



Datalaboratorium Draagkracht en Draaglast

Een data- en modelmatige verkenning naar
factoren die van invloed zijn op de
tevredenheid met de woonomgeving met
CBS-microdata

Q3-2025 (CONCEPT)

SAMENVATTING

Dit rapport beschrijft een experimentele verkenning naar de voorspelbaarheid van tevredenheid met de woonomgeving in wijken op basis van registratiedata en enquêtedata. Doel was om te onderzoeken in hoeverre data en modellen inzicht kunnen geven in de balans tussen draagkracht (hulpbronnen, sterktes) en draaglast (druk, problematiek), en welke factoren daarbij bepalend zijn.

Voor dit onderzoek is gebruikgemaakt van CBS-microdata. Uit Woonbase zijn registratiedata betrokken over woningen, inkomens en huishoudens. Daarnaast zijn gegevens uit de WoON-enquête toegevoegd, die de percepties en ervaringen van bewoners in beeld brengen. Door deze combinatie konden zowel objectieve kenmerken als subjectieve beleving worden meegenomen in de analyses.

In Analyse 1 is gewerkt met registratiedata uit Woonbase. Deze gegevens bleken slechts beperkt in staat om tevredenheid met de woonomgeving te voorspellen. De modellen bereikten een nauwkeurigheid van circa 53%, waarbij vooral de grote groep tevreden bewoners werd herkend, maar ontevreden bewoners nauwelijks.

In Analyse 2 zijn de registratiedata uitgebreid met variabelen uit de WoON-enquête. Daarmee kwam ook de beleving van bewoners in beeld, zoals tevredenheid met woning en buurt, ervaringen, sociale cohesie en verhuishwensen. Deze uitbreiding leidde tot een duidelijke verbetering van de voorspelkracht: de nauwkeurigheid steeg naar bijna 70%. Belangrijkste voorspellers bleken variabelen over ervaren leefbaarheid, tevredenheid met kenmerken in de buurt en verhuisintenties, naast enkele woning- en inkomenskenmerken. Tegelijk blijft het probleem van klasse-onbalans bestaan: ontevreden bewoners worden nog steeds minder goed herkend.

In Analyse 3 is aanvullend gekeken naar de rol van buurtkenmerken. Deze contextvariabelen voegen stabiliteit toe aan het model, maar veranderen de uitkomsten niet wezenlijk. De belangrijkste voorspellers blijven belevingsvariabelen, aangevuld met woningkenmerken.

De drie analyses laten samen zien dat registratiedata alleen onvoldoende zijn om tevredenheid met de woonomgeving te verklaren. Het combineren van objectieve kenmerken met subjectieve beleving biedt duidelijk meerwaarde, terwijl de toevoeging van contextvariabelen dit beeld bevestigt maar niet fundamenteel wijzigt.

INHOUDSOPGAVE

SAMENVATTING	2
INHOUDSOPGAVE	3
1. INLEIDING	4
2. ACHTERGROND	5
3. AANPAK	7
4. MODELLEN	9
5. DATA	15
6. RESULTATEN	17
6.1 ANALYSE 1: REGISTRATIEDATA ALS BASIS	18
6.2 ANALYSE 2: TOEVOEGING SUBJECTIEVE GEGEVENS	22
6.3 ANALYSE 3: TOEVOEGING VAN BUURTKENMERKEN	27
7. INZICHTEN	29
COLOFON	33

1. INLEIDING

De leefbaarheid van buurten, wijken en dorpen staat in toenemende mate onder druk. Beleid op het gebied van wonen, zorg en welzijn wordt vaak gemaakt met beperkte middelen, terwijl de vraag naar passende oplossingen groeit. Vooral de concentratie van kwetsbare groepen in specifieke gebieden, de krapte op de woningmarkt en de veranderende sociaaleconomische context zetten druk op de balans tussen draagkracht en draaglast. Het is voor gemeenten, woningcorporaties en zorgorganisaties steeds lastiger om keuzes te onderbouwen: op basis waarvan wordt besloten waar interventies het hardst nodig zijn?

In de praktijk bestaan er al veel methoden om leefbaarheid te meten. Overheden, corporaties en lokale organisaties werken met eigen data en indicatoren die niet altijd onderling vergelijkbaar zijn. Daarbij is de ervaring van professionals vaak leidend, terwijl de beleving van bewoners slechts beperkt in beeld komt (of vice versa). Het resultaat is een lappendeken aan informatie die moeilijk samen te brengen is in een eenduidig beeld van de staat van leefbaarheid.

Het project 'Draagkracht en Draaglast van buurten, wijken en dorpen', met initiatiefnemers Huurdersbelang Fryslân en Planbureau Fryslân, wil hierin verandering brengen. Het doel is het ontwikkelen van een uniform meetinstrument dat inzicht geeft in de balans tussen draagkracht en draaglast op het laagst mogelijke schaalniveau, met oog voor zowel objectieve registratiedata als de subjectieve beleving van inwoners evenals de inzichten van professionals. Dit instrument moet bijdragen aan een gemeenschappelijke taal en vergelijkbare analyses die beleidsmakers, onderzoekers en uitvoerders kunnen gebruiken bij vraagstukken rond leefbaarheid en huisvesting.

Binnen het bredere Draagkracht en Draaglast project heeft DataFryslân een specifieke rol: het ontwikkelen van een kwantitatieve basis met behulp van registratiedata en voorspelmodellen. Met behulp van CBS-Microdata en aanvullende datasets worden experimentele analyses uitgevoerd die inzicht geven in de factoren die leefbaarheid beïnvloeden, en de mate waarin deze factoren voorspellingen mogelijk maken. De inzichten van de data-analyses kunnen gebruikt worden in de verdere uitwerking van het project Draagkracht en Draaglast. Dit rapport beschrijft de aanpak, modellen, gebruikte data en voorlopige resultaten van deze datagedreven en modelmatige analyses.

In een experimentele setting worden modelmatige analyses uitgevoerd, gericht op het inzichtelijk maken van de bijdrage van afzonderlijke kenmerken van inwoners aan de tevredenheid met de woonomgeving. Daarbij wordt niet alleen het model zelf ontwikkeld, maar ook kritisch gekeken naar de exacte methode en de zeggingskracht van de gebruikte analysetechnieken. Het resultaat is een kwantitatief basismodel dat dient als fundament voor de verdere uitwerking van deze pijler in toekomstige pilotstudies.

2. ACHTERGROND

Het begrip draagkracht en draaglast wordt in dit project gebruikt om de balans in buurten en wijken te duiden. Draagkracht verwijst naar de aanwezige middelen en sterke punten van een gebied, zoals sociaaleconomische positie, woningkwaliteit en sociale cohesie. Draaglast staat voor de druk die op een gebied wordt uitgeoefend, bijvoorbeeld door kwetsbaarheid van bewoners, hoge zorgvraag of druk op de woningmarkt. Leefbaarheid kan worden gezien als de mate waarin draagkracht en draaglast in balans zijn.

Relatie met bestaande instrumenten

Tot nu toe bestaan er verschillende manieren om leefbaarheid in beeld te brengen. Gemeenten werken vaak met eigen beleidsindicatoren, woningcorporaties hanteren kengetallen rond woningvoorraad en leefbaarheid, en zorg-, huurders- en welzijnsorganisaties gebruiken cliëntgegevens of signalen uit de praktijk. Deze maatstaven zijn waardevol binnen de eigen context, maar sluiten vaak niet goed op elkaar aan. Er bestaat, zover bekend is, geen uniforme methodiek, waardoor uitkomsten lastig te vergelijken zijn tussen gebieden of door de tijd heen. Ook ontbreekt in veel bestaande instrumenten de bewonersbeleving, terwijl juist die ervaring bepalend is voor hoe leefbaar een buurt wordt gevonden.

Een bekend landelijk instrument is de Leefbaarometer (<https://www.leefbaarometer.nl/>), die gebruikmaakt van meer dan honderd administratieve indicatoren om de leefbaarheid van buurten en wijken te schatten. De Leefbaarometer biedt waardevolle vergelijkingen en trendinformatie, maar heeft ook beperkingen. Zo richt het zich uitsluitend op data, terwijl de beleving van professionals ontbreekt. Daarnaast is de schaal en context generiek, waardoor de resultaten niet altijd aansluiten op lokale vraagstukken, zeker niet in een regio als Fryslân waar de demografische en sociaal-economische dynamiek afwijkt van stedelijke gebieden.

Dit project legt de nadruk op de balans tussen draagkracht (middelen en hulpbronnen) en draaglast (druk en problematiek), een perspectief dat in de Leefbaarometer minder centraal staat. Het ontwikkelt een model dat beleidsmatig meer richting kan geven door deze balans expliciet te maken. Daarnaast biedt het flexibiliteit om samen met lokale partners nieuwe thema's en indicatoren, zoals sociale cohesie, zorggebruik of energiearmoede, toe te voegen en zo het instrument door te ontwikkelen.

De noodzaak voor de huidige aanpak is meerlaags:

- **Vergelijkbaarheid:** er is behoefte aan een uniforme methodiek en indicatoren die door verschillende partijen op dezelfde manier gebruikt en geïnterpreteerd kunnen worden.
- **Diepgang:** leefbaarheid wordt niet alleen bepaald door objectieve kenmerken zoals inkomen, woningtype of buurtvoorzieningen, maar ook door de subjectieve beleving van bewoners.
- **Professionaliseringslaag:** naast cijfers en bewonersbeleving speelt ook de kennis en ervaring van professionals een rol. Zij signaleren vaak vroegtijdig veranderingen in buurten, bijvoorbeeld op het gebied van sociale problematiek of woontevredenheid.

Drie pijlers in het bredere project

Het project Draagkracht en Draaglast rust op drie samenhangende pijlers. De kracht van het project zit in de samenhang tussen deze drie pijlers. Waar pijler 1 harde data en modellen levert, zorgt pijler 2 voor de noodzakelijke verdieping vanuit de leefwereld van bewoners. Pijler 3 vormt de brug naar toepassing en besluitvorming. Samen moeten ze leiden tot een integraal meetinstrument en een gedragen aanpak voor het monitoren van leefbaarheid.

- **Pijler 1 – Harde cijfers (Data en modellen):** Deze pijler richt zich op het ontwikkelen van een kwantitatieve basis voor het meten van tevredenheid met de woonomgeving. Centraal staat het gebruik van CBS-Microdata en aanvullende registratiedata, die met behulp van voorspelmodellen worden geanalyseerd. Doel is om inzichtelijk te maken welke factoren bijdragen aan tevredenheid met de woonomgeving en hoe deze data kunnen worden ingezet om verschillen tussen buurten en wijken in kaart te brengen.
- **Pijler 2 – Beleving en participatie:** Naast objectieve registraties is het essentieel om de ervaren leefwereld van bewoners in beeld te brengen. Deze pijler richt zich op methoden om inwoners te betrekken, bijvoorbeeld via enquêtes, interviews of lokale werkvormen. Hier gaat het om de subjectieve dimensie van leefbaarheid: hoe ervaren mensen hun buurt, wat vinden zij belangrijk, en welke zorgen of wensen leven er?
- **Pijler 3 – Praktijk en beleid:** De derde pijler verbindt de inzichten uit pijler 1 en 2 met de beleidspraktijk. Hier worden de ontwikkelde instrumenten toegepast in lokale pilots, getoetst op bruikbaarheid en bijgestuurd waar nodig. Ook speelt deze pijler een rol in het creëren van een gemeenschappelijke taal tussen betrokken partijen: gemeenten, woningcorporaties, zorg- en welzijnsorganisaties en bewonersinitiatieven.

Voortborduren op pijler 1

Voor dit rapport staat pijler 1 centraal: de vraag hoe we met data en voorspelmodellen de balans tussen draagkracht en draaglast in kaart kunnen brengen. Daarbij ligt de nadruk op de technische en methodologische onderbouwing, met oog voor de beperkingen en mogelijkheden van deze aanpak. Onze aanpak combineert registratiedata (zoals de Woonbase en CBS Microdata) met enquêtegegevens (WoON) en moderne analysemethoden zoals voorspelmodellen (Random Forest, Gradient Boosting).

3. AANPAK

Gefaseerde opzet

De vraag hoe draagkracht en draaglast in wijken zich tot elkaar verhouden, is complex. Om dit te onderzoeken is gekozen voor een gefaseerde aanpak met drie analyses. Deze analyses bouwen systematisch op elkaar voort: van een eerste verkenning met harde cijfers, via een verdieping met aanvullende factoren, naar een verbreding waarin ook de bredere omgeving wordt meegenomen. De fasering maakt het mogelijk om de waarde van verschillende databronnen en niveaus van analyse afzonderlijk te toetsen, en om geleidelijk te werken naar een bruikbaar instrument. Er is gewerkt met CBS-Microdata (later meer over de data).

Analyse 1 – Registratiedata als basis

De eerste analyse onderzoekt in hoeverre ervaren leefbaarheid te voorspellen is met uitsluitend objectieve registratiedata. De doelvariabele – tevredenheid met de woonomgeving – is afkomstig uit de WoON-enquête, terwijl de verklarende variabelen uit de Woonbase komen. Omdat het aantal Friese respondenten in WoON beperkt is, wordt gebruikgemaakt van alle landelijke waarnemingen. Het model wordt daarmee getraind op een brede dataset, zodat kan worden nagegaan of er algemene patronen bestaan die ook toepasbaar zijn op Fryslân.

Deze analyse heeft vooral een verkennend karakter. Zij laat zien wat harde registraties – zoals woningtype, bouwjaar, WOZ-waarde of inkomenspositie – kunnen bijdragen aan het voorspellen van tevredenheid met de woonomgeving, en waar de grenzen liggen. Daarmee vormt deze stap de basis voor de volgende analyses.

Analyse 2 – Aanvullende verklarende factoren

De tweede analyse voegt informatie toe die niet in registratiedata aanwezig is, maar wel samenhangt met tevredenheid met de woonomgeving. Het gaat daarbij om aanvullende variabelen uit WoON, zoals buurtbeleving, verhuishwensen en woonsituatie. Deze gegevens zijn niet voor de volledige bevolking beschikbaar, maar bieden wel waardevolle inzichten in de achterliggende mechanismen van tevredenheid en ontevredenheid.

De nadruk in deze analyse ligt op verklaren in plaats van voorspellen. Door meer inzicht te krijgen in de factoren die tevredenheid met de woonomgeving beïnvloeden, kan beter worden begrepen waarom sommige wijken en buurten kwetsbaarder zijn dan andere. De analyse levert daarmee kennis op die niet direct generaliseerbaar is, maar wel richting geeft aan het ontwikkelen van indicatoren voor een breder meetinstrument.

Analyse 3 – Buurt- en wijkcontext

De derde analyse introduceert een contextuele laag door buurtkenmerken toe te voegen. Voorbeelden zijn het aandeel sociale huurwoningen, de gemiddelde WOZ-waarde en het inkomensniveau op buurniveau. Deze gegevens zijn breed beschikbaar en maken het mogelijk de balans tussen draagkracht en draaglast niet alleen op individueel niveau, maar ook in de ruimtelijke context te verkennen.

Deze uitbreiding is belangrijk omdat tevredenheid met de woonomgeving niet uitsluitend het resultaat is van individuele kenmerken, maar mede wordt gevormd door de omgeving waarin mensen wonen. Door ook de buurtcontext mee te nemen, kan worden onderzocht in hoeverre verschillen tussen wijken verklaard worden door structurele factoren.

Meerwaarde huidige aanpak

De drie analyses vormen samen een groeimodel. De eerste stap brengt de mogelijkheden en beperkingen van harde registratiedata in beeld. De tweede stap verdiept dit beeld met aanvullende, subjectieve factoren. De derde stap verkent de rol van context op buurtniveau. Deze fasering is nodig omdat leefbaarheid een gelaagd fenomeen is, dat niet in één keer volledig kan worden gemodelleerd. Door stapsgewijs te werken ontstaat een beter onderbouwd beeld van welke data en modellen nodig zijn om draagkracht en draaglast systematisch in kaart te brengen. De drie analyses vormen niet alleen een methodologische opbouw, maar ook een inhoudelijke koppeling met de pijlers die de kern van het bredere project vormen, met pijler 1 'in the lead.'

4. MODELLEN

Waarom voorspelmodellen?

De centrale vraag in dit onderzoek is in hoeverre tevredenheid met de woonomgeving kan worden verklaard en voorspeld met behulp van data. Traditionele statistiek laat verbanden zien, maar schiet tekort wanneer er veel variabelen tegelijk spelen of wanneer verbanden niet lineair zijn. Daarom zijn voorspelmodellen ingezet die in staat zijn complexe patronen te leren uit grote hoeveelheden gegevens.

Een voorspelmodel is een analysemethode die patronen in bestaande data identificeert om een waarde of uitkomst te schatten. Dit type model wordt vaak ingezet om gedrag of behoeften te voorspellen, bijvoorbeeld in het onderwijs, de zorg of mobiliteit. Voorbeelden van toepassingen zijn het voorspellen van leerachterstanden, zorggebruik, verhuizingen, het weer, verkeersdrukke en de kans op energiearmoede. De modellen geven een indicatie van de mate waarin objectieve data geschikt zijn om tevredenheid met de woonomgeving te voorspellen.

Een belangrijk voordeel van voorspelmodellen is dat er vooraf geen selectie van variabelen nodig is: het model bepaalt zelf welke data relevant zijn. Daarbij wordt alle beschikbare data benut, inclusief onverwachte verbanden die anders mogelijk over het hoofd worden gezien. Voorspelmodellen kunnen bovendien omgaan met verschillende soorten data, zoals numerieke en categorische gegevens. Hierdoor is de analyse minder afhankelijk van aannames vooraf.

Daarnaast kunnen modellen zichtbaar maken welke kenmerken het meest bijdragen aan de voorspellingen. Zo kan er een onderscheidt gemaakt worden welke factoren meer of minder essentieel zijn in het voorspellen van tevredenheid met de woonomgeving.

Beslisbomen in gewone taal

De gebruikte modellen zijn gebaseerd op beslisbomen. Een beslisboom werkt als een soort vraag-en-antwoord-spel. Een boom kan bijvoorbeeld beginnen met de vraag: “Is de woning een koopwoning?” Als het antwoord ja is, volgt een ander pad dan wanneer het antwoord nee is. Daarna kan bijvoorbeeld gevraagd worden naar het inkomen, of naar het bouwjaar van de woning. Uiteindelijk leidt dit pad naar een voorspelling, zoals “tevreden” of “ontevreden”.

De boom bestaat uit een reeks knooppunten en vertakkingen. Elk knooppunt vertegenwoordigt een specifiek kenmerk uit de dataset, zoals leeftijd, inkomen of woningtype. Op basis van de waarde van dat kenmerk splitst de boom zich in vertakkingen naar andere knooppunten. Het model leert deze structuur door patronen in de data te analyseren. Tijdens het leerproces bepaalt het model zelf welke kenmerken het meest informatief zijn en kiest het automatisch de knooppunten die leiden tot de beste voorspellingen.

Een enkelvoudige boom is eenvoudig uit te leggen, maar kan instabiel zijn en te sterk aansluiten op de specifieke trainingsdata. Om dit te voorkomen gebruiken we methoden die niet op één boom vertrouwen, maar op een verzameling bomen. Daarmee worden de voorspellingen robuuster en nauwkeuriger.

Beslisbomen vormen de basis voor geavanceerdere machine learning-modellen zoals Random Forest en Gradient Boosting, die meerdere bomen combineren voor hogere nauwkeurigheid en robuustheid. In figuur 4.1 zijn de verschillende modellen gevisualiseerd.

Random Forest

Een Random Forest is een machine learning model. Elke boom in het model maakt een inschatting door vragen te stellen over de gegevens, zoals bijvoorbeeld leeftijd, inkomen of woonomgeving. In plaats van één enkele boom, gebruikt een Random Forest veel bomen tegelijk. Elke boom wordt getraind op een willekeurige subset van de data, zodat ze allemaal iets anders leren. Wanneer het model een nieuwe situatie moet inschatten, zoals of iemand waarschijnlijk ontevreden is over de leefomgeving, dan stemmen alle bomen. De uiteindelijke voorspelling wordt bepaald door de meerderheid van de bomen: als de meeste bomen zeggen dat iemand ontevreden is, dan is dat de voorspelling van het model.

Het voordeel van Random Forest is de stabiliteit. Door het gemiddelde van vele bomen te nemen, worden toevallige patronen uitgevlakt. Het model kan goed omgaan met uiteenlopende soorten data en geeft inzicht in de relatieve belangrijkheid van variabelen. Nadelen zijn dat het minder goed inspeelt op zeldzame klassen, en dat de interpretatie van honderden bomen minder direct is dan die van een enkele boom.

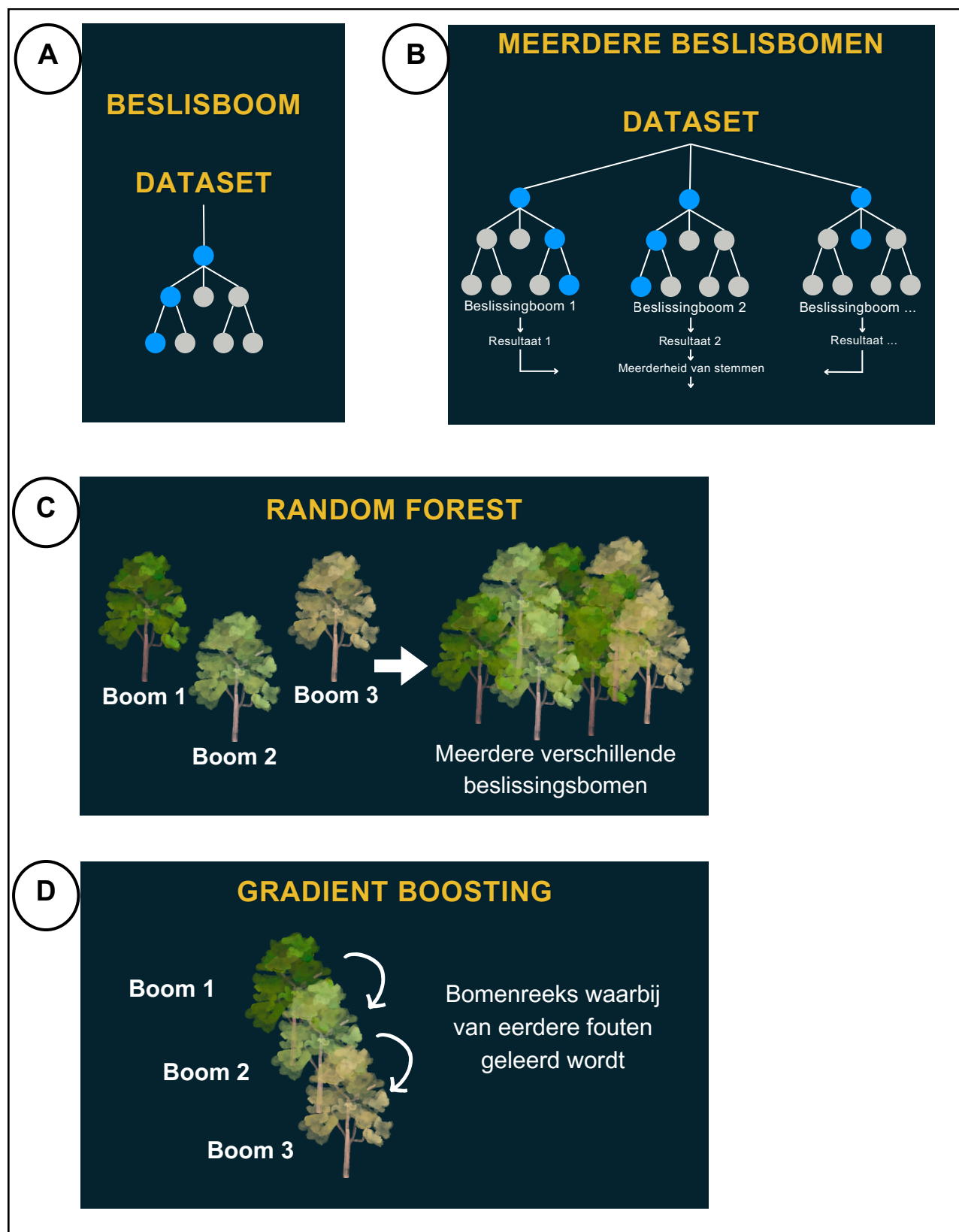
Door deze aanpak kan het random Forest model goed omgaan met complexe en diverse gegevens, en betrouwbare voorspellingen doen in nieuwe situaties.

Gradient Boosting

Random Forest en Gradient Boosting zijn beide modellen die gebruikmaken van beslissingsbomen, maar ze werken op een andere manier. Gradient Boosting bouwt de bomen stap voor stap. Elke nieuwe boom leert vooral van de fouten die de vorige bomen hebben gemaakt. Zo wordt het model steeds beter in het classificeren van de moeilijke gevallen. Bij Gradient Boosting wordt de volledige dataset gebruikt, maar met extra aandacht voor de moeilijk te voorspellen gevallen. Bij Random Forest worden meerdere bomen tegelijk gebouwd, waarbij elke boom een willekeurige subset van de data gebruikt.

Het voordeel is dat Gradient Boosting vaak nauwkeuriger is dan Random Forest. Het nadeel is dat het gevoeliger is voor overfitting: het kan zich te veel richten op de trainingsdata, waardoor de generaliseerbaarheid afneemt. Ook zijn de modellen afhankelijk van zorgvuldige instelling van parameters, zoals het aantal bomen en de leersnelheid.

Figuur 4.1 – Visualisatie van a) enkele beslisboom, b) meerdere beslisbomen gecombineerd, c) random forest en d) gradient boosting.



Uitlegbaarheid van de modellen

Naast het geven van een voorspelling, bieden deze modellen ook inzicht in welke indicatoren het meest bijdragen aan de uitkomst. Dit gebeurt zonder dat vooraf keuzes hoeven te worden gemaakt over welke variabelen belangrijk zijn. Met behulp van SHAP-waardes krijgt elke variabele een kwantitatieve score, die aangeeft hoeveel invloed die variabele heeft op de voorspelling, rekening houdend met de impact van andere indicatoren.

SHAP berekent per variabele hoeveel die bijdraagt aan de uitkomst. SHAP-waarden zijn te zien als een balans: elke variabele duwt de voorspelling een stukje omhoog of omlaag richting een bepaalde categorie. Wat SHAP bijzonder maakt, is dat de methode ook rekening houdt met de onderlinge samenhang tussen variabelen. Dit gebeurt door systematisch alle mogelijke combinaties van variabelen te beschouwen en te analyseren wat er gebeurt als een bepaalde variabele wordt weggelaten. Op die manier wordt de bijdrage van een variabele niet los gezien, maar in de context van andere variabelen.

Een belangrijke eigenschap van beide modellen is dat ze niet alleen voorspellen, maar ook inzicht geven in waarom ze een bepaalde voorspelling doen. Dit gebeurt via SHAP-waarden.

Evaluatie van de prestaties van het model

Om te beoordelen hoe goed de modellen werken, is de dataset opgesplitst in een trainingsdataset (80%) en een testdataset (20%). De modellen zijn getraind op de eerste set en getoetst op de tweede. De trainingsdataset (een random subset van 80% van de data) wordt gebruikt om het model te leren welke patronen er in de data zitten, bijvoorbeeld welke kenmerken samenhangen met tevredenheid met de woonomgeving. Het model leert hieruit verbanden tussen de verklarende variabelen en de doelvariabele.

De testdataset (een random subset van 20% van de data) wordt apart gehouden en pas gebruikt nadat het model is getraind. Hiermee wordt beoordeeld hoe goed het model in staat is om nieuwe, onbekende gevallen te voorspellen. Door deze aanpak wordt voorkomen dat het model alleen maar goed presteert op de data waarop het is getraind, en ontstaat er een realistischer beeld van de kwaliteit van de voorspellingen. Deze werkwijze is essentieel om overfitting te voorkomen en om te zorgen dat het model ook buiten de trainingsdata bruikbaar is, bijvoorbeeld voor het inschatten van tevredenheid over de woonomgeving in andere buurten of toekomstige situaties.

De doelvariabele is echter scheef verdeeld: de meerderheid van de respondenten is (zeer) tevreden, terwijl ontevredenheid zeldzaam is. Dit leidt ertoe dat de modellen de grote groepen redelijk goed voorspellen, maar moeite hebben met de kleine groepen. Daarom is niet alleen gekeken naar de algehele nauwkeurigheid (accuracy), maar ook naar andere maatstaven.

- Precision (precisie): van alle voorspellingen van bijvoorbeeld “ontevreden”, hoeveel waren correct?
- Recall (gevoeligheid): van alle mensen die werkelijk ontevreden zijn, hoeveel herkende het model?

Omgaan met scheve data

Het probleem van de scheve verdeling is niet uniek voor dit onderzoek. Er bestaan technieken om dit te verzachten, zoals:

- het zwaarder laten meewegen van kleine groepen,
- het aanvullen van zeldzame categorieën met kunstmatige voorbeelden,
- of het samenvoegen van naburige categorieën om de balans te verbeteren.

In deze analyse zijn enkele van deze strategieën getest, maar zonder groot succes. Daarom is gewerkt met de data 'as is.'

Robuustheid en generaliseerbaarheid

Een tweede aandachtspunt is de robuustheid van de modellen. Een model dat goed werkt op de huidige dataset, kan minder goed presteren in een andere regio of in een ander jaar. Daarom is het belangrijk om de modellen ook ruimtelijk (in verschillende gemeenten) en temporeel (met data van verschillende jaren) te testen.

Dit is een aandachtspunt voor de volgende fases van het project, waarin de methodiek in pilots verder kan worden gevalideerd.

Confusion matrix uitgelegd

Voorspelmodellen zoals Random Forest en Gradient Boosting doen voorspellingen op basis van patronen in de data, bijvoorbeeld over (on)tevredenheid over de leefomgeving. De modellen geven inzicht in hoe goed ze daadwerkelijk voorspellen door middel van validatie op de testdataset. De nauwkeurigheid van het model is het percentage correcte voorspelde uitkomsten.

Een confusion matrix is een tabel die laat zien hoe de voorspellingen van een model zich verhouden tot de werkelijkheid. De werkelijke categorieën staan meestal op de verticale as, de voorspelde categorieën op de horizontale as.

- De diagonaal van linksonder naar rechtsboven laat de juiste voorspellingen zien. Hoe voller deze cellen zijn, hoe beter het model de werkelijkheid benadert.
- De cellen naast de diagonaal laten de fouten zien: gevallen waarin het model bijvoorbeeld iemand als “tevreden” voorspelt, terwijl die persoon zichzelf “neutraal” of “ontevreden” beoordeelde.

De confusion matrix is een krachtig hulpmiddel: hij laat niet alleen zien hoeveel voorspellingen correct waren, maar ook waar de fouten gemaakt worden. Dit helpt om beter te begrijpen welke groepen wel en niet door het model worden herkend.

Relatie met beleid

Het is van belang om de vertaalslag naar beleid te maken. De modellen voorspellen niet op individueel niveau wat iemand ervaart, maar maken patronen zichtbaar die relevant zijn voor buurten en wijken. Inzicht in de belangrijkste factoren – zoals WOZ-waarde, woningtype of inkomensniveau – kan gemeenten en corporaties helpen om beleid beter te richten.

De uitkomsten moeten daarom vooral gezien worden als signalen: waar kan draagkracht worden versterkt, en waar stapelt draaglast zich op? De modellen geven geen sluitende antwoorden, maar leveren wel een datagedreven basis voor verdere analyse en dialoog.

Belang van data voor modellen

Een belangrijk kenmerk van de modellen in dit onderzoek is dat zij gebaseerd zijn op data op persoonlijk en huishoudniveau. Dit schaalniveau is van grote waarde, omdat het meer nuance geeft dan analyses die alleen op wijk- of gemeentelijk niveau plaatsvinden.

Op buurtniveau kunnen bijvoorbeeld gemiddelde inkomens, woningwaarden of werkloosheidspercentages worden weergegeven. Zulke gegevens zijn nuttig voor een eerste indruk, maar verhullen vaak grote verschillen tussen individuele huishoudens. Twee buurten met hetzelfde gemiddelde inkomen kunnen in werkelijkheid heel anders zijn samengesteld: de ene met veel middeninkomens, de andere met een mix van hoge en lage inkomens.

Door te werken met microdata op persoons- en huishoudniveau kunnen de modellen deze variatie expliciet meenemen. Dit maakt het mogelijk om fijnmaziger verbanden zichtbaar te maken, bijvoorbeeld tussen woningtype, inkomen en ervaren leefbaarheid. Bovendien kunnen patronen die op buurtniveau onzichtbaar blijven, op individueel niveau wel worden blootgelegd.

Het gebruik van microdata verhoogt daarmee de verklarende kracht van de modellen en vergroot de mogelijkheden om beleid gericht af te stemmen. Wel vraagt deze aanpak om zorgvuldige omgang met privacy en om geavanceerde analysemethoden, omdat de hoeveelheid data en variatie aanzienlijk groter is. In het volgende hoofdstuk is meer informatie te vinden over de data.

5. DATA

CBS-Microdata

CBS-Microdata zijn gedetailleerde, op persoonsniveau verzamelde gegevens die door het Centraal Bureau voor de Statistiek (CBS) beschikbaar worden gesteld voor wetenschappelijk en beleidsgericht onderzoek. Deze data bieden een rijke bron van informatie over uiteenlopende maatschappelijke thema's, zoals werk, inkomen, gezondheid en wonen. De data zijn geanonimiseerd en onderzoekers mogen alleen geaggregeerde uitkomsten naar buiten brengen, zodat de privacy van personen en huishoudens gewaarborgd blijft. De analyses in dit onderzoek zijn uitgevoerd met CBS-Microdata.

Het gebruik van microdata heeft voordelen. Allereerst kunnen analyses worden uitgevoerd op een fijnmazig niveau, waardoor verschillen tussen huishoudens zichtbaar worden die in openbare statistieken vaak verborgen blijven. Daarnaast is het mogelijk om gegevens uit verschillende registraties te koppelen, zoals inkomensgegevens uit de belastingadministratie en woningkenmerken uit het Kadaster. Omdat de registraties de volledige bevolking omvatten, ontstaat bovendien een consistent en volledig beeld.

Het gebruik van CBS-Microdata is strikt gereguleerd. De gegevens zijn niet herleidbaar tot individuen en vallen onder strenge privacyregels. Organisatie met een instellingsmachtiging kunnen toegang krijgen, en uitsluitend via een beveiligde remote-access omgeving van het CBS. Onderzoekers moeten vooraf een toets doorlopen om toegang te verkrijgen. Daarnaast wordt alle output die het systeem verlaat, gecontroleerd door het CBS om te waarborgen dat er geen vertrouwelijke informatie wordt gedeeld.

Een belangrijk voordeel van CBS-Microdata is de mogelijkheid tot het koppelen van verschillende datasets (zoals energiegegevens met de woonbase), waardoor diepgaande en domeinoverstijgende analyses mogelijk worden. Deze combinatie van detailniveau, veiligheid en flexibiliteit maakt CBS Microdata tot een krachtige en betrouwbare bron voor onderzoek en beleidsontwikkeling

Woonbase en WoON

Twee bronnen staan centraal: Woonbase en WoON. De woonbase is een samengestelde registratiedataset waarin persoons-, huishoud- en woningkenmerken worden gecombineerd. Het bevat gegevens over leeftijd, inkomen, woningtype, bouwjaar, WOZ-waarde en eigendomsvorm van heel Nederland. Daarmee ontstaat een beeld van de woon- en leefsituatie van huishoudens. Deze dataset is opgebouwd uit koppelingen tussen verschillende registraties, zoals de Basisregistratie Personen (BRP) en de Basisregistratie Adressen en Gebouwen (BAG).

De Woonbase biedt:

- Volledige populatiegegevens (geen steekproef).
- Informatie over woningkenmerken, huishoudens en verhuisbewegingen.
- Mogelijkheid tot koppeling van datasets.
- Toegang via een beveiligde remote-access omgeving.

WoON (WoonOnderzoek Nederland) is een periodieke enquête waarin huishoudens zelf rapporteren over hun woonsituatie en woonomgeving. Het bevat vragen over tevredenheid met de woning en de buurt, ervaringen, sociale cohesie, woonlasten, verhuisplannen en waardering voor voorzieningen. Waar Woonbase objectieve omstandigheden weergeeft, geeft WoON juist inzicht in de subjectieve beleving. In 2024 bevatte WoON 41.295 respondenten.

Belangrijke kenmerken van de WoON-data zijn:

- Representatieve steekproef van huishoudens.
- Informatie over woontevredenheid, verhuishwensen en buurtbeleving.
- Subjectieve beleving als aanvulling op objectieve registratiedata.
- De subjectieve beleving is alleen beschikbaar voor de WoON-respondenten, waardoor deze informatie niet voor iedereen geldt.

Koppeling en analyse-eenheid

De kracht van dit onderzoek ligt in de koppeling van beide bronnen. Voor respondenten in WoON zijn aanvullende kenmerken uit Woonbase beschikbaar gemaakt. Hierdoor kan de tevredenheid met de woonomgeving direct worden verbonden met objectieve kenmerken van woning en huishouden. De analyses vinden plaats op persoons- en huishoudniveau. Dit maakt het mogelijk om patronen in detail te onderzoeken en deze vervolgens te vertalen naar hogere schaalniveaus zoals buurten of gemeenten.

Voorbeelden en schaal

De dataset die zo ontstaat is data-rijk. Een respondent kan bijvoorbeeld een 45-jarige alleenstaande huurder zijn met een huishoudinkomen van €38.000 en een appartement met een WOZ-waarde van €215.000. Uit WoON weten we daarnaast of deze persoon tevreden is met de buurt, zich veilig voelt en verhuisplannen heeft. Zulke combinaties maken het mogelijk om zowel de objectieve omstandigheden als de subjectieve beleving in één analyse mee te nemen.

Omdat WoON een steekproef is, is het aantal Friese respondenten relatief klein. Om voldoende robuustheid te garanderen, zijn de modellen daarom landelijk getraind op ruim 41.000 waarnemingen. De gevonden patronen kunnen vervolgens worden toegepast op de Friese populatie in Woonbase, zodat toch een volledig dekkend beeld ontstaat.

Relatie met de modellen

De gebruikte voorspelmodellen (Random Forest en Gradient Boosting) zijn juist krachtig wanneer zij kunnen leren van veel verschillende kenmerken op laag schaalniveau. Hoe meer variatie de modellen kunnen benutten, hoe groter de kans dat zij patronen vinden die samenhangen met tevredenheid met de woonomgeving. Het gebruik van microdata maakt dit mogelijk: modellen kunnen subtiele combinaties oppikken (bijvoorbeeld het samenspel van woningtype en inkomen) die op geaggregeerd niveau onzichtbaar zouden blijven.

De keerzijde is dat zulke modellen ook gevoelig zijn voor scheve verdelingen, zoals het grote aandeel tevreden respondenten. Juist daarom is het belangrijk dat de kwaliteit, volledigheid en verdeling van de data goed worden begrepen en verantwoord.

6. RESULTATEN

De eerste analyse richt zich op de vraag in hoeverre tevredenheid met de woonomgeving kan worden voorspeld met uitsluitend registratiedata. Dit is een logische eerste stap: registraties zoals inkomen, woningkenmerken en huishoudsamenstelling zijn breed beschikbaar en volledig dekkend voor de bevolking. Ze vormen daarmee een stevig fundament om patronen in draagkracht en draaglast te verkennen.

De doelvariabele in deze analyse is de tevredenheid met de woonomgeving, afkomstig uit de WoON-enquête. Deze informatie is beschikbaar voor ruim 41.000 respondenten in Nederland en is via een koppeling aan Woonbase verrijkt met objectieve persoons- en woningkenmerken. Omdat het aantal Friese respondenten beperkt is, is gekozen voor een landelijk model, dat vervolgens kan worden toegepast op de Friese populatie.

Verdeling doelvariabele

De verdeling van de tevredenheid met de woonomgeving laat een duidelijke scheefheid zien (Tabel 6.1). Van de in totaal 41.295 respondenten geeft meer dan de helft (22.041; 53,4%) aan tevreden te zijn, terwijl bijna een derde (12.653; 30,6%) zelfs zeer tevreden is. Daarmee behoort ruim vier op de vijf respondenten tot de groep die positief oordeelt over de woonomgeving. Slechts 1.674 respondenten (4,1%) noemen zich ontevreden en 560 (1,4%) zeer ontevreden.

In totaal is meer dan 84% van de respondenten tevreden of zeer tevreden. De groepen die een neutraal of negatief oordeel geven zijn klein. Voor de analyse betekent dit dat de dataset ongebalanceerd is: het model krijgt veel voorbeelden van tevreden bewoners, maar nauwelijks van ontevreden bewoners. Dit maakt het voorspellen van ontevredenheid bijzonder lastig en beïnvloedt de prestaties van de modellen.

Tabel 6.1 – Verdeling tevredenheid met de leefomgeving (WoON-data 2024).

Categorie	Aantal	Percentage
Zeetevreden	12.653	30,64%
Tevreden	22.041	53,37%
Neutraal	4.367	10,58%
Ontevreden	1.674	4,05%
Zeetevreden	560	1,36%
Totaal	41.295	100%

6.1 ANALYSE 1: REGISTRATIEDATA ALS BASIS

Voor analyse 1 is gebruikgemaakt van registratiedata uit de Woonbase om de tevredenheid met de woonomgeving (uit WoON) te voorspellen met behulp van verschillende voorspelmodellen. Voorafgaand aan de analyse zijn een groot aantal verklarende variabelen geselecteerd door het bredere onderzoeksteam. Deze variabelen, of subsets daarvan, zijn gebruikt in verschillende machine learning modellen, waaronder Random Forest en Gradient Boosting. Ter vergelijking is ook een klassiek statistisch model toegepast in de vorm van logistische regressie.

Gedurende het analyseproces is de focus komen te liggen op de Gradient Boosting methode. Binnen deze aanpak zijn meerdere variabelenselecties getest, is geëxperimenteerd met verschillende hyperparameters, en zijn de uitkomstklassen samengevoegd om de prestaties van het model verder te onderzoeken.

Het voordeel van deze modellen is dat zij robuust kunnen omgaan met uiteenlopende typen data en relatief ongevoelig zijn voor multicollineariteit. Bovendien zijn ze in staat complexe patronen te leren, zonder dat vooraf strikte aannames over de vorm van de relaties nodig zijn. Voor dit onderzoek zijn beide methoden gebruikt om te toetsen of de resultaten consistent zijn en om de voorspelkracht zo goed mogelijk te benutten.

Voor analyse 1 zijn er maximaal 59 indicatoren (variëaties gebruikte minder variabelen) meegenomen in de analyse.

Technische omschrijving: Implementatie en validatie

De modellen zijn ontwikkeld in Python binnen de beveiligde CBS-omgeving. Daarbij is een reeks stappen doorlopen om de analyses zorgvuldig en reproduceerbaar uit te voeren:

- **Datasplitsing:** de dataset is willekeurig verdeeld in een trainingsdeel (80%) en een testdeel (20%). Dit maakt het mogelijk de prestaties te evalueren op data die het model nog niet eerder heeft gezien.
- **Optimalisatie:** met behulp van GridSearch zijn hyperparameters systematisch afgestemd, zoals het aantal bomen, de maximale diepte van de bomen en het leerpercentage bij Gradient Boosting.
- **Beoordeling:** de prestaties zijn niet alleen beoordeeld op basis van nauwkeurigheid, maar ook met indicatoren die beter rekening houden met de scheve verdeling in de data, zoals precision, recall en balanced accuracy.
- **Class weighting:** verschillende wegeningen voor de klassen zijn getest om de invloed van de oververtegenwoordigde categorieën ('tevreden' en 'zeer tevreden') te beperken.

Deze aanpak leverde waardevolle inzichten op in hoe de modellen functioneren, maar ook in de beperkingen die voortkomen uit de data zelf. Ondanks optimalisatie en het toepassen van technieken om de onbalans te corrigeren, bleven de prestaties bij de minderheidsklassen beperkt. Dit wijst erop dat het probleem niet zozeer in de methode zit, maar vooral in de eigenschappen van de dataset.

Modeluitkomsten Analyse 1

De dataset is opgedeeld in een trainingsdataset en een testdataset. De trainingsdataset bevat 80% van de observaties. Hier is het Gradient Boosting model op getraind. De testdataset bevat 20% van de observaties (20% x 41.295 = 8.259 observaties).

De algehele nauwkeurigheid van de modellen ligt rond de 50–55%. Hoewel dit in eerste instantie redelijk lijkt, moet dit cijfer met voorzichtigheid worden geïnterpreteerd. Het hoge aandeel tevreden en zeer tevreden respondenten betekent dat een model al een behoorlijke score haalt door vooral deze groepen goed te voorspellen. De meer beleidsmatig relevante categorieën — neutraal en ontevreden — worden nauwelijks herkend. Dit blijkt ook uit de confusion matrix: de meeste gevallen in deze groepen worden ten onrechte als 'tevreden' geclassificeerd. De balanced accuracy, die corrigeert voor deze scheve verdeling, ligt daarom aanzienlijk lager dan de ruwe nauwkeurigheid. Zelfs geavanceerde modellen als Gradient Boosting kunnen zonder aanvullende informatie weinig onderscheid maken tussen neutrale en ontevreden bewoners.

Tabel 6.2 laat op de horizontale as de door het model voorspelde waarden zien en op de verticale as de werkelijke waarden uit de testdataset. In 291 gevallen voorspelt het model de uitkomst 'Zeer tevreden' correct en in 4.147 gevallen de uitkomst 'Tevreden'. De overige uitkomsten worden helemaal niet voorspeld door het model. Dit is het gevolg van de ongebalanceerde verdeling van de doelvariabele. De resultaten in Tabel 6.2 laten zien dat het model in $(4.147+291)/8.259 \cdot 100\% = 54\%$ van de gevallen correct voorspelt. Dit bevestigt dat de voorspelbaarheid van ontevredenheid met registratiedata beperkt is.

Tabel 6.2 – Confusion matrix analyse 1.

Werkelijke waarde	Ze	Te	Ne	On	Ze
	te	ve	u	te	te
	vre	re	tra	vre	vre
	de	de	al	de	de
	en	n		n	n
Ze	3	118	0	0	0
	3	307	0	0	0
	24	858	0	0	0
	251	4147	0	0	0
	291	2257	0	0	0
Voorspelde waarde					

Belangrijkste voorspellers (SHAP-waarden) Analyse 1

In Tabel 6.3 en Figuur 6.1 zijn de tien belangrijkste factoren weergegeven die volgens het model samenhangen met de tevredenheid over de woonomgeving. Het figuur is gebaseerd op SHAP-waarden en laat het gemiddelde (absolute) effect per klasse zien. Tabel 6.3 laat het totale gemiddelde effect over alle klassen zien. Daarmee wordt duidelijk welke variabelen het meest bijdragen aan de voorspellingen, en voor welke uitkomsten zij het sterkst doorwerken.

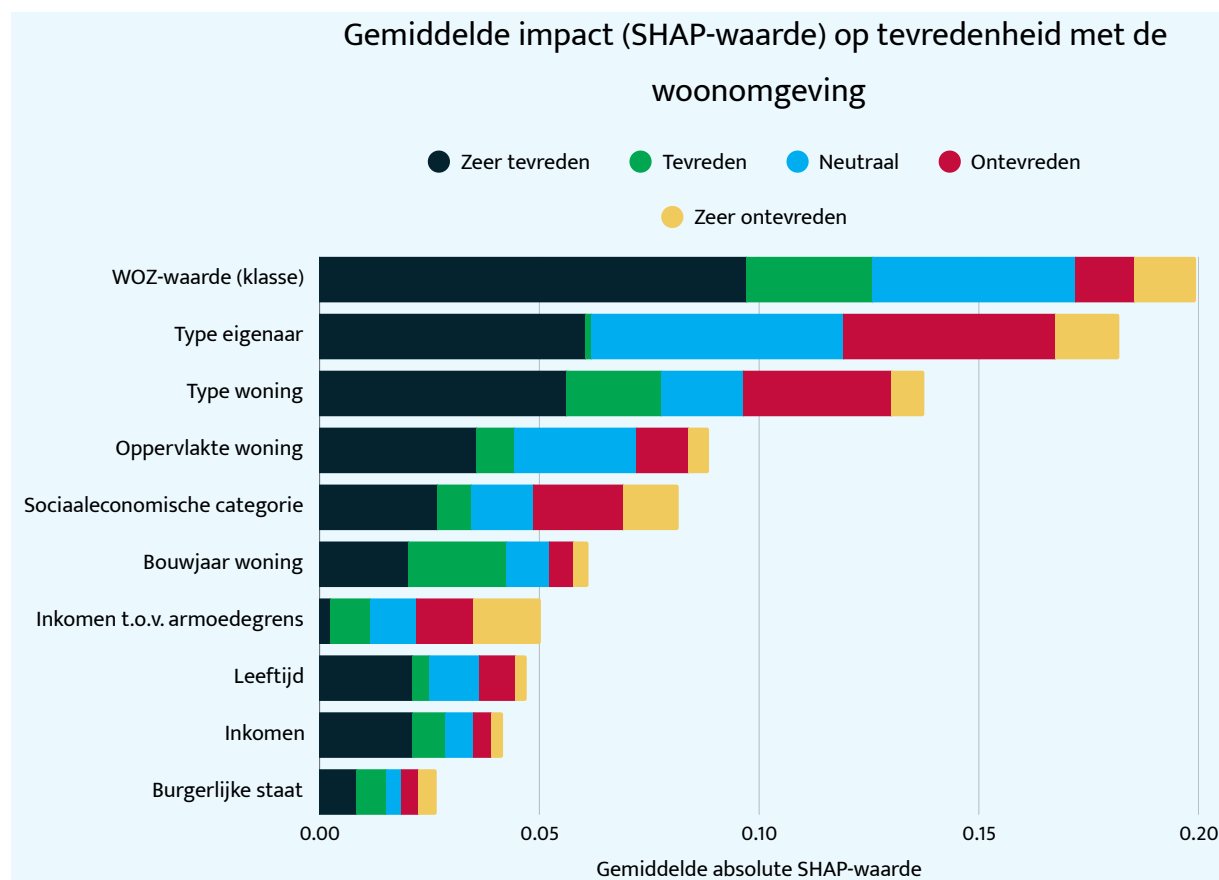
- **WOZ-waarde** van de woning is de meest invloedrijke voorspeller. Huishoudens in woningen met een hogere WOZ-waarde hebben een grotere kans om de woonomgeving als “zeer tevreden” te beoordelen. Dit effect is in de SHAP-resultaten (Figuur 6.1) vooral zichtbaar bij de uitkomst “zeer tevreden” representeert. Voor ontevredenheid is de invloed van WOZ-waarde gering.
- **Eigendomsvorm** (huur of koop) is een structurele factor. Eigen woningbezit hangt vaker samen met hogere tevredenheid, terwijl sociale huur vaker met lagere tevredenheid wordt geassocieerd.
- **Woningtype** staat op de derde plaats en werkt juist de andere kant op. Bepaalde woningtypen, zoals kleine appartementen of portiekwoningen, verhogen de kans dat een respondent ontevreden is. Voor de categorieën “tevreden” en “zeer tevreden” is dit effect minder uitgesproken.
- **Oppervlakte en bouwjaar** van de woning blijken aanvullend relevant: nieuwere en grotere woningen worden vaker geassocieerd met hogere tevredenheid.
- **Sociaaleconomische categorie en inkomen** spelen eveneens een rol. Wat voor werk iemand doet, of personen met een relatief lagere inkomenspositie, scoren gemiddeld lager op tevredenheid.
- **Leeftijd** van de respondent laat een consistent patroon zien: oudere bewoners rapporteren gemiddeld hogere tevredenheid dan jongere bewoners.

De SHAP-waarden maken deze patronen inzichtelijk: ze laten zien hoeveel elke variabele bijdraagt aan een voorspelling, en in welke richting. Daarmee wordt zichtbaar dat vooral de fysieke en financiële kenmerken van woningen doorslaggevend zijn voor het onderscheid tussen tevreden en ontevreden bewoners.

Tabel 6.3 – 10 belangrijkste factoren (SHAP-waarde) om tevredenheid met de woonomgeving te voorspellen.

Variabele	Totale impact	Variabele	Totale impact
WOZ-waarde (klasse)	0.20	Bouwjaar woning	0.06
Type eigenaar	0.18	Inkomen t.o.v. armoedegrens	0.05
Type woning	0.14	Leeftijd	0.05
Oppervlakte woning	0.09	Inkomen	0.04
Sociaaleconomische categorie	0.08	Burgerlijke staat	0.03

Figuur 6.1 – Visualisatie van de 10 belangrijkste factoren (SHAP-waarde) per categorie in de variabele tevredenheid met de woonomgeving (van 'zeer ontevreden' tot 'zeer tevreden' (Analyse 1).



6.2 ANALYSE 2: TOEVOEGING SUBJECTIEVE GEGEVENS

Analyse 1 liet zien dat registratiedata zoals WOZ-waarde, woningtype, eigendom en inkomenspositie een zekere mate van voorspelkracht hebben bij het verklaren van tevredenheid met de woonomgeving, maar dat de verklaringskracht beperkt blijft. De modellen herkenden de meerderheid tevreden bewoners redelijk goed, maar hadden grote moeite met het onderscheiden van de kleinere groepen die neutraal of ontevreden zijn.

Daarom is in Analyse 2 gekozen voor een uitbreiding met gegevens uit de WoON-enquête. Waar registratiedata vooral objectieve kenmerken vastleggen, brengt de WoON-enquête de beleving van bewoners in beeld. Deze beleving is essentieel voor het begrijpen van draagkracht en draaglast in wijken: mensen ervaren hun omgeving niet alleen via stenen, prijzen en inkomens, maar ook via veiligheid, sociale samenhang, voorzieningen en hun persoonlijke woonsituatie.

Het doel van Analyse 2 is tweeledig:

1. Verbeteren van de voorspelkracht door subjectieve belevingsdata aan het model toe te voegen.
2. Beter begrijpen van de drijvende factoren achter tevredenheid met de woonomgeving: welke thema's en variabelen verklaren waarom sommige mensen (zeer) tevreden zijn, terwijl anderen ontevredenheid ervaren?

De grote hoeveelheid beschikbare variabelen is thematisch geordend in clusters zoals algemeen, verhuizen, leefbaarheid, hypotheek, zorg, zorggeschiktheid, VvE en energie. Omdat elk thema veel afzonderlijke variabelen bevat, is binnen elk thema een selectie gemaakt van de meest relevante verklarende factoren. Dit zorgt ervoor dat de modellen niet overbelast raken en dat de analyse zich richt op de variabelen die inhoudelijk het meest betekenisvol zijn.

Voor analyse 2 zijn er maximaal 719 indicatoren (variëaties gebruikte ook 197 variabelen) meegenomen in de analyse.

Verschillen tussen enquête- en registratiedata

Het combineren van registratiedata en enquêtedata betekent dat twee verschillende data bij elkaar worden gebracht. Registratiedata zijn volledig dekkend, objectief en afkomstig uit administratieve bronnen. Zij leggen kenmerken vast, zoals woningtype, WOZ-waarde, huurprijs, inkomen en huishoudsamenstelling. Enquêtedata zijn gebaseerd op een steekproef en richten zich op percepties en ervaringen. De WoON-enquête bevat bijvoorbeeld vragen over de waardering van voorzieningen, gevoel van veiligheid, sociale cohesie en verhuishwensen.

Deze verrijking maakt de analyses inhoudelijk sterker, maar brengt ook nieuwe uitdagingen met zich mee. In de enquête is sprake van routing: respondenten krijgen alleen vervolgvragen die aansluiten bij hun eerdere antwoorden. Daardoor zijn veel waarden in de dataset niet ingevuld. Deze ontbrekende waarden zijn dus niet willekeurig, maar een direct gevolg van de vraagstructuur. Bovendien bestaat er binnen themablokken vaak een sterke samenhang tussen variabelen, bijvoorbeeld tussen verschillende vragen over overlast of voorzieningen. Dit kan de robuustheid van de modellen beïnvloeden en vraagt om aanvullende methoden.

Thematische benadering

Vanwege het grote aantal variabelen is bij deze analyse is gekozen voor een thematische aanpak, waarbij de WoON-variabelen zijn gegroepeerd in blokken als algemeen, verhuizen, leefbaarheid, zorg, zorggeschiktheid, VvE en energie. Binnen elk thema is vervolgens ook een selectie gemaakt van de meest waarschijnlijke kernvariabelen. Op die manier kon systematisch worden onderzocht welke thema's daadwerkelijk voorspellend vermogen hadden.

Eerste fase proces analyse 2: brede opname van variabelen

De analyse is gestart met een brede verkenning, waarin alle potentieel interessante WoON-variabelen tegelijk zijn opgenomen in het model, in combinatie met een brede selectie van de registratiedata. Hiermee werd nagegaan wat de totale winst in voorspelkracht zou zijn en welke thema's het meest zouden bijdragen. Voor analyse 2 is gekozen voor Gradient Boosting (XGBoost), omdat dit model automatisch kan omgaan met ontbrekende waarden. Het algoritme bepaalt tijdens het bouwen van beslisbomen hoe te splitsen wanneer informatie ontbreekt, zonder dat vooraf alle missende waarden hoeven te worden ingevuld.

Parallel is ook een bewerkte dataset gemaakt waarin de ontbrekende waarden wel zijn ingevuld. Voor numerieke variabelen is gebruikgemaakt van KNN-imputatie en voor categorische variabelen van de meest voorkomende categorie. Voor elke variabele is bovendien een extra indicator toegevoegd die vastlegde of een waarde oorspronkelijk ontbrak. Op die manier bleef ook de informatie over de missende waardes behouden.

De uitkomsten van deze eerste fase lieten een duidelijke verbetering zien: de voorspelkracht steeg van circa 55 procent in Analyse 1 naar bijna 70 procent in Analyse 2. Wanneer er een selectie wordt gemaakt van de meest relevante kernvariabelen per thema, verandert de voorspelkracht nauwelijks.

De resultaten lieten zien dat het thema leefbaarheid – met variabelen over tevredenheid met de buurt, ervaren veiligheid en sociale cohesie – het meest bepalend is. Ook verhuisintenties en verhuisbelemmeringen bleken een duidelijke invloed te hebben. Thema's als energie en VvE droegen nauwelijks bij, en ook zorggegevens leverden weinig extra verklaringskracht op.

Databewerking en imputatie

De aanwezigheid van routing en de daarmee samenhangende ontbrekende waarden maakte een zorgvuldige databewerking noodzakelijk. Naast het gebruik van de 'native handling' van missende waarden in XGBoost zijn aanvullende methoden getest. In de datasets zijn missende waarden ingevuld en zijn indicatoren toegevoegd om te markeren dat een waarde ontbrak.

Daarnaast zijn technieken voor dimensiereductie toegepast. Met factoranalyse is geprobeerd onderliggende factoren te identificeren die de samenhang tussen variabelen verklaren, zoals een factor "sociale cohesie" of "woningkwaliteit". Met Principal Component Analysis (PCA) zijn variabelen samengevat in een kleiner aantal componenten. Hoewel deze methoden nuttig bleken om patronen beter te interpreteren en de modellen stabiel te maken, leverden ze geen structurele verbetering van de modelprestaties op. De winst zat vooral in de inhoudelijke toevoeging van belevingsdata, niet in de technische bewerking.

Modeluitkomsten Analyse 2

Bij alle modellen is gebruik gemaakt van Gradient Boosting (XGBoost); de verschillen in voorspelkracht zijn toe te schrijven aan variaties in de inputvariabelen, zoals een selectie, bewerking of het gebruik van samengestelde factoren.

De resultaten laten zien dat de stap van registratiedata alleen naar de combinatie met WoON-variabelen een groot verschil maakt. Met uitsluitend registratiedata blijft de nauwkeurigheid steken rond de 50 tot 55 procent, dit komt logischerwijs overeen met de uitkomsten van analyse 1. Door WoON-data toe te voegen stijgt de nauwkeurigheid naar bijna 70 procent. Daarmee wordt duidelijk dat juist de toevoeging van percepties en belevingsgegevens van bewoners het model krachtiger maakt.

Uit de tabel blijkt ook dat niet alle thema's even zwaar wegen. Vooral het thema leefbaarheid (tevredenheid met buurt, ervaringen en sociale cohesie) is bepalend. Wanneer dit thema wordt weggelaten, daalt de nauwkeurigheid direct terug naar een niveau iets boven die van Analyse 1. Andere thema's, zoals energie en VvE, voegen daarentegen nauwelijks extra verklaringskracht toe.

Tabel 6.4 – Accuracy van verschillende modelvarianten in Analyse 2.

Modelomschrijving	Bewerking	Accuracy
Selectie Woonbase (analyse 1)		52,2%
Selectie Woonbase en algemene variabelen WoON		53,7%
Selectie Woonbase en selectie algemene variabelen WoON		52,8%
Selectie Woonbase en algemene en thema variabelen WoON		69,2%
Selectie Woonbase en selectie algemene en thema variabelen WoON		69,9%
Selectie Woonbase en selectie algemene en thema variabelen WoON	Thema	69,3%
Selectie Woonbase en selectie algemene en thema variabelen WoON	Thema en algemeen	69,8%
Selectie Woonbase en selectie algemene en thema variabelen WoON	Factoranalyse thema	67,7%
Selectie Woonbase en selectie algemene en thema variabelen WoON	PCA thema	67,5%
Selectie Woonbase en selectie algemene en thema variabelen WoON	PCA totaal	67,7%

Vergelijking modellen (accuracy)

Het toevoegen van WoON-variabelen leidt tot een toename in de voorspelkracht van het model. De verbeterde nauwkeurigheid moet echter genuanceerd worden. Hoewel de totale accuracy richting 70 procent stijgt (zie Tabel 6.4), blijft de balanced accuracy duidelijk achter. Dat betekent dat de modellen vooral goed zijn in het voorspellen van de meerderheidscategorieën — tevreden en zeer tevreden bewoners — maar moeite hebben met de minderheidsklassen, zoals neutraal en ontevreden. Deze groepen zijn beleidsmatig juist het meest interessant. Analyse 2 verkleint deze kloof enigszins door subjectieve belevingsvariabelen toe te voegen, maar lost het probleem niet volledig op.

Dit laat zien dat technische verbeteringen in de modelkeuze of parameterafstemming slechts beperkte effecten hebben. De grootste winst komt van de inhoudelijke verbreding van de data. Om ook de kleinere, negatieve groepen goed te kunnen modelleren, zijn aanvullende stappen nodig, zoals het gericht verzamelen van extra data of het inzetten van methoden die speciaal ontworpen zijn voor scheef verdeelde uitkomsten.

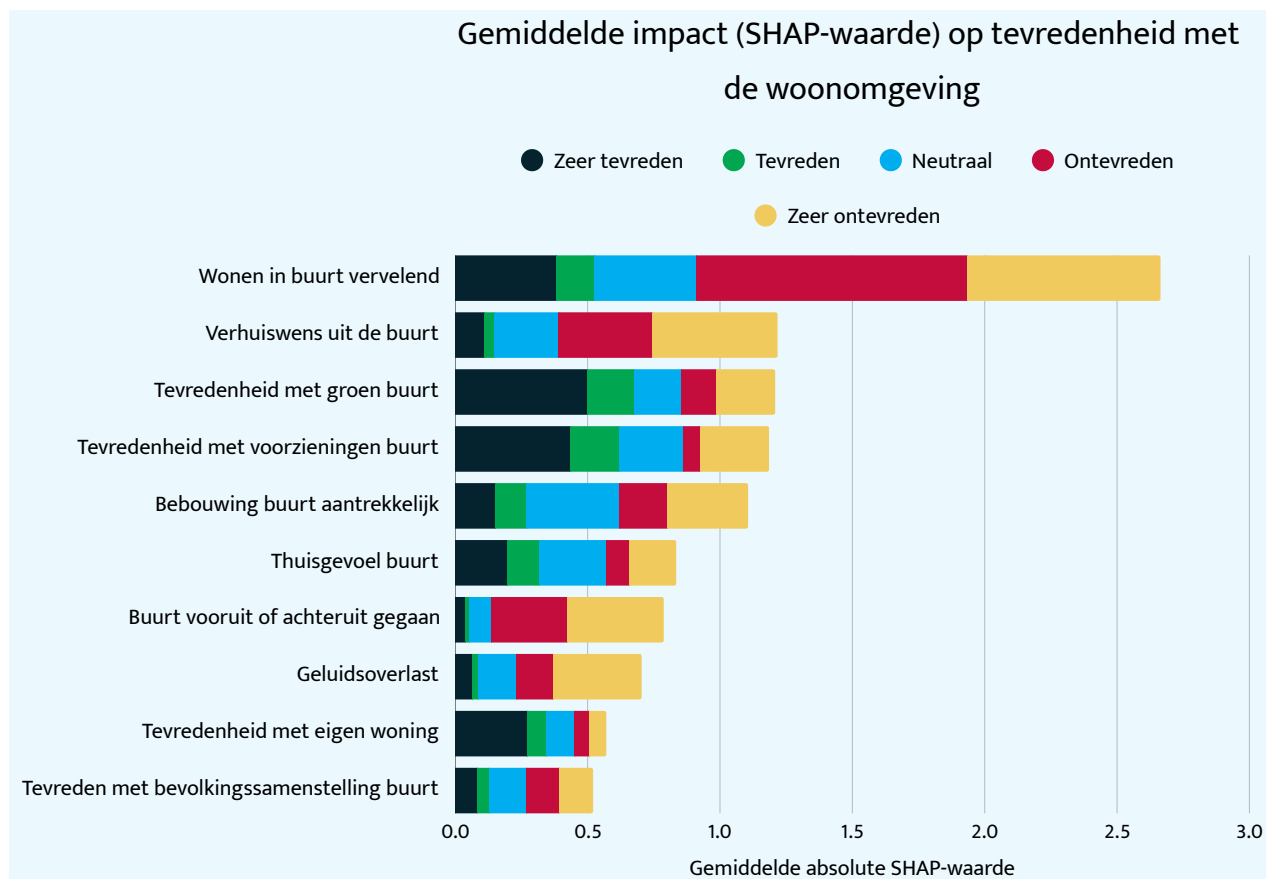
Belangrijkste voorspellers (SHAP-waarden) Analyse 2

Om inzicht te krijgen in de bijdrage van afzonderlijke variabelen is gebruikgemaakt van SHAP-waarden (Figuur 6.2). De resultaten tonen een duidelijke verschuiving ten opzichte van Analyse 1. Waar woningkenmerken toen domineerden, staan nu belevingsvariabelen centraal. Bewoners die hun buurt als positief ervaren, scoren significant hoger op tevredenheid. Ook tevredenheid met de woning en de voorzieningen in de buurt blijkt een doorslaggevende rol te spelen.

Daarnaast blijkt de verhuisdynamiek van belang. Bewoners die verhuishwensen hebben of verhuisbelemmeringen ervaren, geven vaker een lagere waardering voor hun woonomgeving.

Registratiefactoren zoals WOZ-waarde, woningtype en eigendomsvorm blijven relevant, maar hun relatieve gewicht is kleiner geworden. Zij verfijnen de voorspellingen, maar verklaren minder van het verschil dan de percepties van bewoners zelf.

Figuur 6.2 – Visualisatie van de 10 belangrijkste factoren (SHAP-waarde) per categorie in de variabele tevredenheid met de woonomgeving (van 'zeer ontevreden' tot 'zeer tevreden' (Analyse 2).



6.3 ANALYSE 3: TOEVOEGING VAN BUURTKENMERKEN

Inleiding

Waar Analyse 1 uitsluitend registratiedata gebruikte en Analyse 2 deze uitbreidde met belevingsgegevens uit de WoON-enquête, richt Analyse 3 zich op de vraag of buurtkenmerken extra verklaringskracht toevoegen. Het idee is dat leefbaarheid niet alleen wordt bepaald door kenmerken van woningen en huishoudens, maar ook door de context waarin mensen wonen: de dichtheid van de buurt, de inkomensverdeling of het aandeel sociale huur.

De verwachting was dat deze omgevingskenmerken de modellen zouden kunnen versterken of in sommige gevallen zelfs de plaats innemen van individuele variabelen.

Voor analyse 3 zijn er maximaal 205 indicatoren meegenomen in de analyse.

Proces

Voor deze analyse zijn buurtkenmerken uit de StatLine-tabel Kerncijfers Wijken en Buurten 2023 gekoppeld aan de datasets van Analyse 1 en 2. Het gaat om variabelen zoals bevolkingsdichtheid, gemiddelde WOZ-waarde van woningen, percentage koopwoningen, aandeel corporatiewoningen, arbeidsparticipatie, aandeel hoge inkomens en aandeel huishoudens rond het sociaal minimum.

Het proces bestond uit drie stappen:

1. Toevoegen van buurtkenmerken aan het model van Analyse 1 en vergelijken met de oorspronkelijke resultaten. Ook is onderzocht wat het effect is wanneer de samenhangende variabelen uit de Woonbase niet worden meegenomen.
2. Toevoegen van buurtkenmerken aan het model van Analyse 2 en beoordelen of er winst in voorspelkracht optreedt. Daarbij zijn in varianten enkele sterk gecorreleerde variabelen verwijderd, om multicollineariteit te beperken.
3. Opzetten van een model dat uitsluitend met buurtkenmerken werkt, om te onderzoeken in hoeverre deze kenmerken zelfstandig leefbaarheid kunnen verklaren.

Alle modellen zijn uitgevoerd met Gradient Boosting (XGBoost), in lijn met de eerdere analyses.

Accuracy

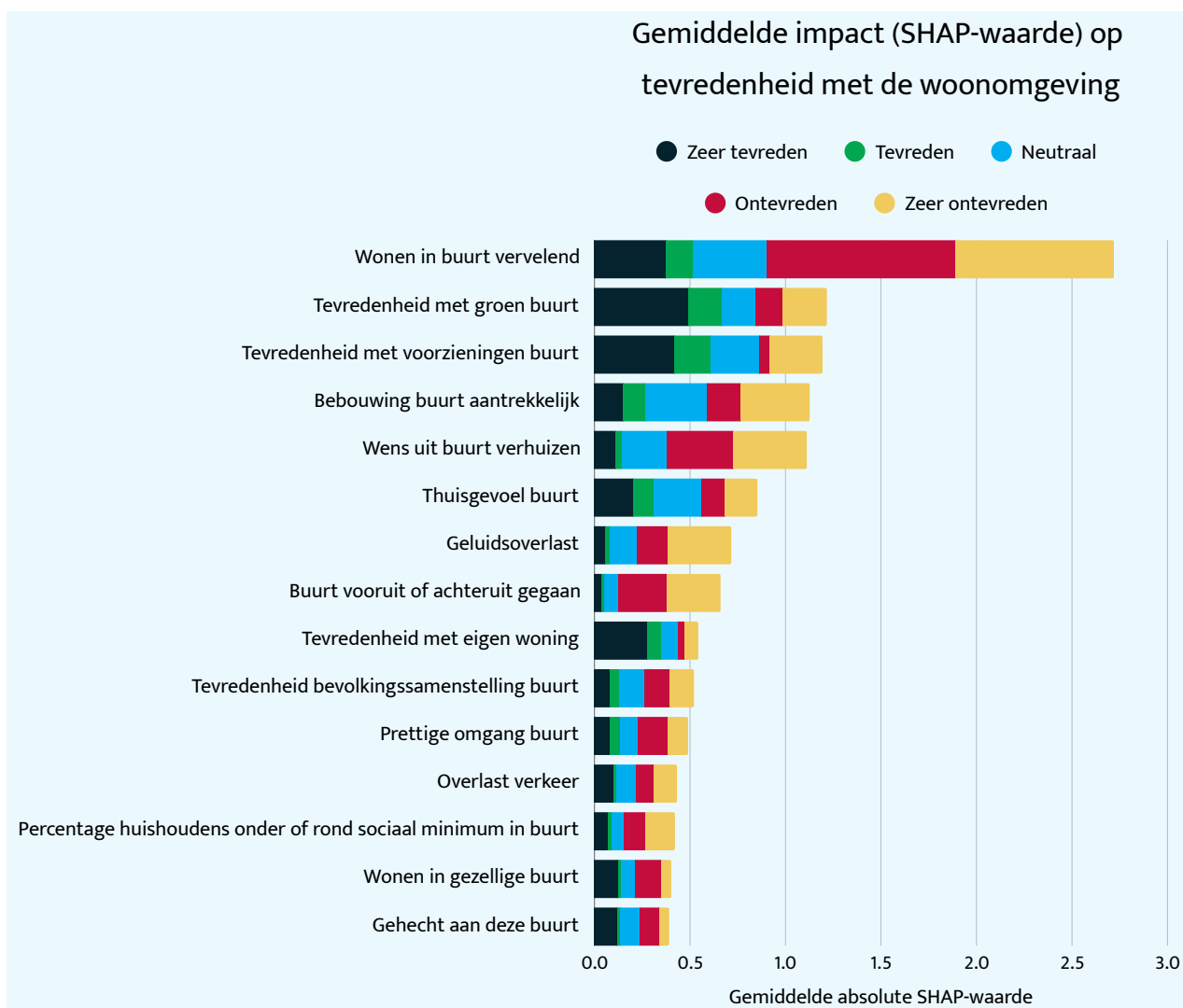
De uitkomsten laten zien dat het toevoegen van buurtkenmerken geen substantiële verbetering oplevert. Voor de basisvariant van Analyse 1 stijgt de nauwkeurigheid marginaal van 52,2% naar 52,4%. Bij Analyse 2 verandert de accuracy een klein beetje van 69,9% naar 69,1%, een marginaal verschil. Een model op basis van uitsluitend buurtkenmerken komt eveneens uit rond 52%, vergelijkbaar met Analyse 1.

Een opvallende variant is de combinatie van buurtkenmerken met Analyse 2, waarbij de leefbaarheidsvariabelen uit de WoON-enquête zijn weggelaten. In dat geval daalt de voorspelkracht fors, naar circa 59%. Dit bevestigt dat subjectieve belevingsdata een veel grotere rol spelen dan contextuele buurtgegevens.

Belangrijkste voorspellers (SHAP) in Analyse 3

De SHAP-analyse laat zien dat door het toevoegen van buurtkenmerken de rangorde van voorspellers slechts beperkt verandert. Individuele en beleavingsvariabelen – zoals ervaren veiligheid, tevredenheid met de buurt en verhuisintenties – blijven dominant. Buurtkenmerken zoals bevolkingsdichtheid of het aandeel corporatiewoningen schuiven in sommige varianten iets op in de top tien, maar hun gemiddelde effect blijft relatief klein.

Figuur 6.3 – Visualisatie van de 15 belangrijkste factoren (SHAP-waarde) per categorie in de variabele tevredenheid met de woonomgeving (van 'zeer ontevreden' tot 'zeer tevreden' (Analyse 3).



7. INZICHTEN

Het hoofdstuk 'Inzichten' brengt samen wat de analyses ons leren over de relatie tussen data, modellen en tevredenheid met de woonomgeving. Waar de eerdere hoofdstukken vooral de technische aanpak en methoden beschrijven, gaat het hier om de duiding: welke patronen worden zichtbaar, welke beperkingen komen naar voren en wat dit betekent voor het begrijpen van draagkracht en draaglast in wijken.

De drie uitgevoerde analyses vormen hierbij het vertrekpunt. In Analyse 1 is onderzocht wat registratiedata kunnen verklaren, in Analyse 2 is dit uitgebreid met subjectieve en thematische gegevens uit de WoON-enquête en in Analyse 3 komt de bredere buurtcontext in beeld. Deze aanpak sluit in eerste instantie vooral aan bij pijler 1: data en modellen. Tegelijkertijd kunnen de resultaten niet los worden gezien van de andere pijlers: de inzichten raken direct aan hoe bewoners hun leefomgeving beleven (pijler 2) en hoe beleidsmakers en professionals deze kennis kunnen benutten (pijler 3).

De inzichten in dit hoofdstuk vormen daarmee de brug tussen de technische uitwerking en de bredere betekenis van dit onderzoek. Ze laten zien hoe data en modellen elkaar aanvullen, waar de grenzen liggen en welke vervolgstappen nodig zijn om de balans tussen draagkracht en draaglast in wijken scherper in beeld te brengen.

Inzichten uit Analyse 1

De eerste analyse laat zien dat registratiedata een basis vormen om verbanden tussen objectieve omstandigheden en tevredenheid met de woonomgeving zichtbaar te maken. Variabelen zoals WOZ-waarde, woningtype en eigendomsvorm blijken belangrijke voorspellers van tevredenheid: ze geven een duidelijk signaal voor wie zich doorgaans prettig voelt in de woonomgeving.

Tegelijkertijd wordt de grens van de aanpak in Analyse 1 zichtbaar. Omdat het overgrote deel van de respondenten tevreden is, lukt het de modellen minder goed om de kleinere groepen met een neutrale of negatieve beleving te herkennen. Juist deze groepen zijn beleidsmatig van belang, maar blijven met registratiedata grotendeels buiten beeld. Het toevoegen van geavanceerdere methoden of het aanpassen van modelinstellingen verandert daar weinig aan. De beperkende factor ligt dus niet in de methode, maar in de aard en kwaliteit van de beschikbare data.

Het belangrijkste inzicht is daarom dat objectieve registratiedata onvoldoende zijn om draagkracht en draaglast goed te vatten. Ze geven vooral informatie over de context van woningen en huishoudens, maar missen de subjectieve laag die essentieel is voor het begrijpen van tevredenheid met de woonomgeving. Objectieve kenmerken zoals woningtype, eigendomsvorm en inkomenspositie verklaren slechts een deel van de ervaren tevredenheid en geven geen volledig beeld van hoe mensen hun woonomgeving beleven. Daarmee blijkt het niet haalbaar om voor heel Fryslân met deze data en methodiek betrouwbaar te voorspellen hoe tevreden mensen met hun woonomgeving zijn.

De analyse maakt duidelijk dat er behoefte is aan aanvullende gegevens die subjectieve beleving en context beter in kaart brengen. Factoren zoals sociale cohesie, veiligheid, buurtkwaliteit en woonwensen zijn noodzakelijk om draagkracht en draaglast in wijken scherper te duiden. Deze uitbreiding is opgepakt in Analyse 2 en 3, waarin de WoON-data over beleving en woonervaringen nadrukkelijker worden betrokken.

Inzichten uit Analyse 2

Analyse 2 laat overtuigend zien dat de toevoeging van subjectieve belevingsdata een meerwaarde heeft bij het voorspellen van tevredenheid met de woonomgeving. Waar de modellen in Analyse 1 beperkt bleven tot woning- en persoonskenmerken, ontstaat door de WoON-enquête een rijker en completer beeld. Met name variabelen over voorzieningen, geluidsoverlast, sociale cohesie, tevredenheid met de buurt en verhuiscriteria blijken essentieel. Deze factoren maken duidelijk dat leefbaarheid niet uitsluitend een kwestie is van objectieve omstandigheden, maar sterk samenhangt met hoe bewoners hun omgeving beleven.

De winst in voorspelkracht is substantieel: de nauwkeurigheid stijgt van bijna 55 procent naar bijna 70 procent. Toch blijft een belangrijke beperking overeind. De modellen voorspellen vooral goed wie tevreden is, maar onderscheiden de kleinere groepen ontevreden bewoners nog onvoldoende. Dit betekent dat juist de bewoners die beleidsmakers het meest willen bereiken — de groepen die minder gelukkig zijn met hun woonomgeving — minder goed in beeld komen.

Wat dit inzichtelijk maakt, is dat verdere verbetering niet alleen gezocht moet worden in de keuze van modellen of technische bewerkingen, maar vooral in de inhoudelijke verbreding van de data. Daarmee laat Analyse 2 zien dat de combinatie van de objectieve structuur van woningen en huishoudens en de subjectieve ervaringen van bewoners informatie geeft.

Inzichten uit Analyse 3

Analyse 3 maakt duidelijk dat buurt- en wijkenkenmerken, zoals bevolkingsdichtheid, inkomensverdeling of het aandeel corporatiewoningen, slechts in beperkte mate bijdragen aan het verklaren van tevredenheid met de woonomgeving. De toevoeging van deze contextvariabelen aan de modellen leverde nauwelijks verbetering in voorspelkracht op.

Dit resultaat onderstreept dat tevredenheid met de woonomgeving in de eerste plaats samenhangt met persoonlijke en huishoudelijke kenmerken en, nog belangrijker, met de subjectieve beleving van bewoners. Het zijn variabelen zoals tevredenheid met de buurt, ervaren veiligheid en verhuiscriteria die de grootste verklarende waarde hebben.

De buurtkenmerken geven wel aanvullende context en kunnen helpen bij de interpretatie van verschillen tussen gebieden, maar blijken niet doorslaggevend voor de voorspellingen. Daarmee bevestigt Analyse 3 dat objectieve en subjectieve gegevens op individueel niveau zwaarder wegen dan bredere omgevingsstatistieken.

Vergelijking met Leefbaarometer

In de herijking van de Leefbaarometer 3.0 (<https://www.leefbaarometer.nl/home.php>) is uitvoerig onderzocht hoe goed het model de ervaren leefbaarheid van bewoners kan voorspellen. Daarbij is gekeken naar de relatie tussen administratieve indicatoren (zoals woningvoorraad, voorzieningen, veiligheid en overlast) en de antwoorden uit bewonersenquête's. Er is gekeken naar verklaarde variantie. Dit geeft aan welk deel van de verschillen (variantie) in de uitkomst door een model kan worden verklaard. Stel dat bewoners in buurten sterk uiteenlopende rapportcijfers geven voor hun leefomgeving: de verklaarde variantie laat zien hoeveel van dat verschil kan worden teruggevoerd op de factoren in het model (bijvoorbeeld woningkenmerken, inkomenspositie of ervaren veiligheid). Een waarde van 70% verklaarde variantie betekent dat het model zeven van de tien verschillen tussen buurten of bewoners kan begrijpen en toeschrijven aan meetbare indicatoren, terwijl 30% onverklaard blijft en mogelijk te maken heeft met individuele voorkeuren, contextfactoren of toevallige variatie.

Uit de validatie blijkt dat de Leefbaarometer in staat is om ongeveer 71% van de variantie in de buurtoordelen van bewoners te verklaren. In sommige stedelijke contexten, zoals in Utrecht, stijgt dit percentage zelfs tot bijna 80% bij buurten met voldoende respondenten. Dit is te zien op pagina 88 in de rapportage 'Leefbaarometer 3.0 Instrumentontwikkeling', beschikbaar op de website van de Leefbaarometer. De methodiek in de Leefbaarometer is gebaseerd op regressieanalyse: subjectieve oordelen worden gebruikt om te bepalen welke objectieve variabelen de sterkste voorspellers zijn. De beleving zelf verdwijnt daarna uit het model; de uiteindelijke index wordt volledig opgebouwd uit administratieve data.

In dit experiment is een andere aanpak gevolgd. Hier zijn drie modellen getest: één gebaseerd op registratiedata, één waarin subjectieve belevingsdata uit de WoON-enquête zijn toegevoegd en een derde waarin ook buurtkenmerken zijn meegenomen. De stap van louter registratiedata naar de combinatie met subjectieve beleving bleek cruciaal. Waar het eerste model bleef steken op circa 55% accuracy, steeg de voorspelkracht in het tweede model naar bijna 70%. Daarmee laat dit experiment zien dat de directe opname van bewonersbeleving in voorspelmodellen een substantiële meerwaarde heeft. Daarbij maken we gebruik van machine learning-methoden (Random Forest, Gradient Boosting), die complexe patronen kunnen leren zonder strikte aannames vooraf. Deze aanpak levert een vergelijkbare verklaarde variantie van rond de 70%, maar via een andere route: niet alleen registratiedata, maar een combinatie van 'harde cijfers' en bewonersbeleving.

Een vergelijking tussen beide benaderingen laat zien dat ze qua verklaarde variantie dicht bij elkaar liggen. De Leefbaarometer bereikt hoge scores door administratieve indicatoren zorgvuldig te selecteren en te wegen op basis van hun samenhang met leefbaarheidsoordelen. Dit experiment laat zien dat het opnemen van subjectieve data in het model tot een vergelijkbare voorspelkracht leidt. Beide methoden zijn dus in staat om een groot deel van de variatie in ervaren leefbaarheid te duiden, maar geen van beide verklaart deze volledig. Dit onderstreept dat er altijd een deel van leefbaarheid samenhangt met unieke lokale context en individuele ervaring die niet volledig te vangen is in indicatoren of modellen.

Conclusie

De drie analyses laten samen zien hoe verschillend soorten data bijdragen aan het begrijpen van tevredenheid met de woonomgeving. Registratiedata vormen een nuttige basis, maar hun verklaringskracht blijft beperkt: ze herkennen vooral de grote groep tevreden bewoners en geven weinig inzicht in de kleinere groepen die neutraal of ontevreden zijn. Belevingsdata uit enquêtes blijken cruciaal. Zij voegen de subjectieve dimensie toe die onmisbaar is om patronen van draagkracht en draaglast zichtbaar te maken, en verhogen de voorspelkracht van de modellen aanzienlijk. Contextvariabelen op buurtniveau dragen nauwelijks zelfstandig bij, maar geven wel extra stabiliteit en achtergrond bij de interpretatie van verschillen tussen gebieden.

Aantal meegenomen indicatoren verschillen ook in de drie analyses: van 59 indicatoren in analyse 1 naar 719 indicatoren in analyse 2. Er zijn tevens varianten gebruikt met ongeveer 200 variabelen (analyse 3). Algemene tendens is dat meer data niet zozeer een beter model genereert.

Daarnaast voorspellen de modellen ongeveer 70% van de data correct. Als gekeken wordt naar de ruwe data van tevredenheid met de woonomgeving is te zien dat ongeveer 83% van de respondenten van de WoON-enquête tevreden of zeer tevreden met de woonomgeving is. De modellen benaderen de 83% niet en laten nog onverklaarde ruimte zien, ook als iedereen als tevreden zou zijn voorspeld. Inzichten vanuit andere bronnen of data zouden hier in de toekomst nog een rol in kunnen spelen. De Leefbaarometer, die administratieve indicatoren gebruikt als proxy's voor bewonersoordelen, en het hier-gebruikte model dat objectieve registraties en subjectieve beleving combineert, leveren beide vergelijkbare resultaten met een verklaarde variantie van rond de 70%.

De gezamenlijke conclusie is dat tevredenheid met de woonomgeving niet goed te verklaren is vanuit objectieve registraties alleen. Pas in de combinatie van objectieve kenmerken met subjectieve beleving ontstaat een robuust model dat de werkelijkheid van bewoners beter benadert. Buurtkenmerken zijn ondersteunend, maar niet bepalend. Voor beleid betekent dit dat het meten en begrijpen van leefbaarheid vraagt om een geïntegreerde benadering waarin harde registraties en zachte belevingen structureel worden gecombineerd.



COLOFON

© 2025 Coöperatie DataFryslân U.A.

Dit rapport is tot stand gekomen dankzij de inzet en bijdragen van het projectteam: Nynke Slagter, Anita Nguyen, Sjaak Moerman, Gerian Kuiper, Siebren Kooistra, Ruud Leerink en Carlos de Matos Fernandes.

Deze publicatie maakt deel uit van het project 'Draagkracht en Draaglast in buurten, wijken en dorpen'. Deze uitgave mag worden geciteerd en gedeeld onder voorwaarde van bronvermelding.