaande regelgeving over de bescherming van persoonsgegevens en best practices voor het implementeren van digitale technologie. De kernwaarden van het toepassingsdomein – in ons geval dat van openbare bibliotheken – spelen ook een belangrijke rol. Er kan dus uit een groot aantal generieke waarden geput worden, maar de ethische waarden die ertoe doen zijn per casus verschillend.

In ons geval ging de casus over het analyseren van uitleengegevens om leners beter te kunnen informeren. Enerzijds doet dit een beroep op de betrouwbaarheid van de bibliotheek inzake de correcte en rechtmatige omgang met de haar toevertrouwde persoonsgegevens en anderzijds op een zekere mate van solidariteit van de leners: leners helpen elkaar indirect door toe te staan dat hun leengedrag beschikbaar te stellen te informeren mijn leengedrag helpt om te voorspellen wanneer een gereserveerd boek waarschijnlijk beschikbaar komt. Met betrekking tot het algoritme en de data, golden de waarden privacy, transparantie, controle over technologie en betrouwbaarheid van het algoritme. De principes van robuustheid en veiligheid van een AI-systeem waren voor de pilot minder relevant omdat er tijdens de pilot geen AI-systeem geïmplementeerd zou worden.

De wijze waarop relevante ethische waarden worden toegepast binnen een (iteratief) AI-ontwikkelproces is afhankelijk van de situatie. Idealiter gebeurt dit in de interactie van een multidisciplinair team waar zeker ook een ethicus deel uit maakt: iemand met expertise en ervaring op het gebied van AI- en data-ethiek. De rol van zo'n expert binnen het team is om kritische vragen te stellen en de discussie over juridisch-ethische afwegingen in goede banen te leiden. Er zijn veel vragen die nooit grondig besproken en beantwoord worden tijdens trajecten zoals deze. Bij de inzet van AI is vanaf het begin al veel winst te behalen door duidelijkheid te scheppen over welke belangen gediend worden met een voorgestelde AI-oplossing: die van de lener, de bibliotheekmedewerker, de pilotpartijen, de subsidieverstrekker? Wat zijn hun motieven? Hun perceptie van het probleem of mogelijkheid? Als het gaat om de verbetering van de dienstverlening, is de bibliotheekgebruiker dan wel betrokken? Is de verwerking van persoonsgegevens wel noodzakelijk? Wat is de noodzaak om AI in te zetten? In hoeverre heiligt het doel de middelen? Wat zijn alternatieven om hetzelfde doel te bereiken? Deze kritische insteek zorgt ervoor dat er niet alleen nagedacht is over het doel van een AI-pilot, maar ook over de belangen van betrokkenen en rechtvaardiging van de voorgestelde aanpak, dat er niet alleen over resultaten en successen gerapporteerd wordt, maar ook over de afwegingen, inzichten en opgedane ervaringen die meegenomen kunnen worden in vervolgpiloten.

Technische aspecten

De bestudering van de juridisch-ethische aspecten was niet louter een theoretische oefening. Het ging gepaard en werd gevoed door de bevindingen van de AI-experimenten en data-exploraties die we tijdens de pilot hebben uitgevoerd. Dit hoofdstuk gaat in op de details van deze experimenten en data-exploraties.

Relevante AI-begrippen

Voordat we in de technische details duiken is het van belang de meest relevante, aan AI gerelateerde begrippen toe te lichten. Zoals eerder, in de inleiding van het rapport aangegeven, behoort het soort AI waar we mee gewerkt hebben tot de sub-categorie van machine learning.

ML richt zich op het bouwen van algoritmische modellen die patronen en relaties in gegevens kunnen identificeren. Het algoritme wordt eerst getraind op een vooraf gedefinieerde dataset en vervolgens gebruikt om bijvoorbeeld (nieuwe) data te classificeren of om voorspellingen te doen. Deze training gebeurt vaak door middel van “supervised learning”. Hierbij wordt het systeem niet alleen gevoed met input (de gegevens), maar ook met het gewenste antwoord. Het doel is dat het systeem zichzelf aanleert om de input te vertalen in de gewenste output.

Er bestaat inmiddels een heel scala aan algoritmische modellen, elk met zijn voor- en nadelen en/of specifieke geschiktheid voor bepaalde datasets. Enerzijds is er de analytische aanpak, waarbij (vooraf, in een hypothese) verbanden worden gelegd tussen input en output in een functioneel model met enkele parameters die worden aangepast

aan – of getraind op – de verkregen data.

In zo'n geval laat het model een voorspelling gemakkelijk verklaren op grond van de input. Anderzijds is er met de opkomst van Big Data en met vorderingen op het gebied van ML en pure rekenkracht, een tendens ontstaan om grote hoeveelheden data massaal en menigmaal aan te bieden aan modellen gebaseerd op neurale netwerken, het zogenaamde Deep Learning (DL). In het geval van DL, worden de algoritmen gestructureerd naar voorbeeld van de werking van de hersenen, die een biologisch neurale netwerk vormen.

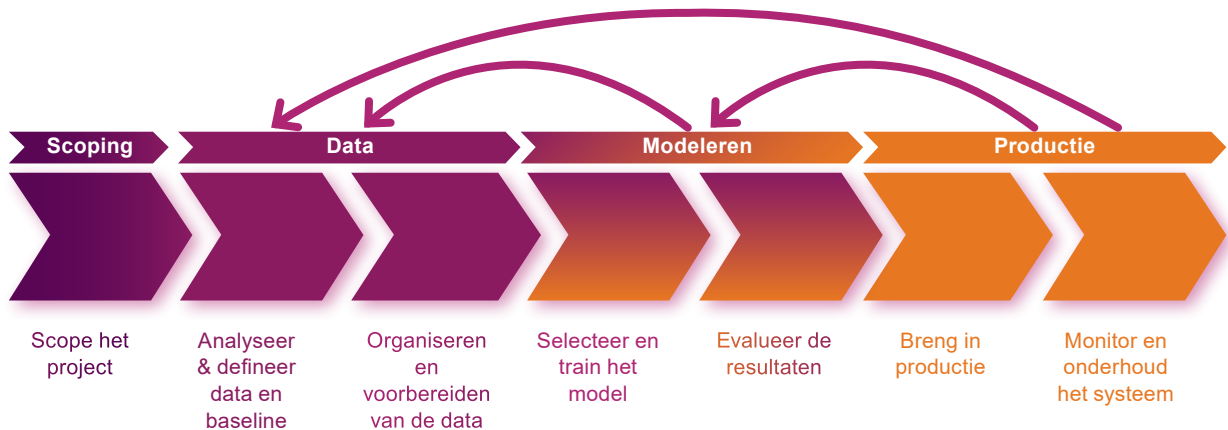
Een manier om deze modellen in te delen is naar de transparantie die ze bieden op het verband tussen de input (data) en de output (data). Een white-box model is transparant (“explainable by design”) en geeft inzicht in hoe een bepaalde uitkomst tot stand komt. Bij black-boxmodellen is dat niet het geval. In het algemeen gesproken zijn de analytische modellen white-box, terwijl de op neurale netwerken gebaseerde modellen black-box modellen zijn, waarbij het lastiger is uit te leggen waarom een model een bepaalde waarde voorspelt.¹¹

Een belangrijk gegeven bij een AI-oplossing is of deze zijn werk kan doen als er nog niet voldoende gegevens beschikbaar zijn, het zogenoemde Cold Start probleem. Denk aan een recommender van een streamingdienst die een advies moet geven op een film die net uit is en nog niet bekeken. Of denk in ons geval aan een voorspelling van het inlevermoment als een nieuw bibliotheeklid voor het eerst een boek leent. Er zijn steeds betere manieren om het Cold Start probleem te omzeilen, waarbij modellen in staat zijn gegevens te groeperen en - ondanks het ontbreken van transactiegegevens - toch een voorspelling te doen op groepsniveau.

¹¹ Zie bijvoorbeeld Transparency, auditability, and explainability of machine learning models in credit scoring, Bucker et al, 2021, <https://www.tandfonline.com/doi/full/10.1080/01605682.2021.1922098>

Machine Learning proces

Het proces om te komen tot een ML-oplossing die aan de verwachting voldoet is een iteratief proces van modelbouw en data-exploratie, zie Figuur 8.



Figuur 8 - Iteratief machine learning proces

Het proces bestaat uit een aantal stappen: scopen van het project, definiëren en voorbereiden van de dataset, selecteren en trainen van het model, en evalueren van de resultaten. Bij tevredenheid met de resultaten wordt het AI-systeem in productie gebracht en continu gemonitord en onderhouden. Meestal is de flow niet zo stapsgewijs van links naar rechts, maar is er sprake van iteratieve stappen. Na het evalueren van het model komen we soms tot de conclusie dat we meer of andere data nodig hebben. Of na een tijd in productie te hebben gedraaid blijkt het model te lijden aan “data drift”. Het model is dan als het ware aan slijtage onderhevig en moet opnieuw getraind worden op meer actuele data.

Voorspellingsdoel en scope

De doelstelling van de pilot was de mate van voorspelbaarheid van het inlevermoment te evalueren met behulp van een AI-model.

Vooraf is het doel geformuleerd om het inlevermoment met een nauwkeurigheid (standaarddeviatie) van +/- 2 dagen te kunnen voorspellen. De gedachte hierbij was dat bij een dergelijke nauwkeurigheid de voorspelling nog nuttig zou kunnen bijdragen aan het informeren van gebruikers bij een reservering.

Bij lagere nauwkeurigheid (een hoger aantal dagen) is er geen of weinig meerwaarde.

We hebben keuzes moeten maken ten aanzien van tijdsbesteding en dataverzameling. We hadden op voorhand het implementeren en in productie nemen van een eventuele AI-oplossing buiten scope geplaatst. We hadden bovendien een doorlooptijdslimiet van een half jaar gesteld aan het iteratief proces van modeleren en data-exploratie. We hebben ervoor gekozen om te werken met de (historische) gegevens waar de pilotbibliotheken over beschikten en die opgeslagen zijn in het bibliotheeksysteem Wise. We besloten geen nieuwe/ additionele (persoons)gegevens te verzamelen. Wel verrijkten we de dataset met data verkregen

via het Centraal Bureau voor de Statistiek (CBS). In de discussies met de pilotgroep werden diverse – mogelijk van invloed zijnde – gebeurtenissen of factoren genoemd waarvan geen data in Wise of bij het CBS beschikbaar is, of waarvan het lastig is deze te verzamelen. Voorbeelden hiervan zijn: openingstijden van de bibliotheek, het hebben van een inleverbrievenbus, weersvoorspellingen, evenementen zoals het WK-voetbal. Deze gegevens zijn niet meegenomen in de dataset van de pilot.

Definitie van de dataset

De dataset van de pilot is in verschillende iteraties opgebouwd, om antwoord te kunnen geven aan de onderzoeksvragen:

- Zijn voorspellingen op basis van niet-persoonsgebonden gegevens nauwkeurig genoeg?
- Welke persoonsgegevens zijn nodig voor (meer) betrouwbare en nauwkeurige voorspellingen?

Basisdataset

De basisdataset bestond uit titel- en exemplaargegevens en geanonimiseerde uitleengegevens van de drie pilotbibliotheken voor fysieke materialen (boeken, DVD's, enzovoorts) in de categorie "Volwassen-Fictie".

Voor het construeren van de dataset werden de gegevens meegenomen die relevant leken voor het voorspellen van de retourdatum. Van de titelgegevens gebruikten we bijvoorbeeld de mediumsoort, en bij boeken het aantal pagina's, omdat dit een maat kan zijn voor hoelang iemand het boek nodig heeft. Van het exemplaar gebruikten we bijvoorbeeld het gegeven hoelang het boek al in de collectie zit, omdat slijtage van oudere exemplaren mogelijk een rol speelt bij de inleverdiscipline.

We gebruikten de standaard circulatietransacties zoals het Wise-systeem deze administreert:

- actie 1: uitleening;
- actie 2: inleveren;
- actie 3: verlengen.

Elke actie kent een datum en tijd en een plaats van handeling. Bovendien registreert Wise bij de transacties ook de administratieve datums die van belang zijn: uitleendatum, uiterste inleverdatum en oorspronkelijke uiterste inleverdatum (bij een verlenging). Op basis hiervan kunnen uitleencasussen worden gevormd die bestaan uit rijtjes acties behorende bij dezelfde lener en exemplaar: uitleening, vervolgens nul of meer verlengingen, tenslotte inname (ofwel in Wise termen: actie 1, nul of meer keren actie 3, gevolgd door een actie 2).

De instellingen voor het uitleenen in Wise per bibliotheek zijn ook van belang. Denk hierbij aan het boetetarief of het maximaal aantal toegestane verlengingen. Wise biedt de mogelijkheid om oude en nieuwe instellingen naast elkaar te bewaren, maar die faciliteit wordt lang niet altijd benut. Het gevolg is dat we referentiegegevens misten uit de tijd van de transacties en bijvoorbeeld niet direct konden zien wat het effect was van het invoeren/veranderen van een bepaalde instelling. Als de datum van invoering bekend is, kan dit natuurlijk soms wel herleid worden als indirecte oorzaak van verandering in inlevergedrag, maar een direct verband kan dan niet makkelijk worden aangetoond op basis van de data. We namen wel de administratieve indeling van exemplaren (gegeven door de plaats-vestiging) en administratieve indeling van transacties (gegeven door de uitleenvestiging) mee omdat die als het ware de grondslag vormen voor dit soort instellingen.

De afgesproken bewaartermijn van transactiegegevens in Wise is 5 jaar, dus de pilot kon beschikken over de betreffende uitleengegevens over de periode 2017 - 2022. Hierin vielen ook de eerste twee jaren van COVID-19 en juist in die periode werden allerlei maatregelen getroffen die gevolgen hadden voor het inleveren van materialen.

We besloten de coronajaren 2020 en 2021 in dit onderzoek verder uit te sluiten vanwege te grote afwijkingen door atypische openingsuren en uitleenafspraken.

Uitbreiding met persoonsgegevens

Tijdens de pilot werd besproken welke persoonsgegevens toegevoegd zouden worden aan de dataset.

De dataset zou stapsgewijs worden uitgebreid met meer persoonsgegevens om te onderzoeken welke gegevens nuttiger waren voor nauwkeurigere voorspellingen.

Eerst gebruikten we gepseudonimiseerde gegevens van leners uit het bibliotheekstelsel, zoals leeftijd, adres, type abonnement en boetes. Later keken we ook naar gegevens van het CBS, zoals afstand tot de bibliotheek en financiële status van huishoudens.¹²

Zo ontstonden er 3 iteraties van data-exploratie:

1. Zonder persoonsgegevens;
2. Met persoonsgegevens (eerste verkenning);
3. Met iets gedetailleerder persoonsgegevens (tweede verkenning).

Verderop vertellen we meer over deze iteraties en de resultaten ervan.

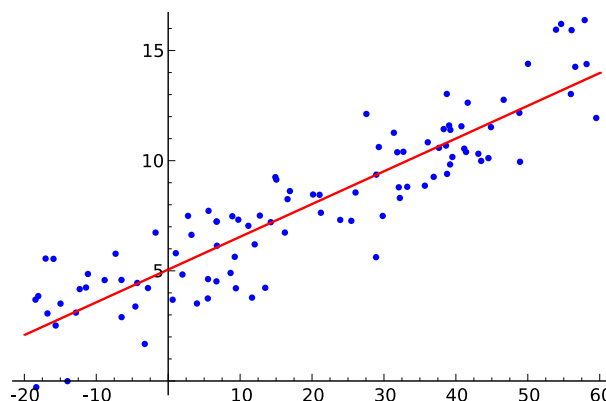
Selecteren van het model

Lineair regressiemodel

Na het bepalen van de data waarmee gewerkt zou worden, was het tijd om een ML-model te kiezen. Op basis van kennis en ervaring binnen het Data Science team van OCLC, werd gekozen voor een analytische aanpak, namelijk het lineair regressiemodel. In Figuur 10 zien we deze techniek toegepast op een simpele eendimensionale dataset, waar de uitkomst – op de verticale as – afhankelijk

is van slechts één variabele (de “voorspeller”), die op de horizontale as is uitgezet. Het model bestaat uit een regressielijn: een rechte lijn die zo goed mogelijk door de punten is getrokken.

Het “trainen” van het model bestaat in dit geval uit het via het regressie-algoritme zo goed mogelijk “passend” maken van de rechte lijn, waarbij de gemiddelde afstand van de punten tot de lijn zo klein mogelijk wordt gemaakt.



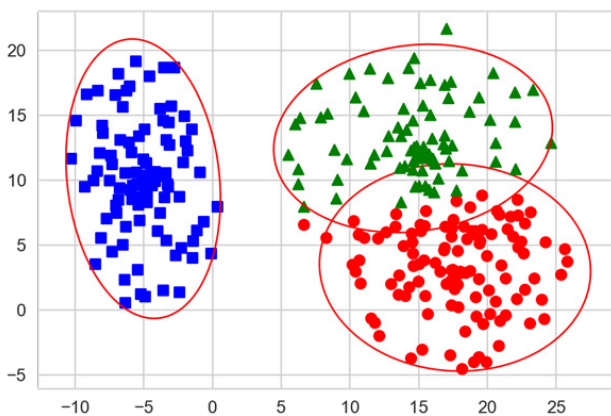
Figuur 9 - Willekeurige gegevenspunten en hun lineaire regressie (bron: https://commons.wikimedia.org/wiki/File:Linear_regression.svg)

In de praktijk is de werkelijkheid natuurlijk complexer en spelen meerdere voorspellers een rol. We spreken dan van meervoudige lineaire regressie. De pilotdataset bevatte veel en verschillende soorten variabelen en complexe correlatiestructuren. Het model dat gebouwd werd gedurende de verschillende iteraties evolueerde naar een **Lineair Regression With Mixed Random Effects Model**. De resultaten met lineaire regressiemodellen zijn voor dit soort data meestal vergelijkbaar met wat de meer moderne technieken met behulp van neurale netwerken en Deep Learning bewerkstelligen: ze zijn minder goed uitlegbaar en neigen meer naar de black-box modellen met lagere transparantie.

¹² CBS-dataset 84799NED: Kerncijfers wijken en buurten 2020

Clusteren

Toen bleek dat de variabelen in de dataset maar een hele beperkte voorspellende waarde hadden voor het inlevermoment, wilden we verder experimenteren om tot betere resultaten te komen. Wat als we voor een meer beperkte groep leners of materialen wel een betere voorspelling kunnen doen? Dan zou het systeem voor bepaalde gevallen een (goede) voorspelling kunnen doen en bij andere gevallen aangeven dat dit helaas niet kan. Een manier om dit aan te vliegen is om te kijken of de data kan worden geclusterd. Door te werken met clusters met vergelijkbare eigenschappen konden we wellicht het inlevermoment beter voorspellen dan voor de set als geheel¹³. Bovendien geeft de clusteringsmethode ons de mogelijkheid om de voorspelling uit te breiden naar groepen leners of materiaalsoorten die we in de dataset nog niet eerder hebben gezien, het eerder genoemde Cold Start-probleem.



Figuur 10 - Voorbeeld van clustering

In Figuur 10 zien we een voorbeeld hoe een dataset wordt ingedeeld in drie clusters. Het blauwe cluster staat op zichzelf, maar het onderscheid tussen groen en rood is minder goed gedefinieerd. Een manier om de kwaliteit van de clustering te meten, is door de totale afstand tot het clustercentrum te berekenen (inter-cluster size) en deze te vergelijken tussen clusters. Hoe lager deze afstand, hoe beter de clustering. Een andere, iets ingewikkeldere, methode is de silhouette-score. Deze score houdt rekening met de kwaliteit van elk punt in een cluster door ook de afstand tot concurrerende clusters mee te nemen.¹⁴ De score varieert van -1 tot 1, waarbij een hogere score betere clustering betekent.

Zoals gezegd, het ontwerpen van een AI-oplossing is een iteratief en exploratief proces. Er wordt steeds aan de dataset en het model gesleuteld om tot betere resultaten te komen. Na evaluatie van de modelresultaten moet je soms terug om andere data te selecteren of de data anders bewerken.

Hierna volgt een beschrijving van de verkenningen en bevindingen van de drie iteraties die tijdens de pilot zijn uitgevoerd.

¹³ In de context van deze clusteringsmethode, is cluster een afgebakende set van datapunten binnen de grotere gegevensset.

¹⁴ [https://en.wikipedia.org/wiki/Silhouette_\(clustering\)](https://en.wikipedia.org/wiki/Silhouette_(clustering))

Eerste iteratie: Zonder persoonsgegevens

Doel: onderzoeken tot hoe nauwkeurig het inlevermoment voorspeld kan worden. Indien nodig en mogelijk: groepen exemplaren of titels herkennen waarvoor dat beter of juist slechter kan.

Dataset: titel-, exemplaar- en transactiegegevens voor uitleenen, verlengen en innemen, met name:

- Titels: mediumsoort, typering, statistische categorie, publicatiejaar, genre, taal, oorspronkelijke taal
- Exemplaren: materiaalsoort, eigenaar, oorsprong, plaatsing, kast, prijs, exemplaar-id
- Transacties: case-id, transactie-id, actie, leendatum, inleverdatum, oorspronkelijke inleverdatum, bibliotheekvestiging, poort, applicatie

Zie [bijlage 2](#) voor de data dictionaire.

Acties om aan de juridische verplichtingen te voldoen

Er worden geen persoonsgegevens gebruikt en de data is geheel anoniem. Dan is de AVG niet van toepassing op deze iteratie.

Aanpak en resultaten

Van elke uitleencasus (uitleen, vervolgens eventueel 1 of meer verlengingen, vervolgens inname) worden sequenties gebouwd. Op de vorm en lengte van de sequenties werd analyse gedaan. Overwegend werd niet-verlengd (90%).

Bij een inlevermoment is zowel de werkelijke datum als de van tevoren uiterste inleverdatum bekend. Als maat/uitkomst van inleveren werd het verschil tussen uiterste inleverdatum en werkelijke inleverdatum genomen in dagen. We noemen deze uitkomst hierna steeds “aantal-dagen-te-laat”: een geheel getal wat (bij voortijdig inleveren) dus ook negatief kan zijn.

In de analyse van aantal-dagen-te-laat werd vooral veel variabiliteit aangetroffen en geen relevante correlaties/verbanden tussen eigenschappen en de uitkomst. Dat wil zeggen dat geen enkele van de gebruikte attributen een uitgesproken sterkere correlatie had met aantal-dagen-te-laat, ofwel geen van de attributen speelt een echte doorslaggevende rol. Dat was ook niet op voorhand te verwachten, maar wel een belangrijke bevinding.

De inlevermomenten liggen ruwweg verspreid rond de inleverdatum, ofwel aantal-dagen-te-laat met waarde 0, met een standaarddeviatie van ongeveer ± 14 . Dit houdt in dat een hele ‘domme’ voorspelling op basis van de uiterste inleverdatum (aantal-dagen-te-laat = 0) een nauwkeurigheid heeft van ± 14 dagen.

Tijdens deze iteratie werd een model gebouwd op basis van *Lineaire Regressie*. Dit leverde een nauwkeurigheid van ± 12 dagen op.¹⁵ Hoewel dit wel een verbetering opleverde, was het resultaat nog lang niet in de buurt van ons voorspellingsdoel (nauwkeurigheid van ± 2 dagen).

Daarna werd een model gebouwd met *Lineair Regression With Mixed Random Effects*. Dit gaf een verbetering tot een nauwkeurigheid van ± 9 dagen.

Er werd besloten geen andere modellen te proberen: de data is dermate variabel dat geen substantiële verbetering mogelijk werd geacht. Het model is weliswaar min-of-meer optimaal voor de bestaande set exemplaren, maar dit model werkt niet goed voor nieuwe exemplaren waarvoor (nog) geen transacties bestaan – het *Cold Start* probleem. Om ook in Cold Start-gevallen een goede voorspelling te kunnen doen, moet een nieuw exemplaar op basis van eigenschappen van titel en exemplaar kunnen worden geclassificeerd in een profiel van exemplaren met betere voorspelnauwkeurigheid / smallere distributie. Er zijn geen eigenschappen gevonden die zo'n profilering ondersteunen.

¹⁵ De stap van 14 naar 12 dagen is uit te leggen aan de hand van het voorbeeld van Lineaire Regressie in Figuur 9: als we in die figuur een horizontale lijn door de punten trekken ter hoogte van de gemiddelde y-waarde dan is de 14 dagen (de standaarddeviatie die hoort bij) de spreiding van de punten rond die horizontale lijn, ofwel een maat voor de afstand van de punten tot die lijn. De 12 dagen hoort dan juist bij de afstand van de punten tot de (slim met Lineaire Regressie) schuin door de punten getrokken lijn. Het is in deze figuur direct duidelijk dat de schuin getrokken lijn een betere ‘voorspelling’ geeft dan een horizontaal getrokken lijn die altijd het gemiddelde ‘voorspelt’, omdat de spreiding rond de schuin-getrokken lijn ‘smaller’ is.

Tweede iteratie: Met persoonsgegevens, eerste verkenning

Doel: onderzoeken tot hoe nauwkeurig het inlevermoment voorspeld kan worden wanneer persoonsgegevens gebruikt worden. Indien nodig en mogelijk: groepen leners herkennen waarvoor dat beter of juist slechter kan. Onderzoeken of clustering mogelijk is.

Dataset: Naast de hierboven genoemde gegevens van titels, exemplaren en transacties werden de volgende persoonsgegevens gebruikt:

- Uit Wise: pseudo-id, adres-id, geboortjaar (afgerond), gender, statistische categorie, abonnementsnummer, thuisvestiging, inschrijfvestiging, boetebedrag over de laatste 3 maanden voor de transactie, verstuurd notificaties
- Uit CBS (op basis van postcode): vermogen per huishouden, afstand tot de bibliotheek

Acties om aan de juridische verplichtingen te voldoen:

De persoonsgegevens werden gepseudonimiseerd en er werd behouden omgegaan met identificerende combinaties. Zo werd in eerste instantie geboortedatum afgerond op 10 jaar (decennium) en werden locatiegerelateerde gegevens (adres, postcode, wijk/buurt gegevens van CBS over inkomens en afstand tot de bibliotheek) afgerond of in grovere “buckets”¹⁶ gecategoriseerd. De gemaakte keuzes werden op intuïtie gedaan, met als grondslag de tellingen c.q. bepaling van het heridentificatie risico van de (afgeronde) combinaties van (geboortedatum, adres).

Zie bijlage 3 voor de verantwoording.

Aanpak en resultaten

Hetzelfde Linear Regression With Mixed Random Effects-algoritme (maar nu met de basisdataset uitgebreid met de persoonsgegevens) werd gebruikt voor het bouwen van een model. Door de toegevoegde persoonsgegevens werd de bereikte nauwkeurigheid wel iets beter dan zonder persoonsgegevens: +/- 7 dagen in plaats van +/- 9 dagen. Dit resultaat was echter nog niet genoeg om bruikbaar te zijn in de praktijk.

De grote variabiliteit in de huidige dataset maakt het onmogelijk om een algemeen model te bouwen dat de vereiste nauwkeurigheid van +/- 2 dagen bereikt.

Vervolgens is met deze dataset gekeken of de leners geclusterd konden worden op basis van de verkregen attributen. Het idee was dat deze clusters dan mogelijk specifiek inlevergedrag vertonen en dat het bepalen van een cluster kan helpen voor het categoriseren van nieuwe leners.

Voor clustering werd een standaard **k-Means clustering algoritme**¹⁷ gebruikt met een variërend aantal clusters. De volgende attributen werden gebruikt:

- Gender, vermogen, afstand-tot-de-bieb, boete, leeftijd (decennium), aantal bewoners op adres

De kwaliteit van de clusters werd gecheckt op basis van de zogenoemde **silhouette-score**. De conclusie was dat deze clustering niet goed bruikbaar was omdat er geen cluster-aantal gevonden werd met voldoende kwaliteit.

¹⁶ In de context van deze de-identificatie methode is een bucket een verzameling bron-waarden die gezamenlijk op 1 en dezelfde doel-waarde worden afgebeeld. Een voorbeeld is het gebruik van “decennia” in plaats van geboortejaren. Hierbij is sprake van buckets van telkens 10 geboortejaren (bijv. 1990-1999, de bron-waarden) behorend bij 1 decennium (1990 in dit voorbeeld, de doelwaarde).

¹⁷ https://en.wikipedia.org/wiki/K-means_clustering

Tenslotte is nog voor een aantal eigenschappen gekeken in hoeverre deze onderscheidend zijn op “te laat” versus “op tijd”, wat aanleiding zou kunnen zijn om nader te onderscheiden op basis van deze eigenschappen (bijvoorbeeld als blijkt dat een bepaalde categorie leners een heel consequent en specifiek inlevergedrag heeft). Zo is met name gekeken naar:

- Vermogen in relatie tot te-laet gedrag
- Boete-opbouw ten opzichte van te-laet gedrag

Deze bleken beide niet of nauwelijks te correleren. Voor boete-opbouw is dat merkwaardig, omdat er een causaal verband is tussen te-laet inleveren en het oplopen van boete. Verwacht werd een reeks pieken bij boete-bedragen >0 voor te laat inleveren, maar deze werden niet of nauwelijks aangetroffen. We besloten onze vraagtekens hierbij te noteren voor eventueel verder onderzoek.

Kortom, tijdens deze tweede iteratie leverde het toevoegen van persoonsgegevens tot de dataset nog onvoldoende resultaat. Ook het slim clusteren van leners volgens hun inlevergedrag, of het zoeken naar onderscheidende eigenschappen van leners bij te-laet inlevergedrag, leverde niets op.

Derde iteratie: Met persoonsgegevens, tweede verkenning

Doel: Onderzoeken wat de invloed is van meer details in de persoonsgegevens. Het verkennen van enkele nieuwe gegevens die mogelijk van invloed zijn op het inlevergedrag en dan gebruikt zouden kunnen worden bij clustering:

Te bekijken eigenschappen:

- Wel/niet inleverbericht gehad voor inleveren: heeft een bericht invloed?

- Aantal geleende exemplaren per jaar: is er verschil in inlevergedrag tussen incidentele leners en grootverbruikers?
- Leeftijdscategorie (nu met een meer gerichte indeling)
- Gezinsgedrag: in hoeverre wordt per gezin geleend en ingeleverd?
- Relatie tussen aantal keren verlengen en relatief te laat zijn

Dataset: In de discussie na de eerste verkenning met persoonsgegevens rees de vraag of de gekozen aanpak niet te veel detailinformatie heeft weggehaald om nog relevante data over te houden voor clustering. Men vond met name dat de leeftjidsindeling op decennium nogal “bot” afgekapt werd in grofmazige categorieën. De suggestie was om met iets meer details te werken om te kunnen bepalen welke details relevant zijn en welke niet. Zo werd door de pilotgroep besloten om te werken met geboortjaar in plaats van decennium.

Op grond van het besluit om met meer details te werken is een aantal aanpassingen gedaan aan de software (queries) die gebruikt werd om de data voor te bereiden op het ML proces:

- Gebruik geboortjaar in plaats van decennium
- Boeteberekening over kortere periode, namelijk 3 maanden in plaats van een vol jaar

Acties om aan de juridische verplichtingen te voldoen

Hoewel we leeftijdsgegevens op een gedetailleerder niveau hebben gebruikt en fijnere categorieën hebben gedefinieerd, was de kans dat de gegevens herleid konden worden naar individuen zeer klein en was er geen (juridisch) risico voor de pilot.

Aanpak en resultaten

Op basis van de aangepaste dataset is allereerst gekeken of het model voor Lineaire Regressie nog een verbeterd resultaat gaf. Dat bleek niet zo te zijn.

Het toevoegen van meer detail in de persoonsgegevens gaf geen verbetering van de nauwkeurigheid, deze bleef steken op +/- 7 dagen.

Daarna is gekeken naar hoe verschillende eigenschappen (mogelijk) van invloed zijn op inlevergedrag. Dit is een verkenning die nodig is om te bepalen of een verbeterde voorspelling mogelijk is voor een bepaalde categorie (gebruikers, uitleningen).

Bekeken eigenschappen en bevindingen:

Wel/niet inleverbericht gehad vóór inleveren

Uit de analyse blijkt een paradoxaal effect: leners die een bericht ontvangen vóór het inleveren blijken relatief vaker te laat in te leveren. We twijfelen aan de juistheid van deze bevinding en moeten eigenlijk terug naar het Data Science team om te kijken of wel naar het juiste berichttype (IV1, inleverbericht, verstuurd vóór de reglementaire inleverdatum) is gekeken en niet ook naar de herinneringsberichten (HRX, herinneringen, verstuurd ná de reglementaire inleverdatum).

Aantal geleende exemplaren per jaar

Vraag: Is er verschil in inlevergedrag tussen incidentele leners en grootverbruikers?

Er werd een simpele categorie ingevoerd waar leners worden ingedeeld naar het gemiddelde aantal leningen per jaar: klein, middel en groot.

Het blijkt dat er wel verschil is in leentermijn per categorie (grootverbruikers lenen gemiddeld korter, wat ook wel moet om aan het grootverbruik te komen), maar een verband met te laat inleveren (aantal-dagen-te-laet) is niet duidelijk geworden.

Leeftijdscategorie

(nu met een meer gerichte indeling)

Gekeken is met name naar de leentermijn zonder verlengingen. Langere uitleentermijnen worden vooral benut door jongeren. De relatie met aantal-dagen-te-laet is niet onderzocht.

Gezinsgedrag: in hoeverre wordt per gezin geleend en ingeleverd?

Doel was om vast te stellen of conclusies kunnen worden getrokken over het groepsgewijs inleveren, bijvoorbeeld dat men eerder op tijd inlevert als men in gezinsverband leent. Dit is nog niet onderzocht. Wel werd gekeken naar of, en zo ja, hoeveel er gezamenlijk wordt uitgeleend en ingeleverd. Het blijkt dat uitleningen in 1 op de 6 keer gezamenlijk gebeuren, terwijl innames vaker, namelijk 1 op de 5 keer, gezamenlijk plaatsvinden. Dit is mogelijk te verklaren doordat het makkelijker is om voor een gezinslid in te leveren dan om te kiezen welk boek meegenomen moet worden. De aantallen zijn niet heel groot en het is nog niet duidelijk of hiermee iets substantieels kan worden gedaan in het kader van het voorspellen van het inlevermoment.

Relatie tussen aantal keren verlengen en relatief te laat zijn

Ook hier is alleen gekeken naar de absolute tijdlijn (vanaf het moment van uitlenen/vorige verlenging) en niet naar hoe de verlenging is gesitueerd ten opzichte van uiterste inleverdatum. Uit de verdeling blijkt echter wel dat de eerste verlenging relatief vaak te laat gebeurt. Er is namelijk een sterke piek juist voorbij de drie weken. Het is te vroeg om conclusies te trekken over de bruikbaarheid, maar wel aanleiding om nog eens beter te kijken, vooral naar de relatie tussen het moment van verlengen en de uiterste inleverdatum.

In zijn algemeenheid kunnen we stellen dat de data-exploraties complexer en tijdrovender waren dan voorzien.

Het bleek lastig om concepten als 'absolute leentermijn' (inclusief verlengingen) en 'relatieve leentermijn' (aantal-dagen-te-laet) in de discussie helder te houden.

De complexiteit van het vraagstuk rondom het voorspellen van de retourdatum, houdt niet alleen verband met individueel lenersgedrag maar ook met individueel bibliotheekbeleid en regels. Neem bijvoorbeeld het verlenggedrag van leners: een verlenging verschuift de leentermijn met enkele weken op en is dus veel belangrijker voor het inlevermoment dan een dag te-laat inleveren. Maar ook het verlengingsbeleid speelt een rol: sommige openbare bibliotheken blokkeren een verlenging voor een exemplaar waar een reservering op staat, en anderen doen dat niet.

Een aantal punten zijn blijven liggen en zouden in een eventueel vervolgonderzoek kunnen worden opgepakt:

- Validatie bevinding inleverbericht en te-laat zijn;
- Relatie tussen hoeveelheid geleende exemplaren per jaar (in gebruikerscategorieën) en verlenggedrag + aantal-dagen-te-laat;
- Idem voor leeftijd;
- Verdere uitwerking van gezins-leengedrag: wordt vaker op tijd ingeleverd als men gezamenlijk inlevert?;
- Verlenggedrag in relatie tot de reglementaire inleverdatum bij eerste en volgende verlengingen;
- Hoe goed kunnen we (voor bepaalde leners) voorspellen of er verlengd gaat worden?

Conclusie technisch aspecten

De doelstelling van een nauwkeurigheid van +/- 2 dagen in de voorspelling van het inlevermoment werd niet gehaald. Er kon op grond van de beschikbare gegevens geen deelcollectie of deelpopulatie worden vastgesteld waarvoor het inlevermoment met een betere nauwkeurigheid dan +/- 7 dagen kan worden bepaald. Hiermee hebben we geen bruikbare oplossing verkregen om het inlevermoment te voorspellen.

Toen we concludeerden dat het nog veel tijd en inspanning zou kosten om tot een doorbraak te komen met onze voorspellingspogingen op basis van AI, besloten we op de pauzeknop te drukken en te evalueren wat we tot dusver van de data-exploraties hadden geleerd. Een voorwaarde van de pilot was dat gegevens alleen worden gebruikt als die in verhouding staan tot het resultaat. Het toevoegen van meer details in de persoonsgegevens heeft hier niet toe geleid. In overleg met de pilotbibliotheken is daarom bij de evaluatie van de derde iteratie besloten de pilot te beëindigen.

Pseudonimiseren of deïdentificeren is een vak apart. De kunst is de juiste balans te vinden tussen onherkenbaarheid – hiding in the crowd - en voldoende specificiteit om relevante patronen te kunnen zien: de populatie moet niet een grijze massa worden. Dit vergt nogal wat voorwerk en uitprobeersels.

Er zijn heel veel variabelen die potentieel het inlevergedrag van de lener kunnen verklaren, maar welke zijn kritisch? En welke brengen alleen maar meer complexiteit aan de dataset zonder enige bijdrage aan de voorspelling te leveren?

Bovendien heeft niet elke variabele uitgebreide en schone historische gegevens met resultaten voor de hele populatie leners. Denk bijvoorbeeld aan variabelen die per bibliotheek verschillen, zoals het boetebedrag, verlengkosten of het maximumaantal verlengingen.

Er werden te weinig werkbare patronen ontdekt voor clustering of voor een goede voorspelling voor het inlevermoment. Dit leidde tot het inzicht dat het bouwen van een goede dataset met bibliotheekgegevens voor AI lastig blijkt te zijn, omdat de gegevens vaak niet consistent, compleet of vergelijkbaar zijn.

Er is veel variabiliteit in de data die niet op grond van de beschikbare gegevens is te verklaren. Intuïtief zeggen we dat de belangrijkste component in het inlevermoment de beslissing van de lener is om te verlengen en deze beslissing is vermoedelijk vaak afhankelijk van factoren die niet in de data zijn verdisconteerd. Een vervolgonderzoek kan zich mogelijk richten op het in kaart brengen van deze (externe) factoren en hun bijdrage afschatten.

Aangezien de meeste uitleningen gewoon regulier zijn - zonder verlenging – en er een grote categorie leners is die nooit verlengen, is het denkbaar dat er een rule-based systeem wordt ontworpen wat in veel gevallen met grote zekerheid kan zeggen dat er niet verlengd zal worden en op tijd ingeleverd. Dit zal zich dan baseren op de leengeschiedenis van de lener die het materiaal in bezit heeft, en geen uitspraak kunnen doen als er geen of niet voldoende leengeschiedenis is. Zo'n oplossing kan dan ofwel een (nauwkeurige) voorspelling doen, ofwel niet, maar wellicht wel goed aangeven hoe groot het vertrouwen is (de waarschijnlijkheid) als een voorspelling wordt gedaan. Verder onderzoek zou moeten uitwijzen of dit in een substantieel aantal gevallen voldoende nauwkeurig voorspelt om werkbaar te zijn bij het aangeven van de levertijd van een reservering.

Bevindingen en aanbevelingen

Conclusie

Probiblio, OCLC en de openbare bibliotheken van Kennemerwaard, Bollenstreek en Krimpenerwaard, hebben onderzocht of en hoe AI kan helpen met het nauwkeurig voorspellen van de retourdatum van uitgeleend bibliotheekmateriaal. Zij hebben in 2022 een AI-pilot uitgevoerd om de juridische, ethische en technische aspecten van deze casus onder de loep te nemen. De eindconclusie is dat er nog geen AI-oplossing is gevonden die het inlevermoment nauwkeurig genoeg (+/- 2 dagen) kan voorspellen. De voorspelling op basis van materiaal- en uitleentransactiegegevens leverde wel een verbetering op ten opzichte van de nulmeting (+/- 14 dagen), namelijk +/- 9 dagen. Het inzetten van gepseudonimiseerde persoonsgegevens van leners en clusteringsmethoden, verbeterde de nauwkeurigheid van de voorspelling naar +/- 7 dagen.

Omdat de resultaten er niet noemenswaardig beter op werden door toevoeging van persoonsgegevens, besloten de pilotdeelnemers hier niet mee verder te gaan. Dit was meer een ethische overweging dan een juridische. De pilot is ruim binnen de juridische kaders gebleven van de AVG, het Europese AI-wetsvoorstel en de verwerkingsovereenkomsten van de pilotpartijen.

Nieuwe inzichten

De juridische verkenning was leerzaam. Dankzij de juridische expertise in het team konden de pilotdeelnemers met kennis van zaken aan de slag gaan. De regulering van het gebruik van persoonsgegevens is met name behoorlijk complex. De specifieke juridische invulling van begrippen zoals “profilering” zijn lastig te doorgronden. ***Het is niet voldoende om advies te vragen aan een juridisch adviseur, je moet als teamlid en ongeacht je rol (projectleider, AI-ontwikkelaar, data-expert, of bibliotheekmedewerker), de AVG begrijpen en naar de geest van de wet handelen.*** De ethische verkenning was een eye-opener. Het

dwong de teamleden langer stil te blijven staan bij de probleemdefinitie van de casus en te kijken naar alternatieve oplossingsrichtingen. Met een ethicus/kritische meedenker aan boord van het team, kun je het proces beter begeleiden, bijsturen en documenteren. Je kunt aannames, waarden en keuzes expliciteren en bespreekbaar maken. Er is meer aandacht voor validatie en dialoog met de stakeholders.

Ethische richtlijnen en principes zijn vaak erg algemeen geformuleerd en gebaseerd op grondrechten waardoor ze moeilijk te concretiseren zijn voor een specifieke casus. Werken met domein-specifieke waarden die relevant zijn voor de casus, geeft meer houvast en richting aan de bespreking van ethische kwesties. ***Voor de ethische verkenning van deze casus was het belangrijk nauw aan te sluiten bij de waarden van de bibliotheek. Uit deze oefening kwam naar voren welke waarden relevant zijn voor de casus: de lener centraal, de betrouwbare bibliotheek en solidariteit.***

De technische verkenning was misschien minder bevredigend dan gehoopt, maar leverde wel inzichten op. Het bijeenbrengen van de uitleengegegevens van de verschillende deelnemende bibliotheken resulteerde in een dataset die onvoldoende consistent en compleet was. De gegevens waren niet altijd vergelijkbaar, wat voor een AI-toepassing lastig werken is.

Het aanvullen en opschonen van de dataset, het anonimiseren, danwel pseudonimiseren en de-identificeren van de persoonsgegevens, vergden veel voorwerk tijdens de iteratieve data-exploraties. De evaluatie van de data-exploraties bleken ook tijdrovender dan voorzien. Bij elke iteratie werd opnieuw nagedacht over de gegevens en de methode, en de juridische en ethische consequenties. Kortom, ***een pilot of experiment klinkt als een kortdurende inspanning, maar dat gaat niet op als je met AI wilt experimenteren en uitzoeken of het wel of niet een oplossingsrichting biedt voor een bepaald knelpunt.***

Tijdens de ontwikkeling van de algoritmische modellen en de data-exploraties, zijn een aantal punten en vragen blijven liggen, als gevolg van tijdgebrek. Deze zouden in een vervolgonderzoek opgepakt kunnen worden. **Het inzicht is echter gegroeid dat er waarschijnlijk beter resultaat geboekt kan worden door data-analyses te doen van processen of deelverzamelingen waar men meer grip op heeft.** Denk bijvoorbeeld aan:

- De grote categorie leners die nooit verlengt en de grote categorie uitleningen die op tijd ingeleverd worden, waarvoor een voorspellingsmodel ontworpen zou kunnen worden dat in veel gevallen met grote zekerheid kan zeggen dat er niet verlengd zal worden en op tijd ingeleverd.
- Het in kaart brengen van het inlevergedrag per bibliotheek, waar eenzelfde set uitleenregels geldt en waardoor de uitleengegevens consistent zijn;
- De achterliggende logistieke processen en de administratie rondom uitleenhandelingen – processen die zeer complex zijn, maar die wel onder controle staan van de bibliotheek. De bibliotheken kunnen in principe hier gemakkelijker verbeteringen in aanbrengen.

Aanbevelingen

Het gebruik van data wordt steeds belangrijker voor bibliotheken, om werkprocessen te verbeteren, beleidskeuzes te informeren en de dienstverlening steeds aan te passen aan de veranderende

verwachtingen van gebruikers. Daarbij is het van belang dat bibliotheken en hun partners, een datavisie ontwikkelen rond het gebruik van data en de bijkomende risico's. **De theorie en praktijk van juridisch-ethisch verantwoord werken met bibliotheekdata behoort een vast onderdeel te zijn van opleidingstrajecten, jaarlijkse trainingsprogramma's en werkbesprekingen.**

Het is nuttig om te onderzoeken hoe bibliotheekwaarden uitpakken in een veranderende omgeving. Het werk van bibliotheekmedewerkers en het gedrag van leners veranderen doordat de dienstverlening meer online plaatsvindt. Er is minder contact met leners over reserveringen, en meer behoefte om de digitale processen uit te leggen. De verwachtingen van leners zijn meer gedreven door gemak en sneller voldoening krijgen. **Er is meer onderzoek nodig naar de vraag: wat betekenen de klassieke bibliotheekwaarden in een online omgeving en hoe gaan we daarmee om?**

Er komen steeds meer AI-Labs die domeinspecifiek en langdurig onderzoek doen naar en experimenteren met de toepassing van (schaalbare) machine learning benaderingen.¹⁸ **Voor bibliotheken zou een gezamenlijke AI-Lab omgeving waarin de data over hun collecties en gebruiksdata bijeen gebracht is, een effectievere innovatie instrument kunnen zijn dan het uitvoeren van afzonderlijke, kleinschalige AI-projecten.** In het Lab kun je de kennis en expertise concentreren die nodig zijn, zoals domeinkennis van de data, juridische en ethische expertise en data science vaardigheden.

¹⁸ Zie bijvoorbeeld in Nederland: het Cultural AI Lab (<https://www.cultural-ai.nl/>) of het AI, Dmedia and Democracy Lab (<https://www.aim4dem.nl/>).



Dankbetuigingen

Graag willen de auteurs van het rapport Vincent Jordaan bedanken voor zijn uitgebreide review van het manuscript. Zijn frisse blik en suggesties voor verbetering hebben het eindrapport net een slag consistenter en leesbaarder gemaakt. Het rapport kon worden gepubliceerd dankzij de inzet en professionaliteit van het communicatieteam van OCLC EMEA, en de auteurs willen met name Gemma Burke hiervoor bedanken.

Bijlage 1: Lijst met gebruikte afkortingen

AI	Artificial Intelligence
ALTAI	The Assessment List for Trustworthy Artificial Intelligence for self-assessment
AVG	Algemene Verordening Gegevensbescherming
BZK	Ministerie van Binnenlandse Zaken
CBS	Centraal Bureau voor de Statistiek
COVID-19	Coronavirus
DEDA	De Ethische Data Assistant (DEDA) van Utrecht Data School
DL	Deep Learning
EC	Europese Commissie
EU	Europese Unie
KB	Koninklijke Bibliotheek
ML	Machine learning
NLAIC	Nederlandse AI Coalitie
POI	Provinciale ondersteuningsinstelling

Bijlage 2: Gebruikte data, data dictionary

Tabel met data-elementen gebruikt in fase 1 (alleen titel-, exemplaar- en transactiegegevens)

case_id	id that binds the txs of the case together, a case is defined by (item, patron) combination and normally starts with tx_type 1 = checkout and ends with tx_type 2 = checkin with N>=0 renewals in between - the case id is defined by the txs_id of the first txs
txs_id	id of the transaction
tx_type	type of the transaction, 1= checkout, 2= checkin, 3 = renewal - there are others but they have been suppressed
tx_date	date of the transaction in ISO notation
checkout_date	date of the checkout in ISO notation, a constant for all txs in the case.
due_date	after this transaction: when is the item due? this changes with renewals
orig_due_date	after a renewal, what was the original due date?
tx_branch	branch_id of the branch where the tx took place
tx_library	library id of the library the tx_branch belongs to
placement_branch	branch_id of where the item is normally on the shelf
placement_library	library id of the library the placement_branch belongs to
owner_branch	the owner is the original acquirer of the item (who does it belong to legally?)
owner_library	the owner is the original acquirer of the item (who does it belong to legally?)
tx_terminal	id of the terminal (SIP2 station, staff desk pc, web-page) the tx took place at
tx_client	id of the application, I=web page, J=App, B=staff desk client, S=SIP2 machine
item_id	
title_id	
medium	
youth	Y / N
adult	Y always in this file
nonfiction	Y / N
fiction	Y always in this file
stat_cat_title	statistical category, probably to fine-grained to work with, 4 digits, based on a mapping that has some of the other categories as input, so overlaps

publ_year	year of publication
publ_year_end	last year of publication (in case of multiple prints/editions in the same title)
genre_codes	at most 3 genre-codes concatenated, genre-codes are 2 or 3 digits, they are PADRed to 3 pos in this concat.
language_codes	at most 3 language codes of languages used in the material, concatenated; a language code has 3 chars so no spaces here
orig_lang_codes	at most 3 language codes of languages in the original print (before translation), idem
item_material_type	
item_shelf	shelf codes where the item is expected to reside when checked back in (placement branch)
item_placement	where would the item be if not temporarily on a different item_shelf
price	price of the item in Euro cents - based on ACQuisition information
collation	"Raw" collation information, contains page count in several forms (still to be extracted), like "300 pagina's ; 22 cm" => 300, or "XVIII, 515 blz. ; 4" => 515

Tabel met persoonsgegevens gebruikt in fase 2:

pseudo_id	unique hash for this actor, not used anywhere else
adres_id	hashed adres_id for this actor; can be used to find actors that form a household - might work for 1<N<10
decenium	year-of-birth rounded to a decenium
gender	gender, M = male, V = female, I = institution
distance_to_library	distance from home address (area/neighbourhood) binned into bins labelled with 500m - folds
capital	(monetary value, scaled to 50k-folds, binned) average capital per household in area/neighbourhood
statistical_category	statistical category of the subscription in Wise
registration_type	registration type in Wise
home_branch	home_branch in Wise
registration_branch	registration_branch in Wise
fines_in_cents	amount (in euro cents) of fines over the last full 4 Q for this actor

Bijlage 3: Tellingen unieke combinaties en buckets voor CBS-data

Tellingen unieke combinaties

Hoe groot is de kans dat 1 record herleid wordt? Hoe groot is de gemiddelde kans dat records herleid worden? Hoe lager het aantal unieke combinaties, hoe kleiner de kans op her-identificatie (“hiding in the crowd”).

Onderstaande tellingen werden gedaan op het ledenbestand van de drie participerende bibliotheken.

Wat?	Jaar + postcode (1974,1234 AB)	Jaar + 4 cijfers pc (1974,1234)	Decade + 4 cijfers pc (1970,1234)
Totaal actoren	149978	149978	149978
Waarvan unieke (p,j)	93294	3411	1155
Beneden CBS-norm (5)	134098	11060	2781

Buckets voor de afstand tot de bibliotheek

Omdat de precieze afstand (zelfs gemiddeld op buurtniveau) een indirect identificerend gegeven kan zijn werd de afstand tot de bibliotheek afgerond in zogenaamde buckets, volgens onderstaand gevalsonderscheid:

```
update _ai_cbs_afstand_tot_bieb
  set val := case
    when AfstandTotBibliotheek_92 between 0 and 0.5 then 0
    when AfstandTotBibliotheek_92 between 0.5 and 1 then 1
    when AfstandTotBibliotheek_92 between 1 and 1.5 then 2
    when AfstandTotBibliotheek_92 between 1.5 and 2.5 then 4
    when AfstandTotBibliotheek_92 between 2.5 and 3.5 then 6
    when AfstandTotBibliotheek_92 between 3.5 and 4.5 then 8
    when AfstandTotBibliotheek_92 between 4.5 and 5.5 then 10
    when AfstandTotBibliotheek_92 between 5.5 and 8.5 then 14
    when AfstandTotBibliotheek_92 between 8.5 and 11.5 then 20
    else 30
  end
where AfstandTotBibliotheek_92 > -1
;
```


Buckets voor vermogen per huishouden

Gemiddeld vermogen per huishouden is weliswaar praktisch moeilijk te gebruiken als (indirect) identificerend gegeven, toch werd ook hier gebruik gemaakt van buckets:

```
update _ai_cbs_vermogen_per_wijkbuurt
  set val := case
    when MediaanVermogenVanParticuliereHuish_82 between -50.0 and 0.0 then -1
    when MediaanVermogenVanParticuliereHuish_82 between 0.0 and 50.0 then 1
    when MediaanVermogenVanParticuliereHuish_82 between 50.0 and 100.0 then 2
    when MediaanVermogenVanParticuliereHuish_82 between 100.0 and 150.0 then 3
    when MediaanVermogenVanParticuliereHuish_82 between 150.0 and 200.0 then 4
    when MediaanVermogenVanParticuliereHuish_82 between 200.0 and 300.0 then 5
    when MediaanVermogenVanParticuliereHuish_82 between 300.0 and 500.0 then 8
    else 15
  end
where MediaanVermogenVanParticuliereHuish_82 > -50
;
```

Bijlage 4: Relevante functionaliteit en data beschikbaar in OCLC Wise

Het bibliotheekstelsel Wise wordt bij een meerderheid van de Nederlandse openbare bibliotheken gebruikt voor de uitleening van fysieke materialen en alle processen die hier (in de loop van de lange historie van openbare bibliotheken) omheen zijn georganiseerd: bestellen/collectiebeheer, reserveren en interbibliothecair leenverkeer, lenersadministratie, beheer van financiën betreffende uitleening (te laat- uitleen- en verlenggeld en tegoeden) en de abonnementen, de bijbehorende berichtenstromen, rapportages en verslagen, et cetera. Marketing en kaartverkoop voor evenementen zijn optionele modules die bij een deel van de gebruikers in gebruik zijn.

Het Wise-systeem staat niet los en gesloten in de wereld, maar is via diverse interfaces gekoppeld aan andere systemen in de keten van bibliotheekprocessen. Denk hierbij aan de boekenleverancier, een centraal zoekplatform, uitleenautomaten en eindgebruikerapplicaties (apps).

Bij al deze processen is data betrokken van een of meer van de relevante entiteiten: titels, exemplaren, leners. Processen bestaan uit transacties op deze entiteiten en die transacties worden in het algemeen in het bibliotheekstelsel geregistreerd. Naast transacties zijn per entiteit een aantal (vaste) eigenschappen geregistreerd die voor de diverse processen van belang zijn.

Op titel- en exemplaarniveau worden de aantal indelingen vastgelegd die collectiemanagement, uitleenadministratie en zoeken (catalogus) ondersteunen. Verder de transacties van het aanschaffen en afschrijven van materialen en het onderhoud/wijzigingen op de indelingen (vaste rubrieken).

Op lenersniveau geldt eigenlijk hetzelfde. Alleen is hier sprake van informatie die onder privacywetgeving valt. Vaste gegevens zijn onder meer naam en adres, transacties zijn onder meer het inschrijven, verlengen en opzeggen van abonnement, verhuizen en wijzigen van filiaal.

Dan zijn er nog de transactionele gegevens in de uitleenadministratie en de financiële administratie: het reserveren, het uitleenen, het verlengen en het inleveren van materialen; het openen en sluiten van financiële posten.

Voor een concreet overzicht van de relevante rubrieken zie de Data Dictionary in bijlage 2.