



WAGENINGEN UR
For quality of life

Wageningen IMARES

Habitatmodellen in het beheer: Zijn state-of-the-art modellen voor kokkels in de Westerschelde bruikbaar voor beheer en beleidsbesluiten?

J. Steenbergen E. Meesters

Rapport nr. C091/06
December 2006



Wageningen IMARES is een
samenwerkingsverband tussen
Wageningen UR en TNO



Wageningen IMARES

Institute for Marine Resources & Ecosystem Studies

Vestiging IJmuiden
Postbus 68
1970 AB IJmuiden
Tel.: 0255 564646
Fax: 0255 564644

Vestiging Yerseke
Postbus 77
4400 AB Yerseke
Tel.: 0113 672300
Fax: 0113 573477

Vestiging Den Helder
Postbus 57
1780 AB Den Helder
Tel.: 022 363 88 00
Fax: 022 363 06 87

Vestiging Texel
Postbus 167
1790 AD Den Burg Texel
Tel.: 0222 369700
Fax: 0222 319235

Internet: www.wageningenimares.wur.nl
E-mail: imares@wur.nl

Rapport

Nummer: C091/06

Habitatmodellen in het beheer: Zijn state-of-the-art modellen voor kokkels in de Westerschelde bruikbaar voor beheer en beleidsbesluiten?

J. Steenbergen E. Meesters

Opdrachtgever:

Ministerie van LNV
Directie regionale zaken (zuid)
Postbus 6111
5600 HC Eindhoven

Project nummer:

4392500001

Aantal exemplaren:	25
Aantal pagina's:	59
Aantal tabellen:	2
Aantal figuren:	13
Aantal bijlagen:	5

Wageningen IMARES is een samenwerkingsverband tussen Wageningen UR en TNO. Wij zijn geregistreerd in het Handelsregister Amsterdam nr. 34135929 BTW nr. NL 811383696B04



De Directie van Wageningen IMARES is niet aansprakelijk voor gevolgschade, alsmede voor schade welke voortvloeit uit toepassingen van de resultaten van werkzaamheden of andere gegevens verkregen van Wageningen IMARES; opdrachtgever vrijwaart Wageningen IMARES van aanspraken van derden in verband met deze toepassing.

Dit rapport is vervaardigd op verzoek van de opdrachtgever hierboven aangegeven en is zijn eigendom. Niets van dit rapport mag weergegeven en/of gepubliceerd worden, gefotokopieerd of op enige andere manier zonder schriftelijke toestemming van de opdrachtgever.

Inhoudsopgave

Inhoudsopgave	2
Samenvatting	4
Summary.....	5
1 Inleiding	6
1.1 Belang habitatmodellen voor beheer.....	6
1.2 Case study: Kokkels in de Westerschelde.....	6
1.3 Veel gebruikte habitatmodellen.....	7
1.4 Gewenste eigenschappen van een goed model.....	8
2 Materiaal en methoden	9
2.1 Gegevens	9
2.1.1 Biotisch	9
2.1.2 Abiotische data	9
2.2 Beschrijving van de gegevens	10
2.3 Modelontwikkeling	10
2.3.1 Regressie- en classificatiebomen	11
2.3.2 Generalized Linear Models	12
2.3.3 Quantielregressie.....	13
2.3.4 General Additive models (GAM)	14
2.4 Vergelijking van de verschillende technieken.....	15
2.4.1 Mate van voorspelbaarheid: Hoeveelheid verklaarde variatie.....	15
2.4.2 Ruimtelijke patronen: GIS-plaatjes.....	15
2.4.3 Multi dimensional Scaling (MDS) om modeluitkomsten te vergelijken.....	16
2.4.4 Sensitiviteit & specificiteit	16
3 Resultaten.....	18
3.1 Verkenning van de data	18
3.1.1 Kokkels	18
3.2 Verkenning van de data	19
3.3 Regressie- en classificatieboom analyses	20
3.3.1 Log getransformeerde data (DN).....	20
3.3.2 Poisson Regressieboom	21
3.3.3 Classificatieboom	22
3.4 Generalized linear models (GLM)	23

3.4.1	Kokkeldichtheid (Poisson regressie)	23
3.4.2	Aanwezig- afwezigheid	23
3.5	Quantiel regressie	24
3.6	General Additive models (GAM)	24
4	Vergelijking van de verschillende technieken	25
4.1	Beschrijvend vermogen	25
4.2	Voorspellend vermogen.....	27
4.2.1	MDS	30
4.2.2	Sensitiviteit & Specificiteit	31
5	Discussie	32
6	Literatuur	40
7	Bijlagen	43

Samenvatting

Habitatmodellen zijn een veel gebruikt middel om de impact van inrichtingsmaatregelen op soorten in te kunnen schatten. Bij dergelijke modellen wordt de verspreiding van een organisme verklaard aan de hand van verschillende omgevingsfactoren. Een vereiste is dat ze ook in veranderende situaties nog een goede voorspelling kunnen geven van de verspreiding van soorten. Over het algemeen wordt echter weinig aandacht besteed aan het zogenaamde voorspellende vermogen. Wageningen IMARES heeft van LNV-Directie Regionale zaken (zuid) nu de opdracht gekregen om (een visie) op beleidsondersteunende tools in de vorm van habitatmodellen te ontwikkelen. Kookeldichtheden uit het zuiden van de Westerschelde zijn gebruikt om modellen mee te maken die de dichtheden relateren aan abiotische factoren. Met deze modellen is daarna onderzocht in hoeverre zij toepasbaar zijn in andere gebieden door de kookeldichtheden in het noorden van de Westerschelde te voorspellen. Verschillende statistische modelleringstechnieken zijn gebruikt: "Generalized Linear Models", "Generalized Additive Models", "Regression and Classification Trees" en Quantiel Regressie. De verschillende modellen zijn gebruikt om dichtheden, dichtheidsklassen en aan/afwezigheid van kokkels te voorspellen. De meeste modellen geven redelijke tot goede beschrijvingen van kookeldichtheden voor het gebied waaruit de data stammen (zuiden van de Westerschelde). Jammergenoeg blijkt voor alle modellen dat het voorspellend vermogen laag is indien de modellen gebruikt worden om de kookeldichtheden in het noorden van de Westerschelde te voorspellen. Verschillende criteria zijn gebruikt om het voorspellend vermogen van de verschillende modellen te evalueren. Hieruit blijkt dat er niet 1 methode is die duidelijk superieur is: hoe nauwkeuriger de gewenste voorspelling, hoe groter de onzekerheid die met de voorspelling gepaard gaat. Een nauwkeurige aantalsvoorspelling heeft een hoge mate van onzekerheid, terwijl dit voor een voorspelling op basis van aan/afwezigheid minder is.

Voor de criteria die gebruikt zijn om de verschillende modelleringstechnieken te vergelijken bleek dat de specificiteit/sensitiviteitsmethode de beste evaluatie gaf voor een voorspelling van aan/afwezigheid van kokkels. Voor meer kwantitatieve schattingen (dichtheidsklassen of dichtheden) lijkt een multivariate aanpak door middel van "non-metric MultiDimensional Scaling", kortweg MDS genoemd een nieuwe aanpak te zijn die goede mogelijkheden biedt voor het vergelijken van verschillende technieken. Beide methoden geven aan dat een poisson regressie (GLM) met kwadratische termen en een poisson "Regression Tree" zonder kwadratische termen de beste voorspellingen geven.

Gezien de beschikbare gegevens moet geconstateerd worden dat modellen voor accurate voorspellingen van kookeldichtheden moeilijk te maken zijn. Met de beschikbare gegevens is het verstandiger om voorspellingen te doen in de vorm van 2-4 dichtheidsklassen of in de vorm van aan/afwezigheid.

Summary

The use of habitat models is an important tool in predicting the human-induced impacts on the distribution of species. These models predict the likely occurrence or distribution of species in relation to environmental or habitat variables. A prerequisite is that the models are also valid in changed environments. However, in most studies little attention is paid to the so called predictive performance of these models. In this study (a vision on) tools in the form of habitat models for the use of conservation planning is developed on request of the Ministry of LNV. We used information on densities of cockles in the southern part of the Westerschelde for the development of the models and several statistic modelling techniques are being evaluated: Generalized Linear Models, Generalized Additive Models, Regression and Classification Trees and Quantiel regression. Cockle data from the northern part of the Westerschelde is used to test the predictive performance.

Most of the models gave a reasonable or good description of the density of cockles in the southern part of the Westerschelde. Unfortunately when the models were used to predict the densities of cockles in the northern part of the Westerschelde the predictive capacity appears to be quite low for all models. We used several different criteria to evaluate the predicting capacity of the models. It turned out that there was not one model that was clearly better than the rest: the more precise the desired prediction, the larger uncertainty. A precise prediction of density has a high degree of uncertainty, whereas a prediction of the presence or absence is much more convenient.

With the sensitivity/specificity method it is possible to evaluate the predictive performance of the models in terms of presence/absence of cockles. For more quantitative estimates (density classes or densities) a multivariate approach through non-metric MultiDimensional Scaling (MDS) seems to be a good way to compare several techniques. Both methods indicate that a poisson regression (GLM) and a poisson regression tree give the best predictions.

Given the available data we have to conclude that models for accurate prediction of cockle densities are still difficult to make. It would therefore be wiser to make predictions for 2-4 density classes or in the form of presence/absence.

1 Inleiding

1.1 Belang habitatmodellen voor beheer

In de Nederlandse kustwateren zijn veel beleidsdoelen gericht op het in stand houden van soorten en habitats. Voor een goede uitvoering van het beheer van gebieden, vooral waar het gaat om inrichtingsvraagstukken, is het dan ook van groot belang om de effecten van habitatveranderingen op het voorkomen van soorten goed in te kunnen schatten. Dit betekent dat kennis nodig is over de relatie tussen habitat en voorkomen van soorten zodat veranderingen in het milieu kunnen worden doorvertaald naar veranderingen in dichtheden, verspreiding en biodiversiteit.

Een eerste stap in de richting van deze kennis is vaak een studie waarbij de verspreiding van organismen bestudeerd wordt ten opzichte van de verspreiding van omgevingsfactoren, waarvan wordt verondersteld dat ze van belang zijn voor de desbetreffende soort en die ook relevant zijn voor beheer. Een formele samenvatting van deze kennis wordt vaak gegeven in de vorm van habitatmodellen (Rykiel, 1996; Guisan & Zimmerman, 2000).

De vraag die echter over het algemeen onderbelicht wordt is: hoe bruikbaar zijn deze modellen eigenlijk en hoe betrouwbaar zijn ze wanneer ze worden toegepast in een andere situatie?

Wageningen IMARES heeft van LNV, Directie Regionale Zaken, binnen een beleidsondersteunend onderzoeksprogramma (BO-02, thema EHS) de opdracht gekregen om (een visie op) beleidsondersteunende tools in de vorm van habitatmodellen te ontwikkelen. Doel van deze rapportage is meer inzicht te verschaffen in de mogelijkheden en beperkingen van habitatmodellen in relatie tot effectenstudies van inrichtings- en herstelmaatregelen en menselijke activiteiten zoals in de Westerschelde. Verschillende methodes komen aan bod en zullen met elkaar worden vergeleken. Kokkels vormen hierbij een ideale doelsoort omdat er een uitgebreide dataset beschikbaar is bij Wageningen Imares. Bovendien zijn kokkels qua biomassa dominant en een belangrijke voedselbron voor steltlopers zoals de scholekster (Graveland, 2005).

1.2 Case study: Kokkels in de Westerschelde

De analyse zal zich in dit onderzoek beperken tot kokkels in de Westerschelde. De Westerschelde maakt onderdeel uit van het Schelde-estuarium. Waar de overige Deltawateren deels of geheel zijn afgesloten van zee en/of rivier staat de Westerschelde nog steeds in verbinding met zowel de zee als de rivier de Schelde. De hoge dynamiek (het getijverschil bedraagt tot 6m), de overgang van zoet naar zout en de hieraan gebonden flora en fauna,

maken het Schelde-estuarium tot een belangrijk natuurgebied. De Schelde is echter ook economisch van grote betekenis als verkeersslagader voor het Antwerpse havengebied. De toegankelijkheid van de haven van Antwerpen is al lang een gevoelig onderwerp tussen Nederland en België (Vlaanderen). Langzamerhand is wederzijds erkend dat het beleid en beheer van het estuarium een gemeenschappelijke zaak is die in goed nabuurschap moet worden vorm gegeven.

1.3 Veel gebruikte habitatmodellen

Wereldwijd wordt er veel aandacht besteed aan het maken van habitat-/habitatgeschiktheidsmodellen (Guisan & Zimmerman, 2000). Veel gebruikte technieken voor univariate (d.w.z. enkelvoudige) responsvariabelen zijn gewone lineaire regressie, logistische regressie, Poisson regressie, regressie en classificatie bomen, neurale netwerken, fuzzy-logic modellen en, meer recent, additieve modellen, generieke additieve modellen ('Generalised Additive Models' (GAM), generieke lineaire modellen ('Generalised Linear Models' (GLM), quantiel regressie en, meest recent, de 'mixed-effect' versies van deze modeltechnieken, GAMM en GLMM (e.g. Amar & Redpath 2005), die rekening houden met de ruimtelijke of temporele variatie van de monsters. Daarnaast zijn er nog multivariate technieken die gebruikt kunnen worden om het voorkomen van bepaalde ecotopen/ecosystemen/clusters te voorspellen (e.g. (Zeman P, Lynen G 2006).

De meeste modellen voor bodemdieren zijn gemaakt met behulp van Poisson regressie (*Generalized Linear Models, GLM*) met zowel lineaire als kwadratische termen voor de verklarende variabelen. Het idee achter dergelijke modellen is dat van een dosis-respons relatie, waarbij het voorkomen van een enkele soort (aantal per m²; g.m⁻²; presence/absence) wordt gezien als respons waarin men een duidelijke (veelal klokvormige, hyperbolische) curve kan waarnemen onder invloed van 1 onafhankelijke/abiotische variabele (§ 2.3.2). Een dergelijke response kan voor meerdere abiotisch variabelen bestaan, waardoor de waargenomen respons een optelsom wordt van meerdere responscurves.

In deze rapportage besteden wij tevens aandacht aan het gebruik van regressie & classificatie bomen, quantiel regressie en 'generalized additive models' (GAMs).

Regressie en classificatie bomen verklaren de variatie van een responsvariabele (voorkomen van een soort) door de data herhaaldelijk op te splitsen in homogene groepen (§ 2.3.1). Hierbij wordt gebruik gemaakt van een combinatie van verklarende (in dit geval abiotische) variabelen (De'ath & Fabricus, 2000).

Quantiel regressie is een methode om functionele relaties tussen variabelen te schatten voor verschillende delen (quantielen) van de distributie van de afhankelijke variabele (§ 2.3.3; Cade & Noon, 2003).

Generalized Additive Models (GAM) zijn vergelijkbaar met GLMs, maar zonder een lineaire relatie te veronderstellen tussen de respons en de onafhankelijke variabelen (§ 2.3.4). GAMs maken

gebruik van zogenaamde 'smoothers' waarbij een vloeiende lijn tussen de datapunten wordt geschat door steeds subselecties van de data te nemen (Thuiller et al, 2003) .

1.4 Gewenste eigenschappen van een goed model

Een goed model moet niet alleen de relaties tussen een afhankelijke variabele en de onafhankelijke variabele(n) zo goed mogelijk **beschrijven**, maar het ook mogelijk maken om goede **voorspellingen** te doen over de verwachte grootte van de afhankelijke variabele bij nieuwe waarden van de onafhankelijke variabelen. De modellen zullen in deze rapportage dan ook worden getest op:

- **Beschrijvend vermogen:** hierbij wordt gekeken of er in het gebied waar het model ontwikkeld is sprake is van een goede fit met de observaties. Het proces om te komen tot een zo goed mogelijke beschrijving wordt ook wel calibratie genoemd.
- **Voorspellend vermogen:** toepasbaar in veranderende situaties: er wordt gekeken of er op een redelijke afstand van het domein waar de modellen op gebaseerd zijn een goede voorspelling van de dichtheden wordt gegeven door de modellen. Het model wordt dus in een ander gebied toegepast als waarvoor het ontwikkeld is. Het analyseren of een model ook een goed voorspellend vermogen heeft wordt ook wel validatie genoemd.

2 Materiaal en methoden

2.1 Gegevens

Voor het ontwikkelen van de modellen zijn biotische gegevens (voorkomen van de kokkels) en abiotische gegevens uit de Westerschelde gebruikt.

2.1.1 *Biotisch*

Voor de modelontwikkeling is gebruik gemaakt van de aantallen van 1-jarige kokkels, zoals waargenomen tijdens de voorjaarssurveys van Wageningen IMARES in het litoraal van de Westerschelde (2002-2006). Tijdens deze surveys worden met behulp van een kokkelschuif of een steekbuis (Figuur 1) monsters uit de bodem gehaald. De inventarisatie is vooral gericht op de droogvallende platen en slikken. Jaarlijks worden zo'n 250 stations bemonsterd in de periode maart-mei (Figuur 2). Het litoraal van de Westerschelde is daarmee volledig gedekt. De aldus verzamelde monsters worden gezeefd en uitgezocht. Indien nodig wordt een deelmonster genomen op basis van volume. Aanwezige kokkels worden gesorteerd naar leeftijd (1-jarig, 2-jarig, meerjarig), geteld en gewogen (g versgewicht). Voor een uitgebreidere beschrijving van deze opnamemethode wordt verwezen naar Bult et al (2004).

Omdat de dichtheden van met name oudere kokkels beïnvloed zou kunnen zijn door visserij, is besloten om enkel 1-jarige kokkels voor de analyse te gebruiken (zie ook Steenberg et al (2004) voor effecten van visserij op 1-jarige kokkels).



Figuur 1 steekbuis (links) en kokkelschuif (rechts) die worden gebruikt voor de IMARES-survey (Bult et al, 2004).

2.1.2 *Abiotische data*

De volgende abiotische gegevens zijn in de vorm van gisbestanden (Arcgis) beschikbaar gesteld door het RIKZ (van der Male & Schouwenaar, 2003):

- Hoogteligging (cm t.o.v. NAP) in 2002, resolutie van de gridcellen: 20 bij 20 meter. Gegevens over de hoogteligging variëren van +233 NAP tot -1206 cm. Om te voorkomen dat er negatieve waarden van diepte in de modellen moesten worden gebruikt hebben we

de oorspronkelijke dieptegegevens getransformeerd op de volgende wijze. De hoogteligging is omgerekend naar diepte middels $\text{diepte} = -(\text{hoogte} + 233)$.

- Gemiddeld zoutgehalte (Saliniteit in psu) in 1992, resolutie van de gridcellen: 100 bij 100 meter.
- Maximale stroomsnelheid bij springtij (cm.s^{-1}) in 2001, berekend met model scalwest 1996, resolutie van de gridcellen: 20 bij 20 meter.

Er bleken geen voor onze data accurate sedimentgegevens beschikbaar te zijn en daarom hebben wij besloten om in het voorjaar van 2006 zelf sedimentmonsters te nemen op dezelfde locaties als waar ook de kokkels worden bemonsterd.

De monsters zijn gevriesdroogd en geanalyseerd op het NIOO CEME in Yerseke met de Malvern Mastersizer 2000. Belangrijke instellingen zijn:

Calculation: Mie theory- general purpose – enhanced sensitivity

Settings: Refractive index (R.I.) = 1.52

Absorptie factor = 0.1 with Blue light correction

Oplosmiddel: water – RI=1.33

Premeasurement: 10 % ultrasoon – 60 sec

Voor deze studie is gekozen voor de mediane korrelgrootte (in phi) als indicator voor de eigenschappen van het sediment. Het voordeel van deze metriek is dat het beide kanten van de korrelgrootteverdeling (van grof zand tot slib) met een gelijke resolutie weergeeft (Steenbergen & Escaravage, 2006).

2.2 Beschrijving van de gegevens

Voor de jaren 2002-2006 zijn per jaar de totale dichtheden van 1-jarige kokkels berekend. Om een beeld te krijgen van de relaties tussen de omgevingsfactoren onderling en de kokkeldata zijn deze grafisch tegen elkaar uitgezet (§ 3.2).

2.3 Modelontwikkeling

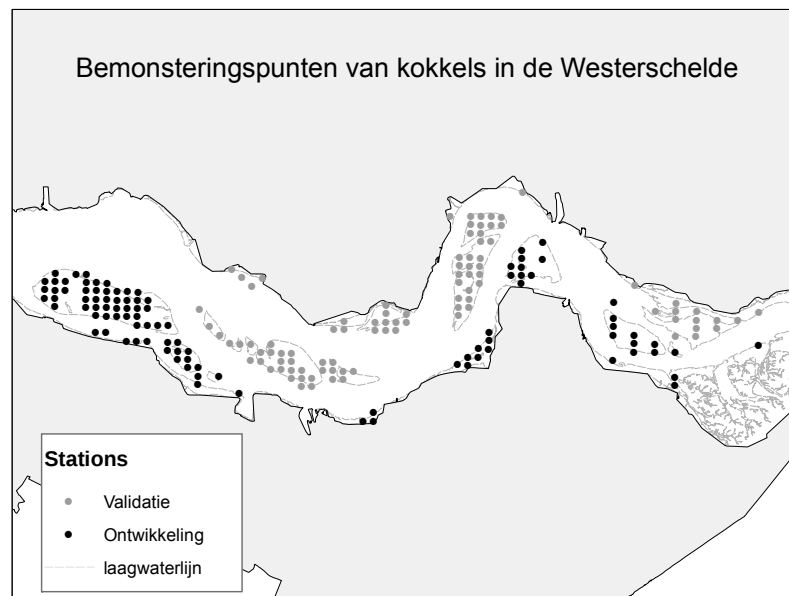
Voorafgaand aan de analyses werden de gemiddelde aantallen 1-jarige kokkels in de periode 2002-2006 per station uitgerekend. Abiotische gridbestanden zijn met behulp van ArcGIS 9 aan de stations gekoppeld. Locaties waar geen informatie was van een of meer abiotische variabelen zijn niet meegenomen in de analyse. Omdat tijdens de modelontwikkeling bleek dat één station in het zuiden, waar een zeer hoge gemiddelde dichtheid ($>1000/\text{m}^2$) aan kokkels was gemeten onevenredig veel invloed had op de uitkomsten van de modellen is besloten deze niet mee te nemen in de analyse.

Uiteindelijk bleven er 194 stations met een complete abiotische dataset over voor de analyses.

Om te kunnen testen hoe goed de modellen werken in een ander gebied zijn de data in tweeën gesplitst (Figuur 2). Hierbij werden de stations in het zuiden (98 stations) gebruikt voor modelontwikkeling en de punten in het noorden (96 stations) voor validatie.

Modellen zijn geconstrueerd door middel van verschillende technieken:

- Regressie en classificatiebomen voor zowel dichtheden als aan- of afwezigheid (§ 2.3.1)
- Generalized lineair models voor zowel dichtheden als aan- of afwezigheid (§2.3.2)
- Quantiel regressies voor dichtheden (§ 2.3.3)
- General Additive Models (§ 2.3.4)



Figuur 2 Overzicht van de bemonsterde stations van kokkels in de Westerschelde. Zwarte punten: modelontwikkeling (zuid; 98 stations), grijze punten: validatie (noord; 96 stations)

2.3.1 Regressie- en classificatiebomen

Regressiebomen delen de responsvariabele steeds op in twee groepen bij een bepaalde waarde van een van de afhankelijke (omgevings-)variabelen waarbij de variatie binnen de twee groepen zo klein mogelijk wordt gehouden en tussen de twee groepen zo groot mogelijk (de splitsing gebeurt op basis van kleinste kwadraten ('ordinary least squares'). De methodiek doet weinig tot geen aannames ten aanzien van de distributie of variatie van de data. Binnen deze analyse methode bestaat de mogelijkheid om als responsvariabele te werken met klassen in plaats van continue waarden. In dat geval spreekt men van classificatiebomen. Met betrekking tot regressiebomen kan voor aantallen ook gewerkt worden met een Poisson verdeling. Voor de verschillende bewerkingen is gebruik gemaakt van de 'rpart' (Thernau & Atkinson, 2006) bibliotheek die beschikbaar is binnen het open-source pakket "R" (R development Core Team 2006).

Regressiebomen zijn berekend voor logaritmisch getransformeerde data en voor niet getransformeerde data waarbij een Poisson verdeling van de data werd verondersteld. Verder zijn classificatiebomen uitgerekend waarvoor alle gegevens <1 op 0 zijn gezet en alle andere waarden op 1. Op deze manier is een matrix van aan- of afwezigheidsgegevens gecreëerd. Voor een uitgebreide uitleg van regressie en classificatiebomen verwijzen we naar De'ath en Fabricius (2000).

2.3.2 Generalized Linear Models

Met de gegevens over gemiddeld voorkomen en de abiotische gegevens zijn habitatmodellen ontwikkeld met behulp van Generalized Lineair models (GLMs). De abiotische variabelen zijn hierbij opgenomen als lineaire en kwadratische termen. Hierdoor kunnen zowel lineaire als optima en minima functies worden gemodelleerd. Poisson regressie (GLM) is toegepast op de kokkelvoorkomens als ratio-scale variabele (#.m²); logistische regressie op de kokkelvoorkomens als nominal scale variabele (aan- /afwezig). Deze regressie-analyses zijn uitgevoerd met de GENMOD procedure in SAS (stepwise backwards regression, type III). De schaalparameter (i.e. overdispersie) is hierbij geschat binnen de gebruikte procedure; Indien het effect van een variabele significant ($p < 0.05$) was werd deze opgenomen in het model.

Aldus werden modellen ontwikkeld van de vorm:

$$dichtheid = e^{(a+b*DIEPTE+c*DIEPTE^2+d*ZOUT+e*ZOUT^2+f*SSH+g*SSH^2+h*SLIB+i*SLIB^2)}$$

2.3.2.1 Correctie voor overdispersie bij GLMs met Poissonverdelingen

Poissonverdelingen treden op als de te beschrijven objecten op een aselechte wijze verspreid zijn over ruimte en tijd. Als daarentegen de objecten geclusterd voorkomen (zoals bij kokkels het geval is) dan gaat een Poissonverdeling eigenlijk niet meer op (Figuur 3). In dit geval komen namelijk extreem hoge of lage uitkomsten vaker voor dan op grond van een Poissonverdeling verwacht mag worden (Oude Voshaar, 1995). De data vertonen met andere woorden een grotere variatie dan verwacht. Hierdoor is het gemiddelde lager dan de variantie. Dit verschijnsel wordt ook wel overdispersie genoemd. Volgens McCullagh & Nelder (1989) kan men er in de praktijk van uitgaan dat overdispersie bij biologische data eerder regel dan uitzondering is.

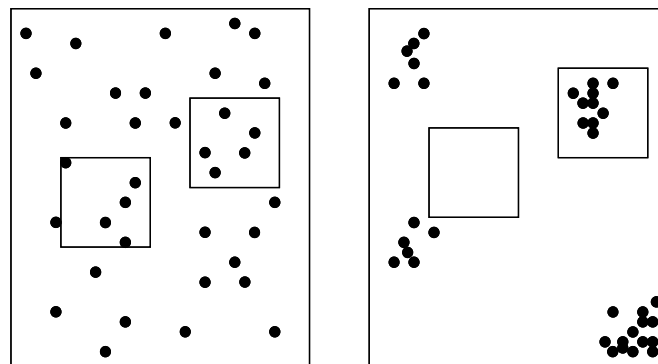
Wanneer er sprake is van overdispersie dan zijn de schattingen van de varianties en de standaardfouten te klein en worden regressieparameters te snel als 'significant' betiteld (Type I fout). Dit heeft weer tot gevolg dat ook niet relevante parameters het model kunnen beïnvloeden. Hierdoor kunnen er foute conclusies worden getrokken.

Overdispersie is een verschijnsel dat erg veel voorkomt in de biologie en hoeft dan ook niet het einde van de analyse te betekenen (Young et al, 1999). Met behulp van een zogenaamde

dispersieparameter kan namelijk worden gecorrigeerd voor dit verschijnsel. Er zijn 2 methodes om de dispersieparameter uit te rekenen:

1. Deviantie: de dispersieparameter wordt berekend door de deviantie te delen door het aantal vrijheidsgraden
2. Pearson: de dispersieparameter wordt berekend met behulp van Pearson's chi-kwadraat gedeeld door het aantal vrijheidsgraden

Deze twee methoden worden ook wel de quasi-Poisson regressie genoemd (Oude Voshaar, 1995). In een simulatiestudie werd aangetoond dat de eerste methode leidt tot grove over- en/of onderschattingen van de significantie (Young et al, 1999). Correctie door middel van de deviantie dispersieparameter (methode 1) bleek minder onderscheidend dan de correctie met behulp van de Pearson dispersieparameter (methode 2). Hierom hebben wij gekozen om de Pearson dispersieparameter te gebruiken. Tegenwoordig wordt de 2^e methode de voorkeur gegeven (o.a. door het R Core Team), hoewel de standaard instelling in SAS nog steeds de 1^e methode gebruikt. Een andere manier om met overdispersie om te gaan, is door in plaats van een Poissonverdeling een negatief binomiale verdeling te modelleren. Deze laatste methode viel buiten het bereik van deze studie.



Figuur 3 Voorbeelden waarin punten aselekt of geclusterd voorkomen. In a zijn de aantallen in aselekt gekozen deelgebieden Poisson verdeeld, in b is sprake van overdispersie (Oude Voshaar, 1995).

2.3.3 Quantielregressie

Quantielregressie is een manier om voor verschillende quantielen van de responsvariabele een lineair model te schatten waardoor een completer beeld ontstaat van de mogelijke causale relaties met andere variabelen. Bij een normale lineaire regressie kijkt men immers alleen naar de relatie ten opzichte van het gemiddelde (veelal overeenkomend met het 50% quantiel). Omdat meestal niet alle verklarende factoren opgenomen zijn in een statistisch model is het gevolg dat bij een normale regressie de relatie tussen (het gemiddelde van de distributie van) de responsvariabele (y) en de gemeten voorspellende factoren (X) zwak of afwezig is. Relaties tussen de afhankelijke variabele en de voorspellende factoren kunnen echter variëren tussen verschillende delen van de distributie van de responsvariabele (opm: maar ook zwakkere?).

Door middel van quantielregressie kunnen deze andere delen van de distributie van de y-variabele onderzocht worden. Op deze manier kunnen dan betere/veiliger modellen gecreëerd worden.

Voor de quantielregressie (Koenker 2006, Cade & Noon 2003) wordt het uiteindelijk ontwikkelde GLM model voor kokkeldichtheid als uitgangspunt gebruikt (§ 3.4.1). Aangezien de gebruikte module voor quantielregressie niet voorziet in data met een Poissonverdeling zijn de kokkeldata eerst log-getransformeerd om de verdeling van de data enigszins te benaderen met een normaal lineair model. Hierna is met deze logaritmisches getransformeerde data een nieuw model opgesteld bestaande uit dezelfde variabelen als voor de GLM gebruikt zijn. Voor de analyses door middel van quantiel regressie is gebruik gemaakt van R (R development Core Team 2006) en de module 'quantreg' (Koenker, 2006).

2.3.4 General Additive models (GAM)

Generalized Additive Models (GAM) kunnen bestaan uit verschillende componenten (Hastie & Tibshirani 1990, Wood 2000, 2001), namelijk:

- 1) Een *random component* die bestaat uit de responsvariabele Y met geassocieerde fout
- 2) Een *link functie* $g(y=a+b*diepte+c*zoutgem+d*ssh+e*sediment)$. De functies van verklarende variabelen kunnen in het algemeen waarden aannemen tussen $-\infty$ en $+\infty$. Afhankelijk van de data kan door middel van de link functie de mogelijke waarden beperkt worden (bv. in 0/1 data of alleen positieve data).
- 3) Een *systematische component* die bestaat uit een lineaire functie van een constante β_0 en afgevlakte (ge-"smooth"-de) termen van de verklarende variabelen $s(X)$. Voor meer detail zie (Hastie & Tibshirani 1990; Wood 2000; Wood 2001). Deze 'smoothing spline' kan variëren van een rechte lijn (aantal vrijheidsgraden=1) tot een interpolatie van alle datapunten ($Df=n-1$, met n het totaal aantal datapunten). Bij de selectie van het beste model dient aan de ene kant het aantal gebruikte vrijheidsgraden te worden geminimaliseerd en aan de andere kant dient het model een zo goed mogelijke fit te hebben. Deze optimalisatie wordt bereikt met behulp van een zogeheten "minimized generalized cross validation (minimized GCV)". Bij deze methode wordt een model ontwikkeld met op een na alle punten en wordt de voorspelling van dat punt vergeleken met de werkelijke waarde van het niet gebruikte punt. Deze vergelijking wordt gekwantificeerd met behulp van de GCV-score. Het beste model is het model met de laagste GCV-score (zie Wood 2000, 2001). Deze methode is rekenintensief, maar met de huidige computers geen beperking meer. Met behulp van deze 3 componenten wordt een model geconstrueerd dat de gegevens zo goed mogelijk volgt. Met andere woorden, waar bij GLMs een lineaire relatie in de parameters aan het model ten grondslag ligt, is deze relatie bij GAMs afwezig en zijn alleen de data richting gevend voor de ontwikkeling van een model.

2.4 Vergelijking van de verschillende technieken

2.4.1 *Mate van voorspelbaarheid: Hoeveelheid verklaarde variatie.*

Voor de dichtheidsvoorspellende modellen kan als maat voor het percentage verklaarde variatie een R^2 als volgt worden berekend:

$$R^2 = SS_{\text{exp}} / SS_{\text{tot}}$$

$$SS_{\text{res}} = \sum (\text{Voorspelde Dichtheid} - \text{Gevonden dichtheid})^2 = \text{residuele}$$

'sum of squares'

$$SS_{\text{exp}} = \sum (\text{Voorspelde Dichtheid} - \mu)^2 = \text{verklaarde 'sum of squares'}$$

$$SS_{\text{tot}} = SS_{\text{res}} + SS_{\text{exp}} = \text{Total 'sum of squares'}$$

Met μ = gemiddelde.

Er wordt een R^2 gegeven voor de modeldata (in het zuiden); dit is een maat voor het beschrijvend vermogen van het model (calibratie). Met het model dat op basis van de data uit het zuiden gemaakt is, wordt vervolgens een voorspelling gedaan voor het noorden aan de hand van de abiotische gegevens van het noorden (validatie). Hiermee kan dan ook via een vergelijking tussen de voorspelde en gevonden waarden van de responsvariabele een R^2 berekend worden. De R^2 voor de data in het noorden zegt dan iets over de correct voorspelde variatie in de gegevens door het model (validatie).

Met behulp van de modeluitkomsten is ook de totale hoeveelheid aan (voorspelde) kokkels in het noorden geschat. Hierbij is gebruik gemaakt van de stratificering zoals gebruikt in de jaarlijkse opnames. Stratificering houdt in dat intensiever wordt bemonsterd in die gebieden waar grotere dichtheden kokkels worden verwacht. Elk bemonsterd station wordt vervolgens representatief verondersteld voor een oppervlak dat varieert per stratum (Kesteloo et al, 2006). De som van voorspelde dichtheden ($\#/m^2$) per station vermenigvuldigd met het representatief geacht oppervlak dat bij het betreffende station hoort is een schatting van de totale (gemiddelde) dichtheid van de kokkels (let wel: Voor deze schatting zijn enkel de stations uit de analyse gebruikt).

Daarnaast is voor ieder model het gemiddelde absolute verschil tussen de gemeten waarden en de voorspelde dichtheden berekend.

Voor de aan- afwezigheidsmodellen is berekend in hoeveel % van de gevallen het kokkelvoorkomen door het model juist wordt voorspeld. Aangezien de aan-afwezigheidsmodellen als uitkomst een kans op voorkomen geven is de grens gelegd bij 0.50. Alle stations waar volgens de modellen de kans op voorkomen kleiner is dan 0.50 kregen de waarde 0 (geen kokkels) en de overige stations kregen waarde 1 (wel kokkels).

2.4.2 *Ruimtelijke patronen: GIS-plaatjes*

R^2 is een veel gebruikte manier om de geschiktheid van een model te kunnen inschatten. Voor het beleid is het echter ook van belang dat de juiste patronen in de verspreiding van soorten

worden voorspeld. Met behulp van Arc-GIS (vs. 9.1) zijn verspreidingskaartjes gemaakt van gemiddelde (waargenomen) kokkeldichtheden en modelvoorspellingen in het noordelijke deel van de Westerschelde.

Hiertoe zijn de kokkeldata eerst opgedeeld in 3 klassen: Dichtheid (DN)=0 #/m² (gold voor 57% van de stations in 2002-2006), DN=0-50 #/m² (30%), DN > 50 #/m² (13%)¹. Per model is ook berekend hoe vaak de modelvoorspellingen in de juiste klasse voorkwamen.

2.4.3 Multi dimensional Scaling (MDS) om modeluitkomsten te vergelijken

Non-metric Multidimensional Scaling (MDS) is een methode om de gelijkenis tussen verschillende kolommen of rijen van een matrix door middel van (dis)similariteiten (bv. tussen stations/monsters) te visualiseren in 2 of 3 dimensies (Shepard 1962, Kruskal 1964). Aangezien de matrix van voorspelde dichtheden op elk station (rijen) tegen methodieken (kolommen) ook gezien kan worden als een soorten bij stations matrix, kunnen we procedures, ontwikkeld voor het vergelijken van stations, toepassen op de gebruikte modelleringsmethodieken. De methoden worden op basis van de uitkomst per station met elkaar vergeleken en er wordt een (dis)similariteit/afstand berekend. In het twee dimensionale vlak kunnen de methodieken dan dusdanig gepositioneerd worden dat hun afstand tot elkaar (overeenkomstig hun gelijkenis met de andere dichtheidsvoorspellingen) optimaal wordt weergegeven. Het kiezen van een geschikte afstandsindex is hierbij wel van cruciaal belang. Bij de vergelijking zijn waarden kleiner dan 0.5 op 0 gezet en alle getallen tot de vierdemachts wortel getransformeerd grote getallen de afstandsindex niet te laten domineren. Daarna is als afstandsindex de euclidische afstand tussen elk paar van technieken uitgerekend.

2.4.4 Sensitiviteit & specificiteit

Om te kunnen inschatten of een model de aanwezigheid cq afwezigheid van kokkels goed voorspelt kan een sensitiviteit en een specificiteit worden berekend (Pearce & Ferrier, 2000). De specificiteit van een model is een maat voor de frequentie waarmee een negatief resultaat wordt gevonden (dat wil zeggen, de testuitkomst is dat datgene waarop getest wordt afwezig is) indien de conditie ook werkelijk afwezig is.

Op analoge wijze geeft de sensitiviteit van een test de frequentie aan waarmee voor dat model een positief resultaat wordt gevonden als de conditie ook werkelijk aanwezig is. De sensitiviteit en de specificiteit worden beide uitgedrukt als fractie, of in procenten, bijvoorbeeld 0,90 of 90%.

Op grond van bovenstaande definities wordt in het ideale geval voor zowel de specificiteit als de sensitiviteit van een test 100% gevonden. In werkelijkheid komt dit niet voor. Meestal daalt het ene als het andere stijgt: een test heeft altijd een twijfelgebied, en een hoge specificiteit wordt bereikt door in dit twijfelgebied negatief te kiezen, terwijl een hoge sensitiviteit juist wordt

¹ Wanneer er meer dan 50 kokkels per m² voorkomen wordt gesproken van een kokkelbank.

bereikt door dit twijfelgebied positief te kiezen. Men kiest afhankelijk van de situatie voor een zo hoog mogelijke specificiteit (afwezigheid is belangrijk) of een zo hoog mogelijke sensitiviteit (aanwezigheid is belangrijk).

Hiertoe is eerst een matrix gemaakt met de mogelijke combinaties van specificiteit en sensitiviteit:

	Aanwezig	Afwezig
Voorspeld aanwezig	A	B
Voorspeld afwezig	C	D
	A+C	B+D

$$\begin{aligned}
 \text{Sensitiviteit} &= \text{aantal locaties waar kokkels aanwezig zijn correct voorspeld} / \text{totaal locaties met kokkels aanwezig} \\
 &= A / (A+C)
 \end{aligned}$$

$$\begin{aligned}
 \text{Specificiteit} &= \text{aantal locaties waar kokkels afwezig zijn correct voorspelt} / \text{totaal aantal locaties kokkels afwezig} \\
 &= D / (B+D)
 \end{aligned}$$

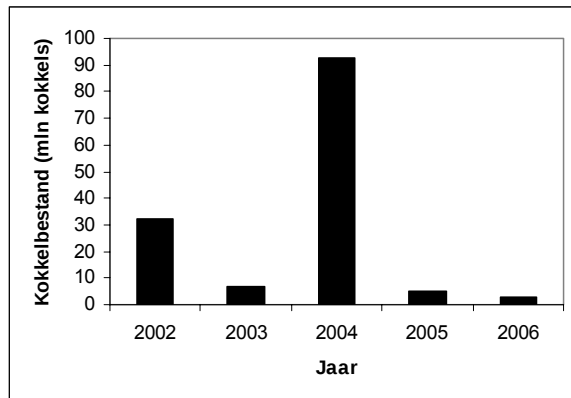
Deze 2 waarden zijn voor ieder model berekend. Er is dus enkel gekeken of het model ook daadwerkelijk kokkels voorspelde in gebieden waar kokkels aanwezig waren en of het model op locaties waar geen kokkels aanwezig waren een ook tot een 0-waarde kwam. Wanneer zowel de sensitiviteit als de specificiteit de waarde 1 nadert kun je spreken van een goed model.

3 Resultaten

3.1 Verkenning van de data

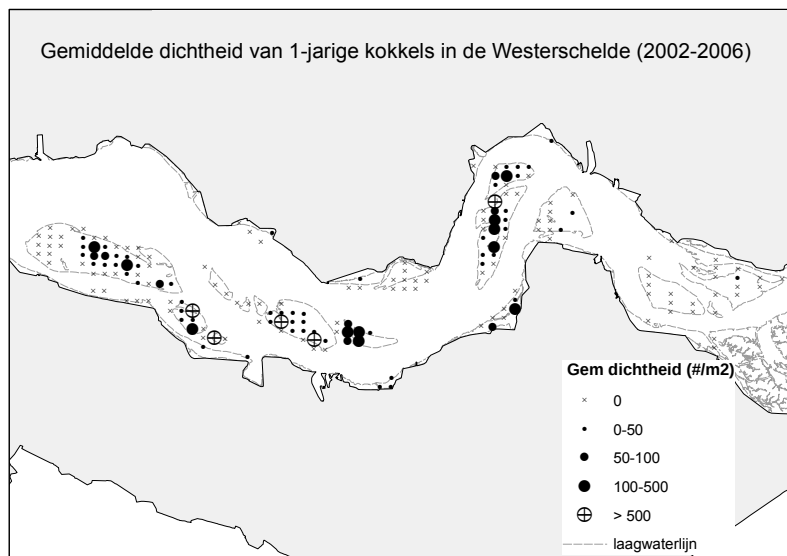
3.1.1 Kokkels

Er is sprake van een grote variatie in het voorkomen van 1-jarige kokkels tussen de jaren. In het jaar 2004 kwamen verreweg de meeste 1-jarige kokkels voor in de Westerschelde (Figuur 4).



Figuur 4 dichtheid 1-jarige kokkels in de Westerschelde periode 2002-2006.

Zoals blijkt uit Figuur 5 is de verspreiding van kokkels in de Westerschelde niet erg homogeen. Op de meeste stations die zijn gebruikt voor de analyse zijn zelfs helemaal geen kokkels aangetroffen. Slechts op 5 stations was de gemiddelde kokkeldichtheid groter dan 500 kokkels per m².

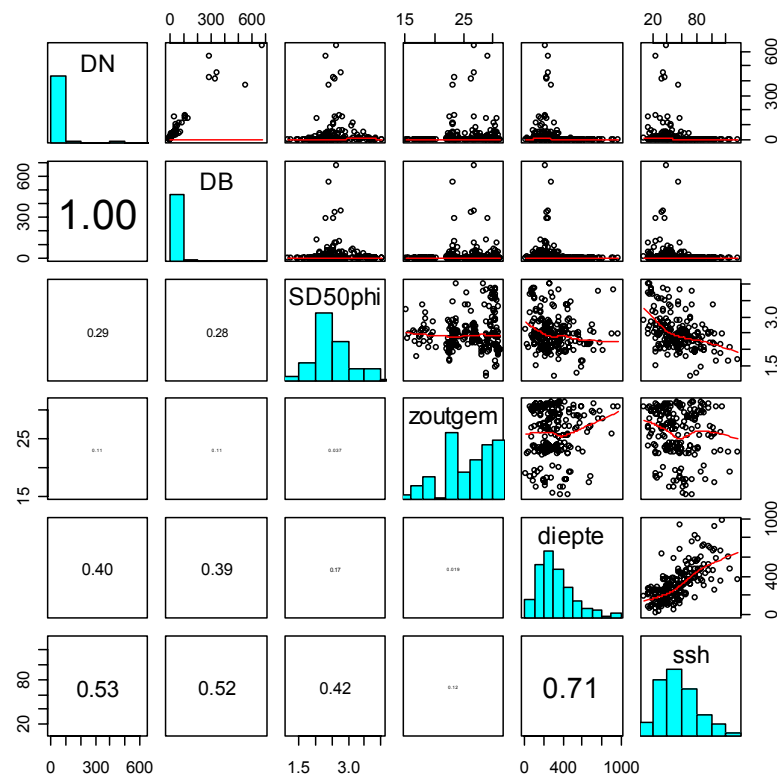


Figuur 5 Gemiddelde dichtheid van 1-jarige kokkels in de Westerschelde van 2002-2006

3.2 Verkenning van de data

Wanneer de dichtheid en de biomassa van de 1-jarige kokkels tegen elkaar worden uitgezet, blijkt dat er sprake is van een hoge mate van correlatie (Figuur 6, bovenste rij, 2^e van links en 2^e rij, 1^e van links). Voor de constructie van een model maakt het dientengevolge niet veel uit of men met dichtheid of met biomassa werkt. Verder blijkt dat kokkels zeer onregelmatig voorkomen en dat er op de meeste stations zelfs helemaal geen kokkels aanwezig waren.

Bij een één op één vergelijking (Figuur 6) is er geen sprake van een duidelijke correlatie tussen de omgevingsfactoren en de dichtheid of biomassa van kokkels (bovenste twee rijen in Figuur 6). Wel lijkt er sprake te zijn van een optimum met betrekking tot de mediane korrelgrootte (SD50phi) en komen er geen kokkels meer voor bij een zoutgehalte lager dan 18. Ook kan (zij het met enige moeite) men een optimum zien bij diepte en stroomsnelheid (ssh). Tussen diepte en stroomsnelheid lijkt tevens een sterk positief verband te bestaan (Spearman Rank correlatie is 0.71). Ook blijkt de verdeling van de verschillende variabelen vaak erg scheef. Dit laatste zou kunnen betekenen dat bij bepaalde analyses de variabelen getransformeerd moeten worden om 1) aan de assumpties van de analyse te voldoen of 2) een beter model te krijgen.



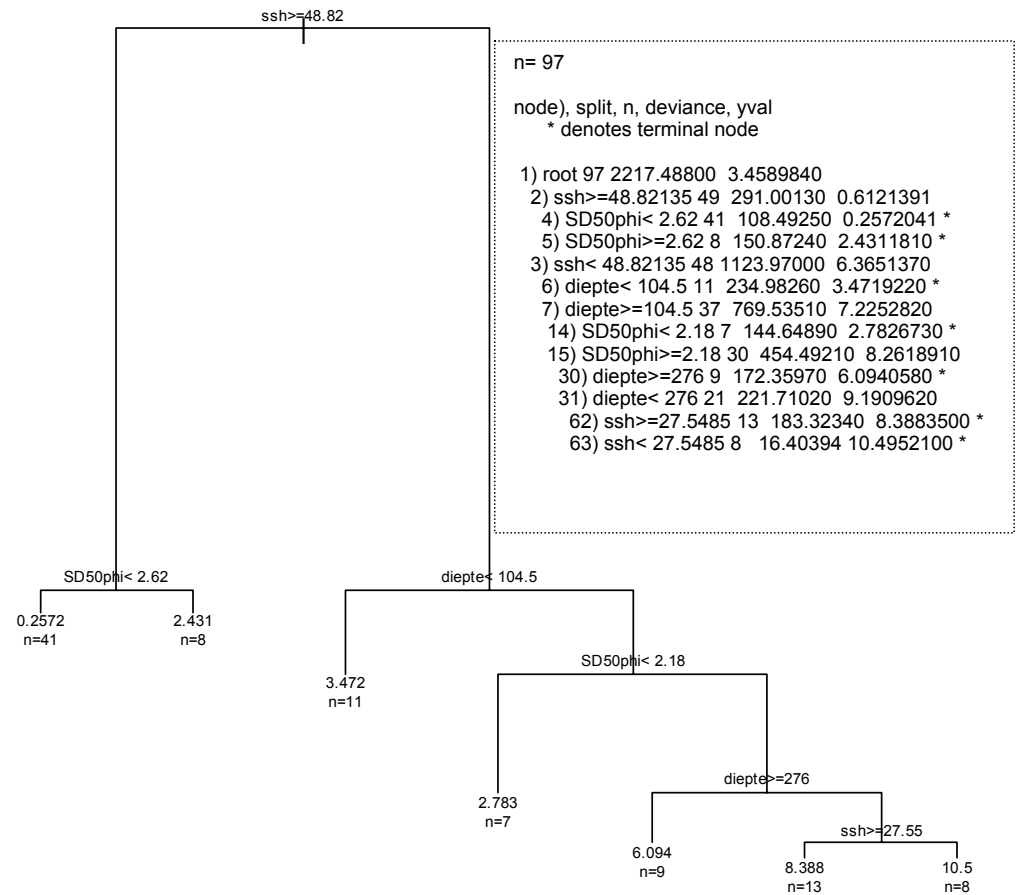
Figuur 6 'Pairplot' van de verschillende variabelen (gemiddelde dichtheid en biomassa (DN & DB) 1-jarige kokkels en de omgevingsfactoren; SD50phi=mediane korrelgrootte, zoutgem=gemiddeld zoutgehalte, diepte en ssh=maximale stroomsnelheid bij springtij). Rechts van de diagonaal: diagrammen waar de variabelen tegen elkaar worden uitgezet (inclusief een 'smoothing' lijn). Diagonaal: histogrammen van de variabelen en links van de diagonaal de correlatiecoëfficiënten (Spearman rank); tekstgrootte van de correlatiecoëfficiënt is proportioneel ten opzichte van de correlatie.

3.3 Regressie- en classificatieboom analyses

In Figuur 7, Figuur 8 en Figuur 9 wordt de regressieboom weergegeven die uiteindelijk de beste schatting gaf (gebaseerd op R^2). Boven elke splitsing staat een grenswaarde van een van de abiotische variabelen. Wanneer aan die conditie wordt voldaan dan moet men de vertakking naar links volgen. Aan het einde van een tak staat de verwachte waarde van de dichtheid cq aan- of afwezig en het aantal gevallen dat aan die conditie voldoet.

3.3.1 *Log getransformeerde data (DN)*

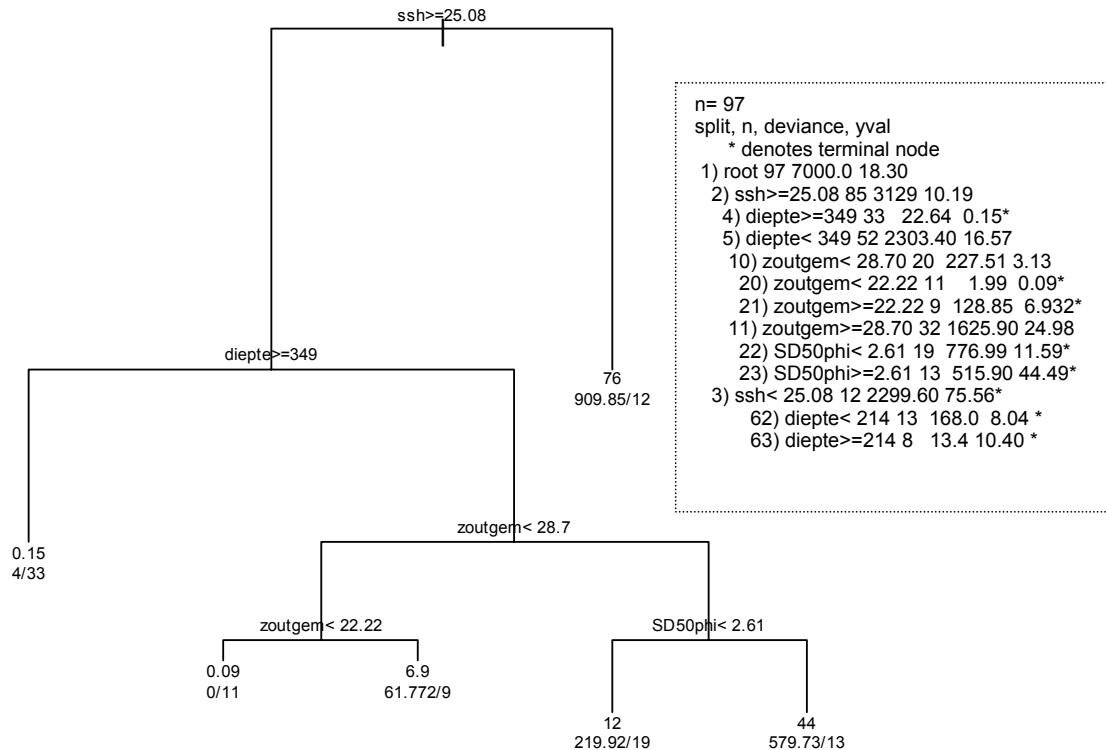
Het model verklaart 54 % van de variatie (R^2) en in het noorden verklaart het model 56 % van de variatie. Echter dit is gebaseerd op log-getransformeerde data (dit is wel de gebruikelijke manier om de R^2 uit te rekenen). Voor vergelijking met de andere technieken transformeren we de waarden terug naar hun oorspronkelijke schaal en worden de verklaarde variaties respectievelijk 8 en 9%. Dit verschil wordt veroorzaakt door de kwadratische termen in de berekening van de R^2 . Grote getallen kunnen een groot effect hebben op de berekende waarde (indien een grote waarde relatief goed voorspeld wordt, zal de R^2 altijd hoog zijn onafhankelijk van rest van de data). De belangrijkste variabelen zijn stroomsnelheid (SSH), mediane korrelgrootte (SD50phi) en diepte (Figuur 7). De lengte van een splitsing is relatief ten opzichte van de verklaarde variatie. Blijkbaar kan met stroomsnelheid al een groot deel verklaard worden.



Figuur 7 Regressie boom voor de dichtheden van kokkels (Data $\log((DN*1000)+1)$ getransformeerd)

3.3.2 Poisson Regressieboom

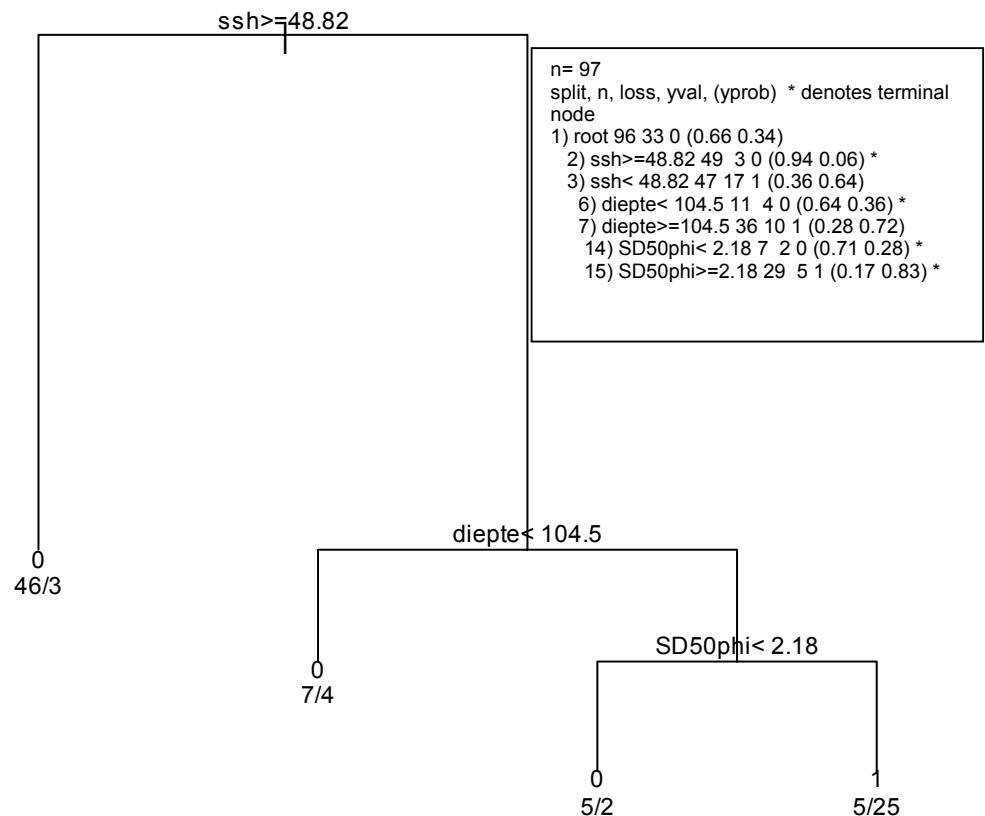
De Poisson regressieboom verklaart 16 % van de variatie van de data in het zuiden (modelvoorspelling) en 9 % van de variatie van de data in het noorden. Stroomsnelheid is wederom de belangrijkste bepalende factor (Figuur 8), maar andere belangrijke factoren zijn diepte en het gemiddelde zoutgehalte (zoutgem). De uitgebreide uitvoer van de analyse (niet opgenomen) geeft verder aan dat als alternatief voor stroomsnelheid bij de eerste splitsing, ook diepte gebruikt kan worden ($diepte < 255$), maar deze keuze resulteert wel in een model dat iets minder van de totale variatie verklaart). Met behulp van mediane korrelgrootte wordt op het einde nog een splitsing aangebracht in de hogere waarden die in twee groepen worden opgesplitst met geschatte dichtheden van 220 en 580 kokkels per vierkante meter.



Figuur 8 Poisson regressieboom van de dichtheden van kokkels. De lengte van de splitsing is relatief ten opzichte van de hoeveelheid verklaarde variatie. Als aan de conditie voldaan wordt, dient men de linker vertakking te nemen. Aan het eind van elke tak staan de geschatte waarde, het totale aantal en het aantal waarnemingen (monsters) in die groep.

3.3.3 Classificatieboom

Het model zoals weergegeven in (Figuur 9) voorspelt in 80% van de gevallen juist. In het noorden werd in 73 % van de gevallen de aan- of afwezigheid van kokkels juist voorspelt. Opnieuw blijkt stroomsnelheid de belangrijkste variabele te zijn, met name voor het voorspellen van afwezigheid, terwijl de mediane korrelgrootte vooral van belang is om de aanwezigheid van kokkels te voorspellen.



Figuur 9 regression tree van kokkels; 0=afwezig, 1=aanwezig. Aan het einde van elke vertakking staan de verwachte uitkomst (0, voor afwezig; 1 voor aanwezig), het aantal in de eerste klasse en het aantal in de tweede klasse.

3.4 Generalized linear models (GLM)

Parameterschattingen worden weergegeven in: Bijlage 1

3.4.1 Kokkeldichtheid (Poisson regressie)

Dit model verklaart 21% van de totale variantie (R^2). Wanneer het model wordt toegepast in het noorden, dan wordt 38 % van de variatie verklaard (R^2). Het model doet het dus blijkbaar beter in het noorden dan in het zuiden, waar het is ontwikkeld.

De belangrijkste significante factoren zijn diepte en gemiddeld zoutgehalte. Mediane korrelgrootte speelt ook een rol, maar is iets minder significant.

3.4.2 Aanwezig- afwezigheid

Het model geeft een juiste voorspelling in 81 % gevallen. In het noorden worden is de voorspelling in 67% van de gevallen juist.

Significante factoren voor dit model zijn zoutgehalte en mediane korrelgrootte, de lineaire factor van zoutgehalte is hierbij het minst significant.

3.5 Quantiel regressie

Parameterschattingen worden weergegeven in: Bijlage 1

Aangezien de methode niet kan werken met een Poissonverdeling, zijn de dichtheden eerst log-getransformeerd en daarna hetzelfde model toegepast als voor de GLM is gebruikt. Uit de resultaten blijkt dat niet alle variabelen significant zijn. Dit is het gevolg van het feit dat de quantielen vaak maar een klein deel van de data bevatten en de regressie dus veel minder kracht heeft om significante effecten te onderscheiden. Voor de voorspelling maakt dat echter niet uit, hoewel de bijdrage van een niet-significante variabele in het algemeen klein zal zijn. De R^2 voor de modellen met data uit het zuiden (calibratie) zijn 6, 5, 4, 6 en 21%. De quantielen onder de 0.5 zijn niet meegenomen omdat deze alleen maar nullen bevatten (de meeste data bestaan uit nullen). De modellen geven steeds meer significante factoren bij hogere quantielen. Dit komt omdat de hogere quantielen de hogere dichtheden bevatten die meer en meer op de veronderstelde respons curve passen. In principe kan men stellen dat de natuurlijke variatie voor 80% zal bestaan uit de grenzen die door de modellen voor het 0.1 en 0.9 quantiel worden afgebakend. De R^2 voor de voorspellingen van de verschillende modellen zijn 9, 9, 8, 6 en 8% voor respectievelijk het 0.5, 0.6, 0.7, 0.8 en 0.9 quantiel.

3.6 General Additive models (GAM)

De resultaten van de GAM's worden weergegeven in Bijlage 2.

Het model verklaarde 75 % van de variatie van de modeldata. In het noorden deed het model het een stuk minder goed, slechts 9 % van de data werd verklaard. Dit kan verklaard worden door het feit dat het GAM model vooral de data volgt en niet alle variabelen in het model aanwezig zijn die de variatie in een afdoende manier beschrijven. Indien dit wel zo was, zou het verschil tussen calibratie en validatie niet zo groot zijn. Voor de calibratie is het GAM model uitermate geschikt, maar bij de validatie blijkt dat het model te specifiek is voor de situatie in het zuiden en slecht gebruikt kan worden voor het noorden.

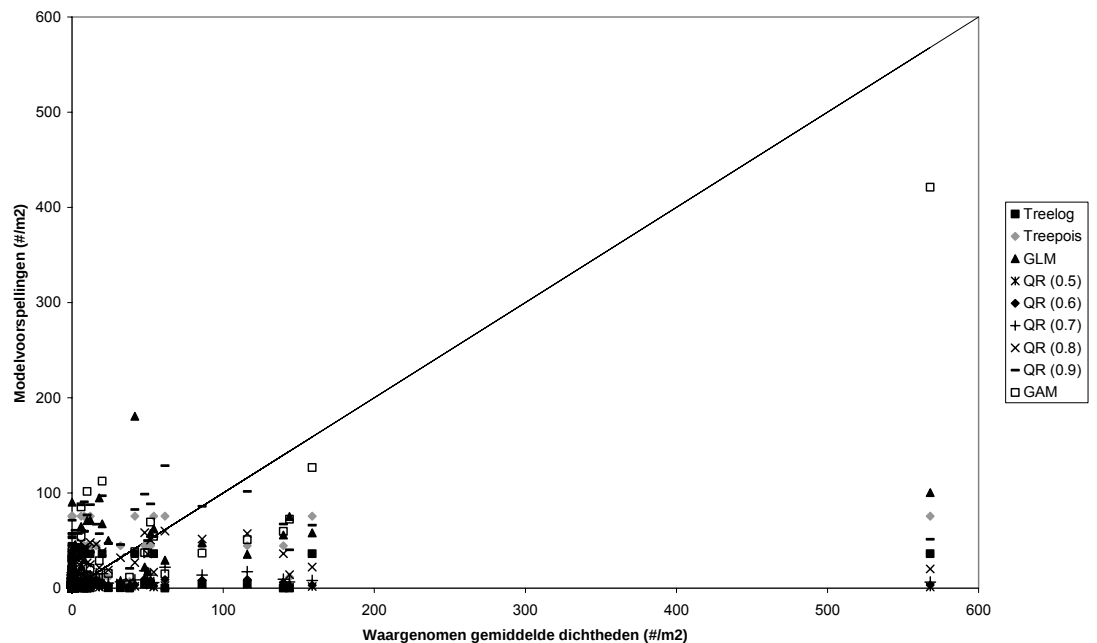
4 Vergelijking van de verschillende technieken

4.1 Beschrijvend vermogen

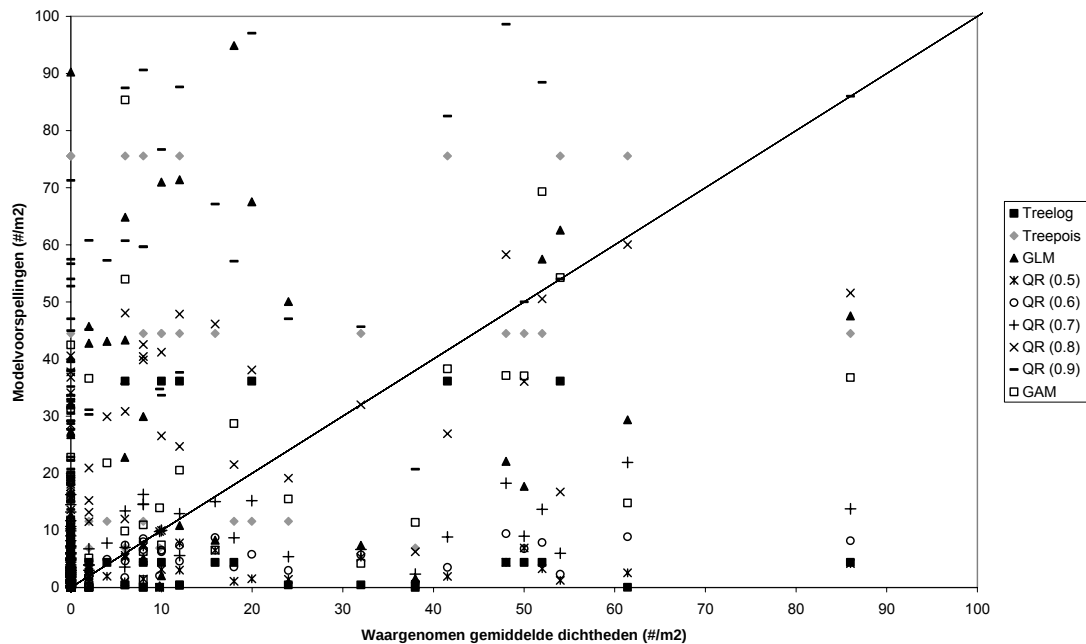
Op basis van de R^2 zouden we moeten concluderen dat het GAM model met een R^2 van 75 % verreweg het beste is (Tabel 2). Zoals aangegeven in § 3.6 komt dit doordat het GAM model vooral de data volgt. Het gebruik van R^2 is echter niet optimaal. Doordat er gekwadeerd wordt, kunnen grote verschillen of getallen een groot effect uitoefenen op de uiteindelijke hoeveelheid verklaarde variatie. Er moet dan niet teveel gewicht worden toegekend aan deze index en men moet het vooral relatief (modellen tov elkaar) beschouwen. Een ander probleem zijn de p-waarden bij de analyses. Door het gebruik van gekwadeerde variabelen in de lineaire regressies ontstaat er een hoge mate van correlatie tussen de factoren en hun gekwadeerde variant. Deze correlatie beïnvloedt de p-waarden.

De modelvoorspellingen en waargenomen dichtheden in het zuiden van de Westerschelde worden weergegeven in Figuur 10a en b. De hogere waarden (Figuur 10a) lijken vooral door de GAM redelijk te worden voorspeld. Geen van de toegepaste methodieken geeft echter een overduidelijk goede fit. Het is dan ook moeilijk om op basis van de grafieken een conclusie te trekken over de geschiktheid van de modellen. Om een objectieve beoordeling van de methodieken mogelijk te maken, hebben we ook nog op andere manieren naar de calibraties gekeken (Tabel 1).

a



b



Figuur 10 calibratie van de modellen; y=voorspellingen, x=gemeten waardes. a) totale range, b) waarden < 100. Bij een goede voorspelling liggen alle punten op de diagonale lijn.

Totale dichtheid. De totale gemiddelde dichtheid aan kokkels in het zuiden van de Westerschelde is geschat op 473 miljoen (Tabel 1). De voorspelling van de totale dichtheid op basis van de uitkomsten van de Poisson regressies (regressieboom en GLM) en de GAM komen hiermee overeen. Dit komt omdat bij deze modellen van te voren geen transformatie is toegepast in tegenstelling tot de regressieboom die voor de log-getransformeerde data is gemaakt. Ook bij de quantielregressies komen de totale modeldichtheden niet overeen met de waargenomen totale dichtheid. Dit komt doordat deze regressies slechts een deel van de data gebruiken. De dichtheid berekend met de quantielregressie met $\tau=0.8$ komt het meest overeen met de waargenomen dichtheid.

Gemiddeld absoluut verschil. Het gemiddelde absolute verschil tussen de gemeten data en de modeldata varieert van 12.75 kokkels/m² (GAM) tot 33.71 kokkels/m² (quantielregressie; $\tau=0.9$). Als we de verschillen eerst transformeren om het effect van de grote getallen te beperken presteren de quantielregressies ($\tau=0.5$ & $\tau=0.7$) met een verschil van resp 0.42 en 0.46 het best en de Poisson regressieboom en de quantielregressie ($\tau=0.8$) met een verschil van resp 1.47 en 1.55 het slechtst.

Klassen (§ 2.4.2). Op basis van de 3 dichtheidsklassen; (Dichtheid=0, Dichtheid=0-50 en Dichtheid> 50 kokkels/m²) kan worden aangetoond voor welke waargenomen klasse de modellen een goede voorspelling geven en voor welke klasse de voorspelling minder goed overeenkomt. Zo valt op dat bijna geen model voorspelt dat er meer 50 kokkels per m² op een

locatie voorkomen. De 0-waarnemingen worden met name door de hogere quantielregressies (0.8 & 0.9) niet goed voorspeld. De GAM geeft voor elke klasse een goede voorspelling (> 70% juist voorspeld) en de zowel de Poisson regressieboom als GLM geven voor elke klasse een goede tot matig goede voorspelling.

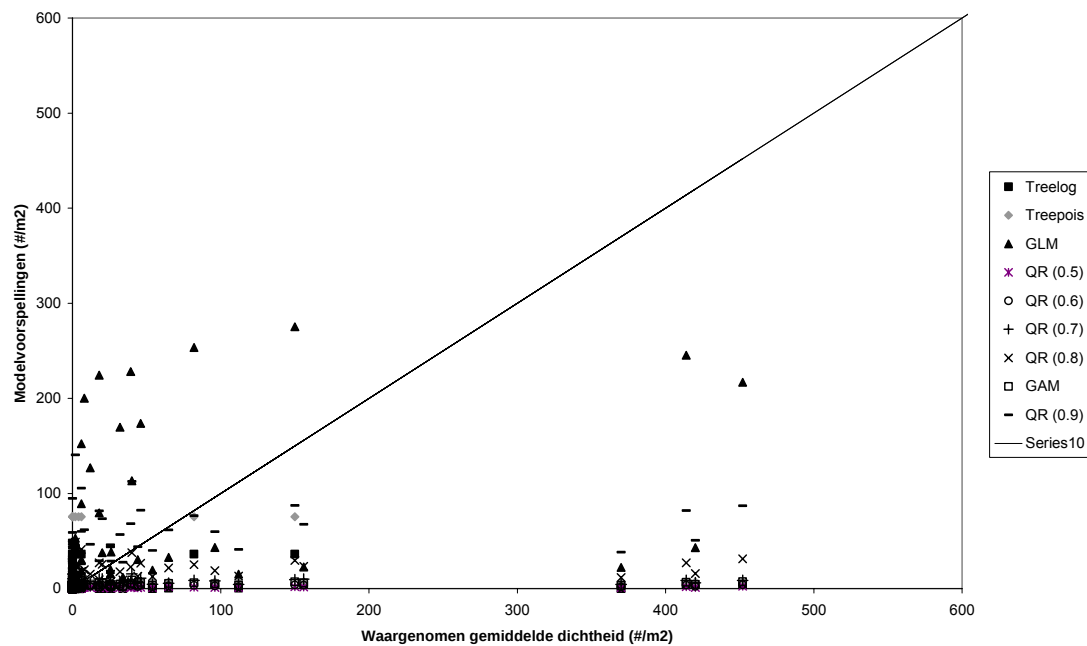
Tabel 1 Modeluitkomsten (zuiden), DN = gemiddelde dichtheid (#/m²) in 2002-2006. % correct: % dat de gemeten dichtheidsklasse dezelfde was als de voorspelde. Gewogen gem. % correct van 69%: in 69% van de gevallen viel de voorspelling in dezelfde klasse als de gemeten waardes.

	R ²	Geschatte totale dichtheid n (#*mln)	Gemiddeld absoluut verschil	Gemiddeld absoluut verschil (4e wortel getransformeerd)	klasse Dichtheid=0, % correct	klasse Dichtheid = 0-50, % correct	klasse Dichtheid > 50, % correct	klassen: gewogen gem. % correct
DN		473						
treelog	0.08	94	16.95	0.60	97%	52%	0%	76%
treepois	0.16	474	20.55	1.47	67%	76%	44%	67%
GLM	0.22	474	21.15	0.96	54%	64%	67%	58%
GAM	0.75	474	12.73	0.96	70%	72%	78%	71%
QR (0.5)	0.06	48	17.97	0.42	76%	88%	0%	72%
QR (0.6)	0.05	83	17.84	1.02	37%	96%	0%	48%
QR (0.7)	0.04	167	18.28	0.46	27%	100%	0%	43%
QR (0.8)	0.06	464	22.03	1.55	19%	96%	44%	41%
QR (0.9)	0.21	946	33.71	0.70	13%	44%	89%	28%

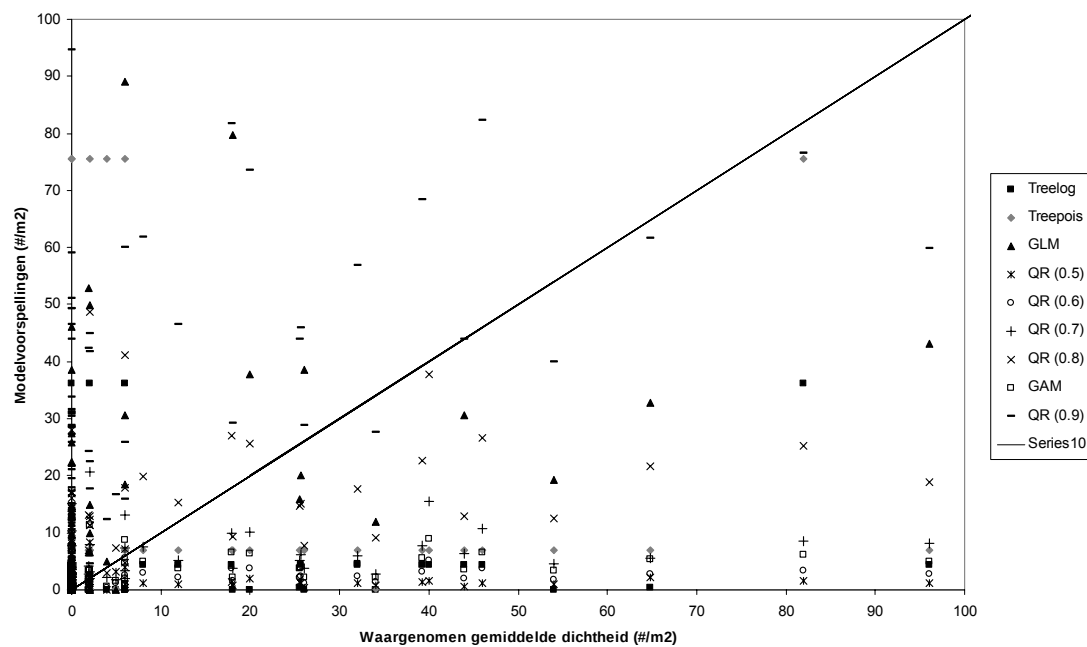
4.2 Voorspellend vermogen

Op basis van de R² zouden we moeten concluderen dat het voorspellende vermogen van de Poisson regressie (GLM) met 38 % verreweg het beste is (Tabel 2). Echter, zoals gezegd, kunnen incidentele hoge waarden een groot effect uitoefenen op de berekende R². Ook in dit geval was dat het geval. Anders even uitleggen dat het is gebeurt ondanks dat we alles > 1000 eruit hebben gehaald) waarde die zeer hoog was t.o.v. de andere dichtheden en als deze waarde buiten het model werd gehouden veranderde de hoeveelheid verklaarde variatie van het calibratie model amper, maar die van het validatie model daalde naar 13%. Blijkbaar kan het behouden van 1 datapunt een dusdanig groot effect op de modelparameters hebben dat de hoeveelheid verklaarde variatie zeer sterk beïnvloed wordt. Dit kan ook als volgt verklaard worden. Indien de hoge waarde in het model wordt gehouden, zal het uiteindelijke model hogere waarden voorspellen. Aangezien het juist die hoge waarden zijn die een (disproportioneel) groot effect op de totale hoeveelheid verklaarde variatie hebben, zal de R² een stuk hoger zijn. Toch is het dan heel goed mogelijk dat de voorspelde waarden over de gehele linie minder goed zijn. Het is duidelijk dat betere criteria nodig zijn om de verschillende technieken met elkaar te vergelijken.

a)



b)



Figuur 11 fit van de modellen in het noorden; y=voorspellingen, x=gemeten waardes. a) totale range, b) waarden < 100. Bij een goede voorspelling liggen alle punten op de diagonale lijn.

De modelvoorspellingen en waargenomen dichtheden in het zuiden van de Westerschelde worden weergegeven in Figuur 11 a en b. De lagere waarden (Figuur 11b) lijken vooral door de 0.8 quantiel goed te worden voorspeld, maar het is duidelijk dat geen van de toegepaste methodieken een echt goede fit geeft. Om een objectieve beoordeling van de methodieken

mogelijk te maken, hebben we ook nog op andere manieren naar de voorspellingen gekeken (Tabel 2).

Totale dichtheid. De totale gemiddelde dichtheid aan kokkels in het noorden van de Westerschelde is geschat op 919 miljoen (Tabel 2). De voorspelling van de totale dichtheid op basis van de uitkomsten van de GLM komt het dichtst in de buurt gevolgd door de quantielregressie ($\tau=0.9$). De overige modellen geven allemaal een onderschatting van de totale dichtheden.

Gemiddeld absoluut verschil. Wanneer we kijken naar het absolute gemiddelde verschil tussen de voorspellingen en de gemeten waarden geven de GLM en de quantielregressie ($\tau=0.8$) de grootste verschillen. Als we de verschillen eerst transformeren om het effect van de grote getallen te beperken lijken de GAM en de regression tree (logdata) het best bruikbaar.

Klassen (§ 2.4.2). Op basis van de 3 dichtheidsklassen; (Dichtheid=0, Dichtheid=0-50 en Dichtheid> 50 kokkels/m²) kan worden aangetoond voor welke waargenomen klasse de modellen een goede voorspelling geven en voor welke klasse de voorspelling minder goed overeenkomt. Zo valt op dat bijna geen model voorspelt dat er meer 50 kokkels per m² op een locatie voorkomen. De 0-waarnemingen worden met name door de hogere quantielregressies (0.8 & 0.9) niet goed voorspelt. De GLM geeft voor elke klasse een goede tot matig goede voorspelling. Als we de modellen over het geheel genomen beoordelen blijkt dat de regressiebomen de beste voorspellingen geven (resp. 69 & 67 % van de voorspellingen vielen in de juiste klasse). Hier moet wel bij opgemerkt worden dat wederom de hoogste dichtheidsklassen door deze modellen niet goed worden voorspeld. Aangezien dit echter maar om een klein deel van de data gaat (op slechts 12 van de 94 locaties waren meer dan 50 kokkels/m² aanwezig) heeft dit niet zo'n grote invloed.

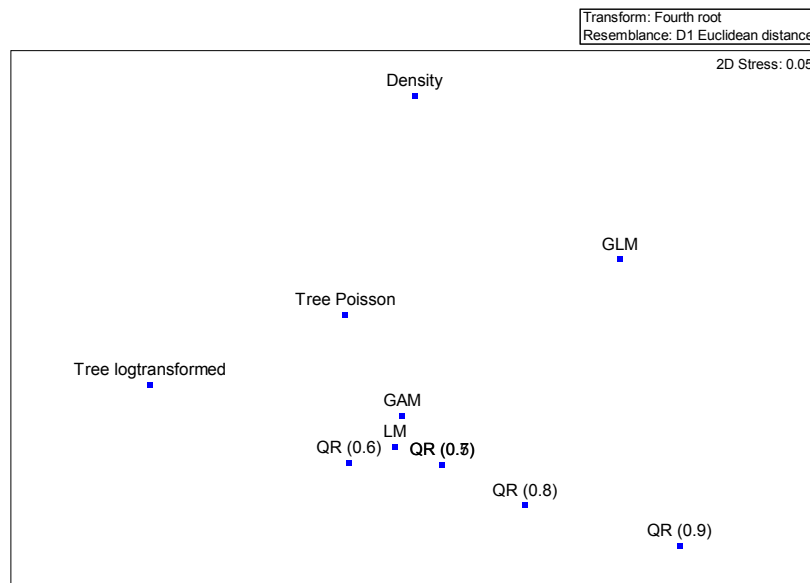
In Bijlage 3 zijn de modelvoorspellingen grafisch weergegeven.

Tabel 2 Modeluitkomsten (noorden), DN = gemiddelde dichtheid (#/m²) in 2002-2006. % correct: % dat de gemeten dichtheidsklasse dezelfde was als de voorspelde. Gewogen gem. % correct van 69%: in 69% van de gevallen viel de voorspelling in dezelfde klasse als de gemeten waardes.

	R ²	Geschatte totale dichtheid n (#*mln)	Gemiddeld absoluut verschil	Gemiddeld absoluut verschil (4e wortel getransformeerd)	klasse Dichtheid=0, % correct	klasse Dichtheid = 0-50, % correct	klasse Dichtheid > 50, % correct	klassen: gewogen gem. % correct
DN		919						
treelog	0.09	69	36.19	1.19	96%	37%	0%	69%
treepois	0.09	239	38.65	1.53	78%	70%	17%	67%
glm	0.38	970	42.45	1.59	65%	47%	42%	57%
gam	0.09	68	36.62	1.01	48%	83%	0%	54%
QR (0.5)	0.09	29	36.25	1.50	78%	57%	0%	62%
QR (0.6)	0.09	44	36.08	1.62	52%	80%	0%	55%
QR (0.7)	0.08	111	36.54	1.81	31%	93%	0%	47%
QR (0.8)	0.06	292	44.71	2.13	19%	93%	0%	40%
QR (0.9)	0.08	864	36.03	1.55	9%	60%	75%	24%

4.2.1 MDS

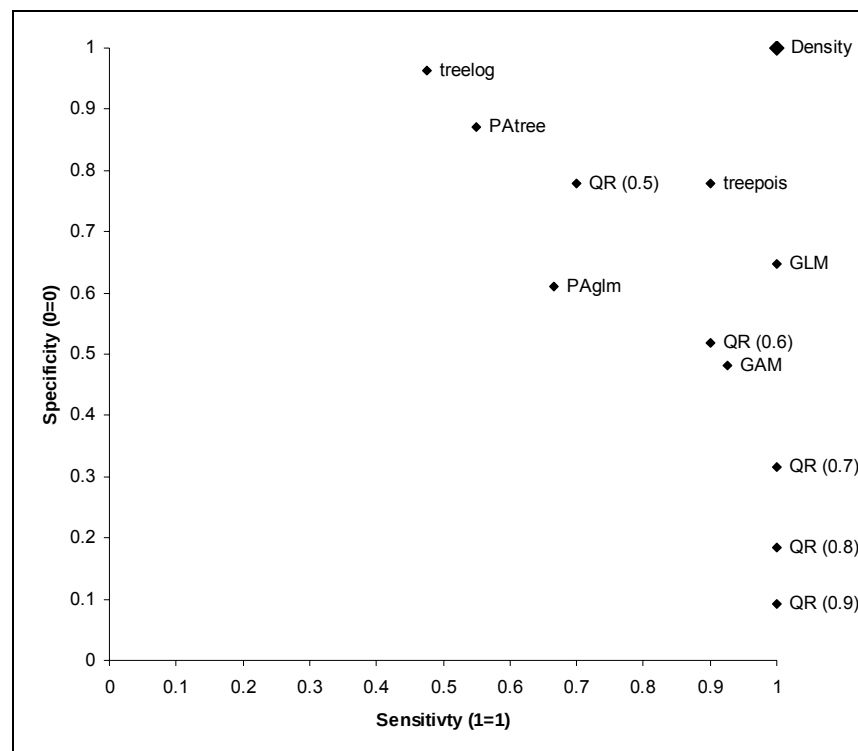
Het MDS plot toont de gelijkenis van de voorspellingen van de verschillende technieken met de echte waarden en met elkaar (Figuur 12. MDS plot). Via een dissimilariteitsindex (Euclidische afstand) worden per techniek de voorspellingen voor alle stations vergeleken met de gevonden waarden. Een techniek waarvoor de voorspellingen het dichtst bij de werkelijke waarden zitten, zal dan de kleinste afstand tot het punt 'Density' in de grafiek hebben. De stress-waarde (rechts boven) geeft aan dat de onderlinge similariteit voldoende kan worden weergegeven in 2 dimensies. Een vergelijking op deze manier laat zien dat de Poisson regressieboom en de GLM het dichtst bij de werkelijke waarden komen.



Figuur 12. MDS plot (het punt Density geeft de positie aan van de originele data uit het noorden)

4.2.2 Sensitiviteit & Specificiteit

Wanneer we de modellen enkel beoordelen op basis van aanwezigheid en afwezigheid van kokkels komen de Poisson regressie (GLM) en Poisson regressieboom het dichtst in de buurt van de waarneming (Figuur 13). De GLM voorspelt vooral goed de aanwezigheid, terwijl de Poisson regressieboom de afwezigheid beter voorspelt, maar ook de aanwezigheid nog tamelijk goed voorspelt. De andere twee regressiebomen zijn goed in het voorspellen van de afwezigheid, maar minder in het voorspellen van de aanwezigheid van kokkels.

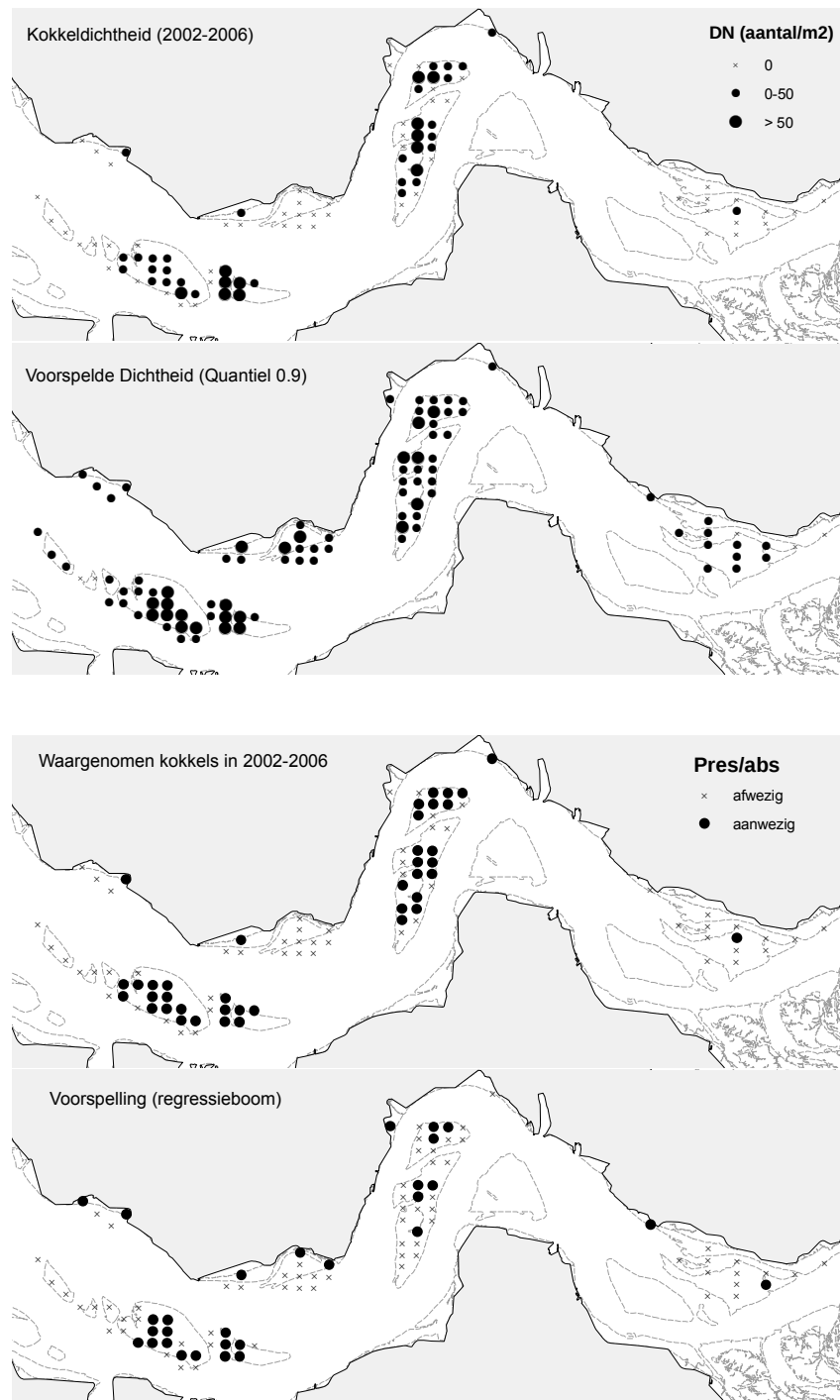


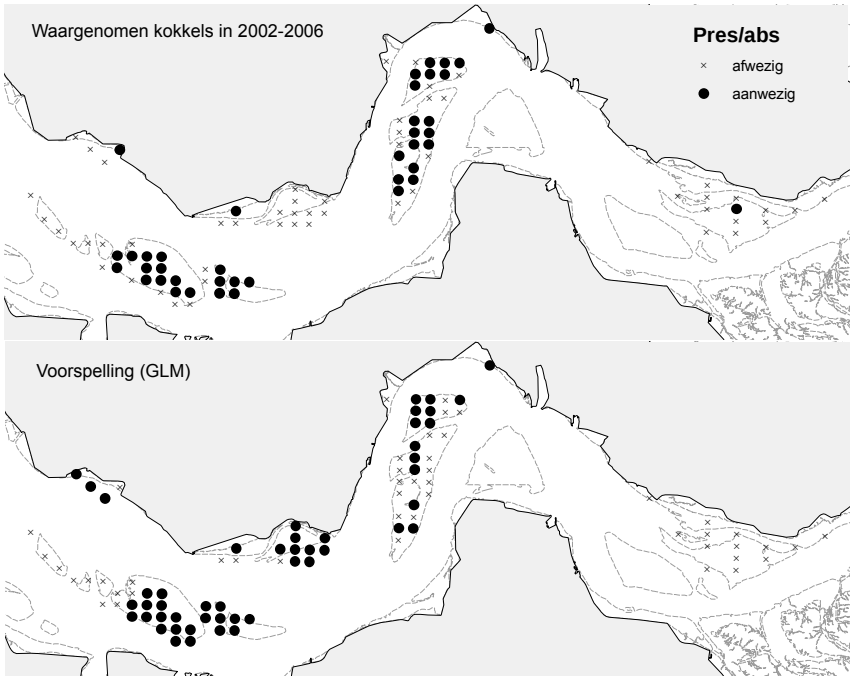
Figuur 13 Sensitiviteit vs specificiteit (het punt Density geeft de positie aan van de originele data uit het noorden). PAglm: GLM op basis van aan-/ afwezigheid, PAtree: regression tree op basis van aan-/ afwezigheid. Overige modellen zijn gemaakt op basis van dichtheden. QR: Quantielregressies, treelog: regression tree op loggetransformeerde data, treepois: regression tree met Poissonverdeling.

5 Discussie

Modelbeoordeling

Om het voorspellende vermogen van een model te kunnen beoordelen is het essentieel dat het model in een ander gebied wordt getest dan waar het is ontwikkeld. In deze studie hebben wij er voor gekozen om een model in het zuiden van de Westerschelde te ontwikkelen en om deze te testen in het noorden. In Nederland zijn voor kokkels in o.a. de Waddenzee, de Westerschelde en Oosterschelde in de afgelopen jaren meerdere habitatmodellen met behulp van Generalized Lineair modelling ontwikkeld (Bult et al, 2003; Kater & Baars, 2002; Ysebaert et al, 2002; Steenbergen et al, 2004;





Bijlage 4). Niet alleen van kokkels zijn in het verleden modellen ontwikkeld (Bijlage 5). Van de meeste, in het verleden ontwikkelde modellen is niet altijd even goed duidelijk hoe hun voorspellende vermogen is. Vaak wordt wel aangegeven hoe het model werkt in het gebied waar het is ontwikkeld (calibratie) maar is niet gekeken naar een onafhankelijk gebied. Zoals ook bleek in deze studie hoeft een goed beschrijvend vermogen echter niet per se te beteken dat een model ook goed werkt in een ander gebied.

Daar komt bij dat in het verleden modellen vaak enkel werden beoordeeld op basis van % verklaarde deviantie of op basis van % verklaarde variatie zoals in deze rapportage. Deze methodes zijn echter erg gevoelig voor grote afwijkingen naar boven of naar beneden door de kwadraten die in de formules zitten en daarom nogal onbetrouwbaar. Tijdens deze studie zijn meerdere methoden toegepast om modellen op basis van hun voorspellende vermogen te beoordelen.

Enkel op basis van R^2 is het moeilijk de modellen te beoordelen. Geen van de modellen bleek op deze basis een goed voorspellend vermogen te hebben. De GLM had dan wel een R^2 van 0.38, maar dit was met name het gevolg van incidentele hoge waarden die een groot effect uitoefenen op de berekende R^2 . Ook wanneer de gemeten waarden worden uitgezet tegen de voorspelde waarden, misschien wel de meest duidelijke vergelijking van de verschillende technieken, kan enkel worden geconcludeerd dat de modelvoorspellingen maar weinig overeenkomen met de gemeten waarden. R^2 is een standaardmaat die vaak wordt gebruikt bij modelbeoordeling, de grote spreiding van de kokkeldata is echter een probleem voor zowel de modelontwikkeling als voor de beoordeling van de modellen. Voor de meeste modellen geldt dat ze over het geheel genomen wel redelijk tot goed voorspellen waar kokkels wel en niet voorkomen, maar dat ze alleen vaak veel lagere dichtheden voorspellen. Dit is waarschijnlijk het gevolg van het grote aantal 0-waarden in de data, waardoor de gemiddelde schattingen omlaag worden getrokken.

In deze studie zijn daarom meerdere criteria gebruikt om het voorspellende vermogen van de verschillende modelleringstechnieken met elkaar te vergelijken en te beoordelen, maar ook deze geven niet een eenduidig antwoord.

Op basis van de totale dichtheden voor het gehele gebied zouden we kunnen concluderen dat de voorspellingen van GLM en quantielregressie ($\tau=0.9$) het best aansluiten bij de waargenomen dichtheden in het noorden. Op basis van het gemiddeld absoluut verschil (zowel niet als wel getransformeerd) bleek de GLM echter een relatief groot verschil te geven. Dat betekent dus dat deze methode op kleine schaal relatief grote verschillen geeft. De absolute verschillen tussen waarnemingen en de voorspellingen van de GAM zijn relatief kleiner dan die van de overige modellen, maar deze techniek geeft echter weer een onderschatting van de totale dichtheid.

De indeling in drie dichtheidsklassen geeft wel duidelijk aan welke modellen bijvoorbeeld alleen de lage dichtheden goed voorspellen en welke modellen ook de hoge dichtheden nog redelijk goed voorspellen. Het is echter met deze methode nog steeds moeilijk om aan te geven welk model het dichtste bij de waarheid komt qua voorspelling. Idealiter gebruiken wij hiervoor een grafiek die in 1 oogopslag laat zien hoe de modellen ten opzichte van elkaar presteren. Hiervoor hebben wij 3 methoden toegepast: het plotten van de voorspelde dichtheden tegen de gevonden dichtheden, de MDS methode en de sensitiviteit-specificiteit methode.

De eerste methode is het meest transparant: de data staan allemaal in 1 grafiek, maar de interpretatie is moeilijk als er geen model duidelijk beter is dan de andere. De vergelijking door middel van een multivariate statistische methode, MDS lijkt redelijk perspectief te bieden voor het vergelijken van modellen onderling als de exacte dichtheden van belang zijn. Nadeel van de MDS methode in vergelijking met de sensitiviteit-specificiteit methode is dat de uitkomst van de MDS methode onder grote invloed staat van de gebruikte voorbewerkingen (transformatie en keuze van afstandsindex). Voordeel van de MDS methode is dat zij niet alleen beperkt is tot aanwezig-/afwezigheidsdata. Sensitiviteit-specificiteit is een methode die rechttoe rechtaan is en kan op niveau van aanwezig-/afwezigheid de prestaties van verschillende modellen goed beoordelen. Deze methode is dan ook het makkelijkste te interpreteren.

Op basis van de MDS en de sensitiviteit-specificiteit bleken de beide Poisson-regressies (GLM & regressiebomen) het voorkomen van de kokkels het beste te voorspellen.

Oorzaken van inaccurate voorspellingen

Uit de vergelijkende studie naar de verschillende technieken om kokkels mee te modelleren kan worden geconcludeerd dat het ene model beter is dan het andere, maar dat er geen model is dat 'de perfecte voorspelling' geeft. Estuaria zijn zeer dynamische ecosystemen en bodemdiergemeenschappen vertonen dan ook grote fluctuaties in ruimte en tijd (Ysebaert & Herman, 2002). Door het gebruik van meerjarige gemiddelden is er tot op zekere hoogte wel rekening gehouden met de natuurlijke dynamiek, maar toch kunnen er nog aanzienlijke fouten op treden bij voorspellingen.

Bij de modelvoorspelling kunnen 2 soorten fouten optreden:

1. Er worden wel kokkels voorspeld, maar er zijn in werkelijkheid geen kokkels aanwezig
2. Er zijn wel kokkels aanwezig, maar model voorspelt dat er geen kokkels aanwezig zijn

Modellen veronderstellen in principe dat er een evenwicht bestaat tussen de doelsoort en zijn omgeving (Barry & Elith, 2006). Dit houdt in dat een model er vanuit gaat dat op elke potentiële (qua habitat) geschikte locatie kokkels aanwezig zullen zijn. Factoren die niet zijn meegenomen in het model (bv. broedval, wintersterfte, extreme omstandigheden, etc) kunnen echter een rol hebben gespeeld (b.v. bij het vestigen of de overleving van broed) waardoor de kokkels op het moment van monsteren niet aanwezig zijn. De dynamiek van kokkels wordt gekenmerkt door enorme verschillen in broedval tussen verschillende jaren (Kamermans et al, 2003). Of een

broedval wel of niet succesvol is, hangt onder andere af van natuurlijke factoren zoals het optreden van een strenge winter en jaarlijkse variatie van predatiedruk op jonge levenstadia (Beukema & Dekker, 2005; Beukema et al, 2001). In 'slechte' jaren is de predatiedruk op de larven waarschijnlijk hoog waardoor er minder broedval plaatsvindt. Dit betekent dat op een locatie die qua abiotische factoren wel geschikt is voor kokkels niet per definitie ieder jaar ook veel (1-jarige) kokkels worden aangetroffen. In de periode 2002-2006 waren er relatief veel 1-jarige kokkels in 2004. De rest van de jaren was het aantal kokkels aanzienlijk minder. Voor deze studie werd per bemonsterd station een gemiddelde van de kokkels berekend gebaseerd op monsters van 2002 tot 2006.

Een oplossing voor dit probleem zou kunnen zijn om modellen te construeren die of meer variabelen bevatten en/of meer rekening houden met stochastische fluctuaties in de relevante variabelen. Een andere mogelijkheid zou kunnen zijn om meer dynamische modellen te maken (bv. vergelijkbaar met ECOWASP of ERSEM). Deze modellen zijn echter wel een stuk arbeidsintensiever.

Tevens kunnen oNjuiste voorspellingen kunnen worden gedaan doordat ontstaan wanneer niet alle te veel variatie in de modellen data niet is wordt verklaard door de gebruikte modellen. Dit kan onder andere komen doordat de gebruikte onderliggende abiotische gegevens niet accuraat en/of volledig zijn (Barry & Elith, 2006). De abiotische gegevens zijn vaak uit een andere periode dan de biotische en er wordt vaak vanuit gegaan dat er abiotisch niet veel verandering is geweest opgetreden, maar voor sommige variabelen zou is deze aanname misschien niet houdbaar zijn. Ook is het mogelijk dat de abiotische variabelen onder sommige omstandigheden zeer goede proxies (i.e. indicatoren, maar geen echte causale factoren) zijn voor de feitelijke causale factoren, maar dat dit onder enigszins gewijzigde omstandigheden niet meer hoeft te gelden.

Een belangrijke vraag in het onderzoek naar de verspreiding van (benthische) organismen in relatie tot hun omgeving is, welke parameters, of welke combinatie van parameters moeten worden gemeten (Constable, 1999). Constable (1999) concludeert dat sediment, nutriënten en voedselvoorziening en waterbeweging de belangrijkste bepalende fysische factoren zijn voor de verspreiding van macro-benthos. Daarnaast zijn in estuaria zoals de Westerschelde saliniteit en droogvalduur (diepte) belangrijke factoren (Constable, 1999; Ysebaert et al, 2002). Voor deze studie zijn de parameters stroomsnelheid, diepte, zoutgehalte en sediment gebruikt om het voorkomen van kokkels te verklaren. Met betrekking tot de accuraatheid van deze gegevens valt nog wel wat winst te behalen. De gegevens zijn allen gemeten cq berekend (stroomsnelheid) voor 1 jaar. De platen in het intergetijdengebied zijn echter mobiel en de randen kunnen in een paar jaar wel tientallen tot soms zelfs honderden meters verplaatsen (Ysebaert, 2002). Door deze dynamiek veranderen diepte, stroomsnelheid en sedimentsamenstelling voortdurend en het verdient dan ook de aanbeveling om deze

parameters minimaal jaarlijks te monitoren. Zoutgehalte is voor kokkels een beperkende factor. Jaarlijks sterven in de winter de meeste kokkels in het oosten van de Westerschelde door hoge afvoeren van de rivier de Schelde. In jaren met extreme afvoeren zullen meer kokkels sterven dan in jaren met lage afvoeren. Voor zoutgehalte wordt dan ook aanbevolen om gedurende het jaar te meten om zo de extremen te kunnen bepalen. Sedimentgegevens zijn nu enkel aanwezig van de bemonsterde stations waardoor het niet mogelijk is goede gebiedsdekkende kaarten te maken.

Representativiteit van kokkeldata (schaal)

Kokkeldichtheden worden tijdens de jaarlijkse bemonstering geschat door middel van kleine hapjes op ca 250 vaste locaties. Voor het doel van de bemonstering in de Westerschelde, namelijk het schatten van de totale biomassa van kokkels in dit gebied, is dit een efficiënte methode. Wanneer deze monsterpunten voor toepassingen op kleinere schaal worden gebruikt zijn de genomen hapjes misschien niet altijd even representatief. Kokkels komen namelijk niet homogeen in een voor hen geschikt gebied voor. Stratificatie van de bemonstering kan veel betere dichtheidsbepalingen opleveren: idealiter zou er eerst een grove monsterring plaats moeten vinden om gebieden met hoge en lage dichtheden te bepalen (dit zou eventueel ook op basis van gegevens uit het verleden kunnen gebeuren of met behulp van Sidescan Sonar). Daarna kan men bepalen waar en hoe dicht men gaat bemonsteren om een zo goed mogelijk beeld van het totale bestand te krijgen. Het is overigens aannemelijk dat stations die dicht bij elkaar liggen gecorreleerd zijn. Hierdoor kunnen schattingen negatief beïnvloed worden en indien men voor deze ruimtelijke autocorrelatie corrigeert, kunnen de modelvoorspellingen beter worden. Bij multivariate analyses is dit probleem nog amper aangepakt, maar bij univariate analyses bestaan er thans mogelijkheden om ruimtelijke of temporele correlatie mee te nemen in modelberekeningen (zogenaamde Generalized Linear Mixed Models en Generalized Additive Mixed Models; deze vielen thans buiten de scope van dit onderzoek).

Omgaan met onzekerheden

Het uitvoeren van effectenstudies evenals andere waarnemingsstudies gaat altijd gepaard met onzekerheden. Ook na toepassing van statistische habitatmodellen zijn deze onzekerheden niet geheel weg te nemen. Een manier voor beheerders en beleidsmakers om met deze onzekerheden om te gaan is door te kijken naar verschillende modellen en beslissingen te baseren op de uitkomsten van meerdere modellen. Wijzen alle modellen in dezelfde richting dan is de uitkomst waarschijnlijker. Uit deze studie blijkt, mede gezien de natuurlijke dynamiek, dat het onwaarschijnlijk is dat er ooit een perfect habitatmodel gemaakt zal worden. Dit is ook in ander onderzoek al eens geopperd (Stewart-Oaten, 1996). Het is echter niet zo dat de modellen die ontwikkeld zijn niet bruikbaar zijn voor het nemen van (beleids)besluiten. Dankzij de verschillende methodieken van modelbeoordeling kan de kwaliteit van modellen op

verschillende niveaus van nauwkeurigheid worden beoordeeld. Zo is gebleken dat naarmate je het voorkomen van kokkels nauwkeuriger wilt voorspellen, de onzekerheid groter wordt. Op niveau van exacte dichtheidsvoorspellingen dient een ruime marge te worden genomen, terwijl op niveau van aan-/afwezigheid behoorlijk goede voorspellingen kunnen worden gedaan.

De modellen in deze studie zijn gemaakt voor kokkels, slechts een van de vele soorten die voorkomen in het betreffende gebied. Voor andere soorten gelden misschien andere variabelen en het moge duidelijk zijn dat niet alle mogelijke abiotische variabelen gemeten kunnen worden. Daarom verdient het aanbeveling om monitoringsprogramma's breed en flexibel op te stellen en ook langere tijd te meten om een idee te krijgen van de natuurlijke variatie van de betreffende variabelen. Voor het nemen van beslissingen zijn flinke veiligheidsmarges nodig zolang men de natuurlijke fluctuaties nog niet afdoende begrijpt en/of kan voorspellen. Tot het zover is, blijven langlopende monitoringsprogramma's essentieel om het systeem beter te begrijpen en effecten van veranderingen te kunnen beoordelen.

6 Literatuur

- Amar, A, Redpath, SM (2005) Habitat use by Hen Harriers *Circus cyaneus* on Orkney: implications of land-use change for this declining population. *Ibis* 147, 37–47
- Barry S, Elith J (2006) Error and uncertainty in habitat models. *Journal of applied ecology* 43: 413-423
- Beukema JJ, Dekker R, Essink K, Michaelis H (2001) Synchronized reproductive success of the main bivalve species in the Wadden Sea: causes and consequences. *Mar Ecol Prog Ser* 211: 143-153.
- Beukema JJ, Dekker R (2005) Decline of recruitment success in cockles and other bivalves in the Wadden Sea: possible role of climate change, predation on postlarvae and fisheries. *Mar Ecol Prog Ser* 287: 149-167.
- Bult T, Kater B, Baars D (2003) Habitatmodellen voor de commerciële schelpdieren in de Westelijke Waddenzee. RIVO-rapport: C026/03
- Bult TP, Ens BJ, Baars D, Kats R, Leopoldt M, (2004) B3: Evaluatie van de meting van het beschikbare voedselaanbod voor vogels die grote schelpdieren eten. RIVO rapport nr: C018/04.
- Cade BS, Noon BR (2003) A gentle introduction to quantile regression for ecologists. *Front Ecol Environ* 1:412–420
- Constable AW (1999). Ecology of benthic macro-invertebrates in soft-sediment environments: a review of progress towards quantitative models and predictions. *Austr J Ecol* 24: 452-476
- De'ath G, Fabricius KE (2000) Classification and regression trees: A powerful yet simple technique for ecological data analysis. *Ecology* 81:3178-3192
- Graveland J (2005) Fysische en ecologische kennis en modellen voor de Westerschelde, wat is er beleidsmatig nodig en wat is er beschikbaar voor de m.e.r. verruiming vaargeul? RIKZ-rapport: 2005.018
- Guisan A, Zimmermann NE (2000) Predictive habitat distribution models in ecology. *Ecol Model* 135: 147- 186.

Hastie, T. and R. Tibshirani, 1990. Generalized Additive Models. Chapman & Hall, London.

Kamermans P, Bult T, Kater B, Baars D, Kesteloo J, Perdon J, Schuiling E (2003) Invloed van natuurlijke factoren en kokkelvisserij op de dynamiek van bestanden aan kokkels (*Cerastoderme edule*) en nonnen (*Macoma Balthica*) in de Waddenzee, Ooster- en Westerschelde. RIVO-rapport: C058/03

Kater BJ, Baars JMDD (2002) De Oosterschelde werken en de relatie tussen abiotische factoren en biomassa kokkels. RIVO-rapport: C055/02.

Kesteloo et al 2006.

Kruskal J (1964) Multidimensional scaling by optimizing goodness of fit to a nonmetric hypothesis. Psychometrika 29:1-27

Koenker R (2006). quantreg: Quantile Regression. R package version 3.82.

<http://www.econ.uiuc.edu/~roger/research/rq/rq.html>

Male K van der, Schouwenaar B, (2003). Stroomsnelheden en andere fysische parameters in de Oosterschelde en in de Westerschelde. Werkdocument RIKZ/AB/2003.813x.

McCullagh, P, Nelder JA (1989). Generalized linear models (2nd ed). London, Chapman and Hall.

Oude Voshaar, JH (1995) Statistiek voor onderzoekers. Wageningen, Wageningen pers.

Pearce J, Ferrier, S (2000) Evaluating the predictive performance of habitat models developed using logistic regression. Ecological modelling 133: 225-245.

Shepard R (1962) The analysis of proximities: multidimensional scaling with an unknown distance function. Psychometrika 27: 125-140

Steenbergen J, Baars JMDD, Bult TP (2004) Habitatmodellen voor kokkels in de Westerschelde. RIVO-rapport: C055/04.

Steenbergen J, Escaravage V (2006) Baseline studie MEP-MV2. Lot 2: bodemdieren, eindrapportage campagnes 2004-2005. Wageningen IMARES-rapport: C053/06.

Stewart-Oaten A (1996) Problems in the analysis of environmental monitoring data. In: *Detecting ecological impacts: Concepts and applications in coastal habitats* (eds. R.J. Schmitt & C.W. Osenberg), pp 109-131. Academic press, San Diego.

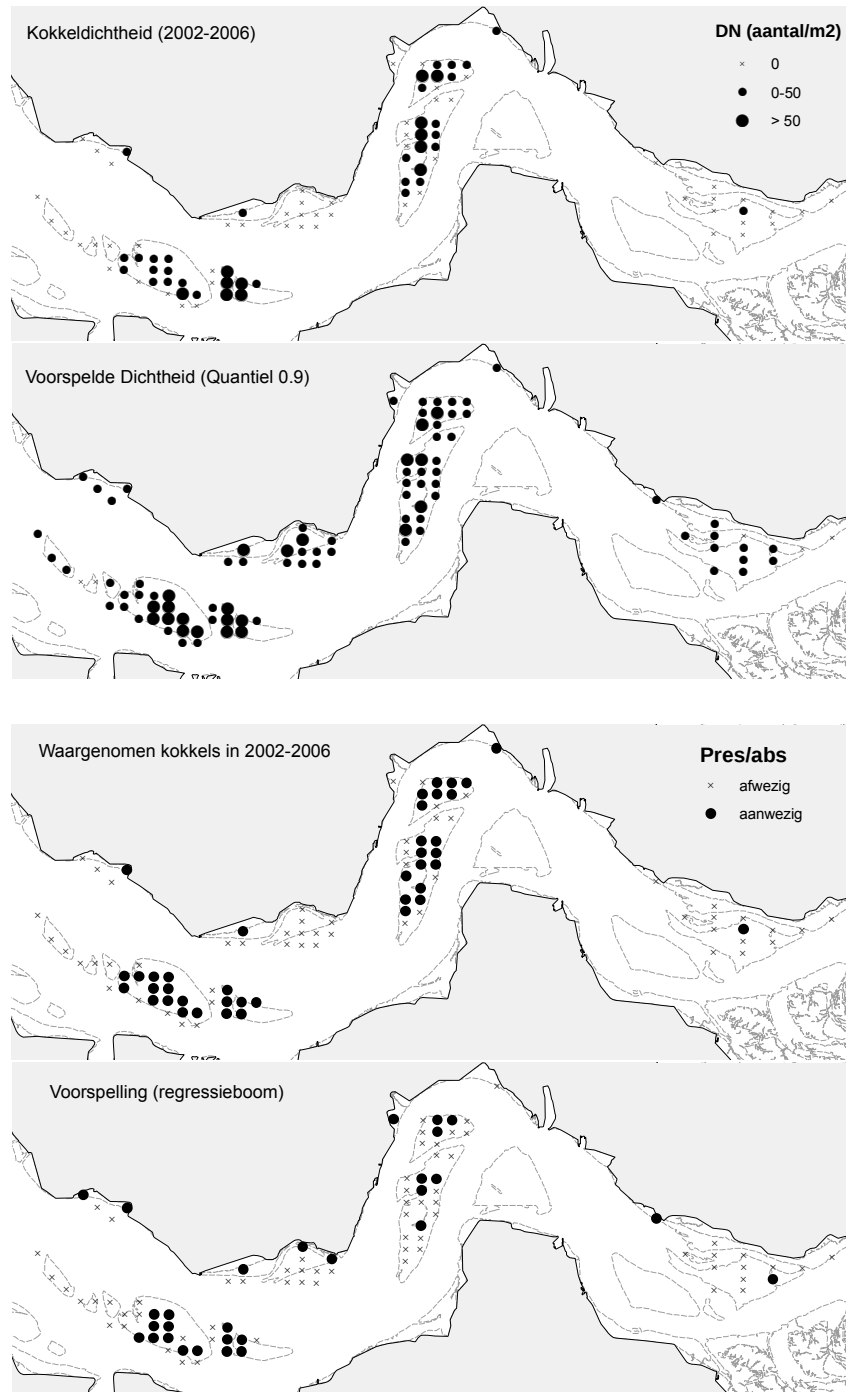
-
- Therneau TM, Atkinson B, R port by Brian Ripley <ripley@stats.ox.ac.uk>. (2006). rpart: Recursive Partitioning. R package version 3.1-32. S-PLUS 6.x original at <http://mayoresearch.mayo.edu/mayo/research/biostat/splusfunctions.cfm>
- Thuiller W, Araújo MB, Lavorel S (2003) Generalized models vs. classification tree analysis: Predicting spatial distributions of plant species at different scales. *Journal of vegetation science* 14: 669-680.
- R Development Core Team (2006) R: A language and environment for statistical computing. R Foundation for Statistical Computing, Vienna, Austria. ISBN 3-900051-07-0, URL <http://www.R-project.org>.
- Rykiel EJ (1996) Testing ecological models: the meaning of validation. *Ecological modelling* 90: 229-244.
- Wood (2000) Modelling and Smoothing Parameter Estimation with Multiple Quadratic Penalties. *JRSSB* 62(2):413-428
- Wood (2001) mgcv: GAMs and Generalized Ridge Regression for R. *R-News* 1(2):20-25
- Young LJ, Campbell NL, Capuano GA (1999) Analysis of overdispersed count data from single-factor experiments: a comparative study. *J. Agr. Biol. Env. Stats* 4: 258-275.
- Ysebaert T, Meire P, Herman PMJ en Verbeek H (2002) Macrobenthic species response surfaces along estuarine gradients: prediction by logistic regression. *Mar. Ecol. Prog. Ser.* 225: 79-95.
- Ysebaert T, Herman PMJ (2002) Spatial and temporal variation in benthic macrofauna and relationships with environmental variables in an estuarine, intertidal soft-sediment environment. *Mar. Ecol. Prog. Ser.* 244: 105-124.
- Zeman P, Lynen G (2006) Evaluation of four modelling techniques to predict the potential distribution of ticks using indigenous cattle infestations as calibration data. *Exp App Acar* 39, 163-176

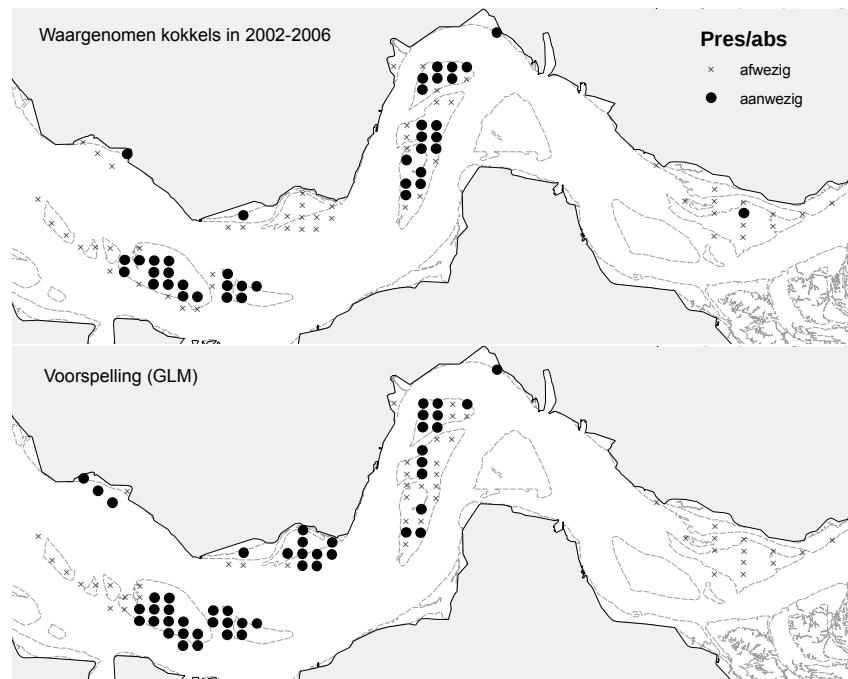
7 Bijlagen

Bijlage 1 Modeluitkomsten GLM & Quantielregressies

Bijlage 2 Modeluitkomsten GAM

Bijlage 3 GIS-plaatjes van de waargenomen dichtheden kokkels (gemiddeld 2002-2006) en de voorspelde dichtheden.





Bijlage 4 Overzicht bestaande habitatmodellen voor kokkels in de Nederlandse kustwateren.

Bijlage 5 Algemene discussie over bestaande habitatmodellen

Bijlage 1 Modeluitkomsten GLM & Quantielregressies

Uitleg bij de tabellen:

Estimate: De geschatte waarde of richtingscoëfficiënt

Std. Error: de standaard fout van de geschatte waarde

t value: de t-waarde die bij deze schatting hoort, hoe hoger hoe meer significant.

Pr (>|t|): de waarschijnlijkheid dat de waarde afwijkt van 0.

Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

GLM: Dichtheden van kokkels

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	-1.04 E+02	3.19 E+01	-3.27	0.0015 **
zoutgem	6.73 E+00	2.31 E+00	2.91	0.0045 **
zoutgem2	-1.24 E-01	4.25 E-02	-2.92	0.0045 **
diepte	4.46 E-02	1.57 E-02	2.85	0.0054 **
diepte2	-1.09 E-04	3.54 E-05	-3.08	0.0027 **
SD50phi	1.04 E+01	4.70 E+00	2.21	0.0299 *
SD50phi2	-1.91 E+00	8.61 E-01	-2.22	0.0289 *

GLM: aan- cq afwezigheid kokkels

	Estimate	StdErr	ChiSq	Prob ChiSq
Intercept	-96.8913	30.6595	9.99	0.0016 **
zoutgem	5.315126	2.2307	5.68	0.0172 *
zoutgem2	-0.09935	0.0417	5.67	0.0173 *
SD50phi	18.36358	5.6089	10.72	0.0011 **
sd50phi2	-3.00214	0.9545	9.89	0.0017 **

GLM Dichtheden kokkels (data=loggetransformeerd)

	Estimate	Std. Error	t value	Pr(> t)
(Intercept)	-2.07 E+01	6.16 E+00	-3.36	0.0011 **
zoutgem	9.81 E-01	4.66 E-01	2.11	0.0380 *
zoutgem2	-1.76 E-02	9.33 E-03	-1.89	0.0622 .
diepte	-2.08 E-03	2.60 E-03	-0.8	0.4252
diepte2	-5.80 E-07	2.95 E-06	-0.2	0.8446
SD50phi	6.63 E+00	2.06 E+00	3.22	0.0018 **
SD50phi2	-1.12 E+00	3.78 E-01	-2.96	0.0039 **

Quantielregressies (dichtheid kokkels)

De Tau-waarde geeft het quantiel weer dat gebruikt is voor het berekenen van het model. De lagere quantielen worden niet getoond omdat deze alleen maar nullen bevatten en dus geen modelvoorspelling opleverden.

Tau=0.5

	Value	Std. Error	t value	Pr(> t)
(Intercept)	-8.28821	6.57909	-1.25978	0.2110
zoutgem	0.45949	0.36966	1.243	0.2171
zoutgem2	-0.00771	0.00774	-0.99599	0.3219
diepte	0.00229	0.00289	0.78987	0.4317
diepte2	0	0	-1.73196	0.0867 .
SD50phi	0.76504	2.06903	0.36976	0.7124
SD50phi2	0.04609	0.39311	0.11725	0.9069

Tau=0.6

	Value	Std. Error	t value	Pr(> t)
(Intercept)	-6.76069	4.4292	-1.52639	0.1304
zoutgem	0.07672	0.18568	0.41319	0.6805
zoutgem2	0.00022	0.00405	0.05322	0.9577
diepte	-0.00242	0.00214	-1.12854	0.2621
diepte2	0	0	0.02505	0.9801
SD50phi	4.36024	2.83084	1.54026	0.127
SD50phi2	-0.69246	0.56804	-1.21903	0.2260

Tau=0.7

	Value	Std. Error	t value	Pr(> t)
(Intercept)	-10.5017	8.404	-1.24961	0.2147
zoutgem	0.45045	0.66512	0.67724	0.5000
zoutgem2	-0.00754	0.01366	-0.55217	0.5822
diepte	-0.00543	0.00489	-1.11112	0.2695
diepte2	0	0.00001	0.42594	0.6712
SD50phi	4.93182	2.42805	2.03118	0.0452 *
SD50phi2	-0.82469	0.44832	-1.83951	0.0691 .

Tau=0.8

	Value	Std. Error	t value	Pr(> t)
(Intercept)	-18.1538	7.09723	-2.55787	0.0122 *
zoutgem	0.71839	0.54873	1.3092	0.1938
zoutgem2	-0.01162	0.01192	-0.97518	0.3321
diepte	-0.00173	0.00414	-0.41747	0.6773
diepte2	0	0	-0.57303	0.5681
SD50phi	7.46918	2.39858	3.114	0.0025 **
SD50phi2	-1.22333	0.44998	-2.71866	0.0079 **

Tau=0.9

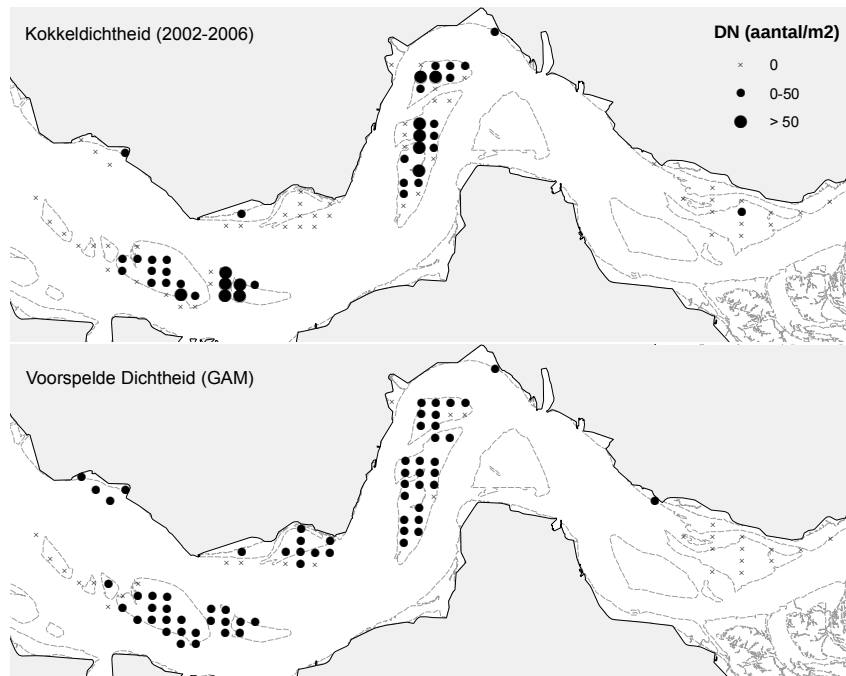
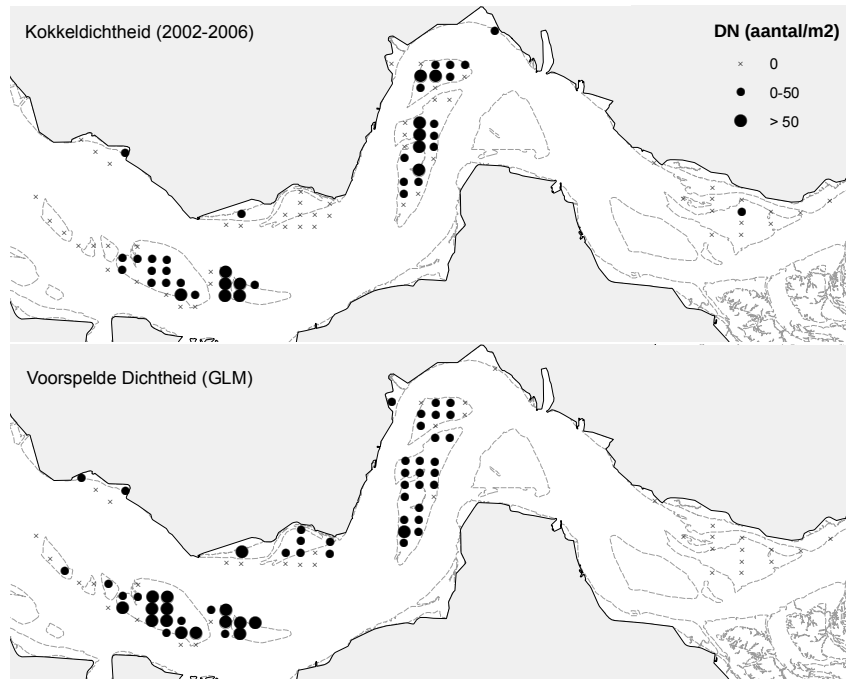
	Value	Std. Error	t value	Pr(> t)
(Intercept)	-25.19951	7.05428	-3.57223	0.0006 ***
zoutgem	1.39731	0.57069	2.44847	0.0163 *
zoutgem2	-0.02541	0.01173	-2.16563	0.0330 *
diepte	-0.00192	0.00406	-0.47316	0.6373
diepte2	0	0	-0.60015	0.5499
SD50phi	7.62049	2.85253	2.67148	0.0090 **

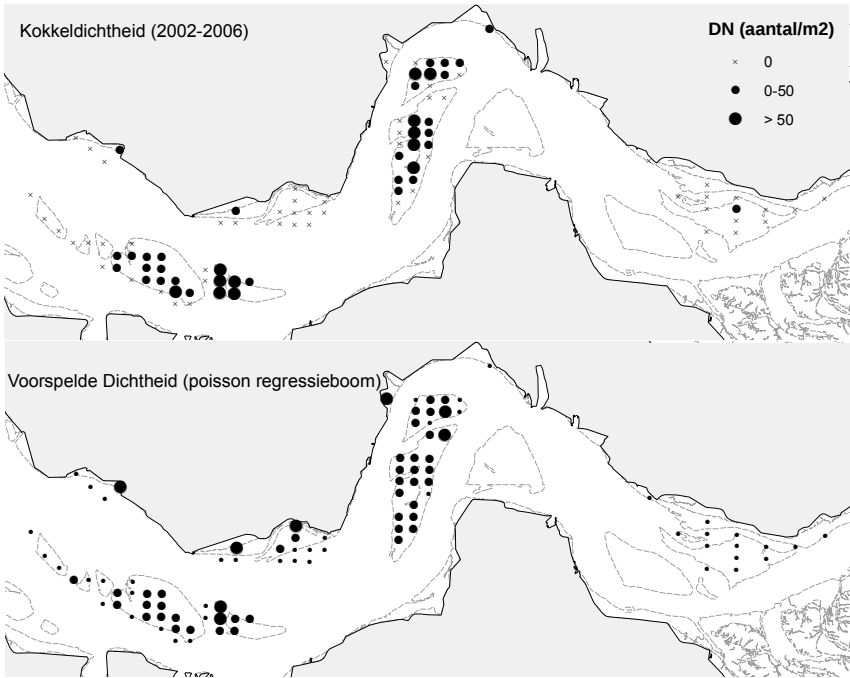
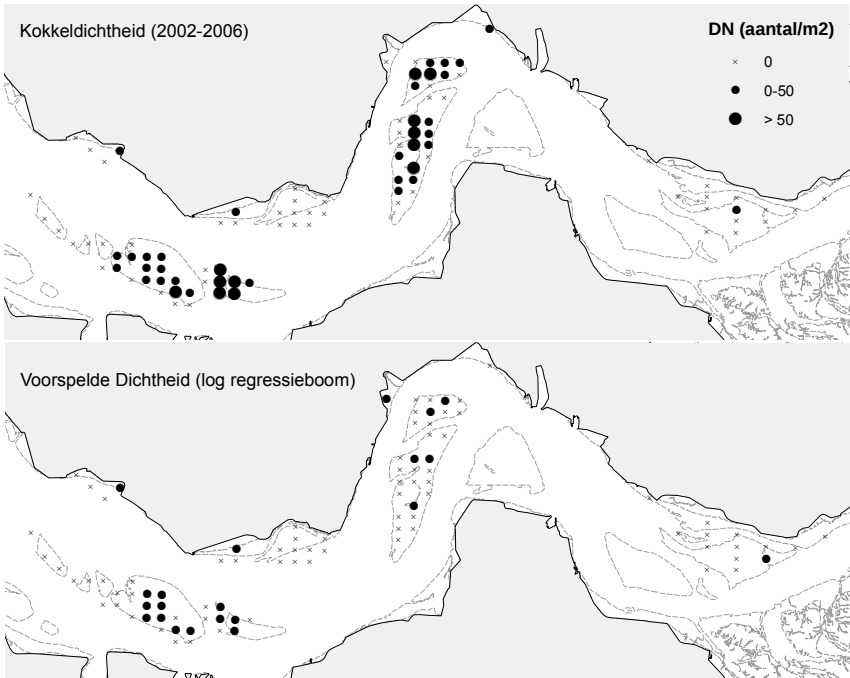
SD50phi2	-1.30179	0.5415	-2.40404	0.0183 *
----------	----------	--------	----------	----------

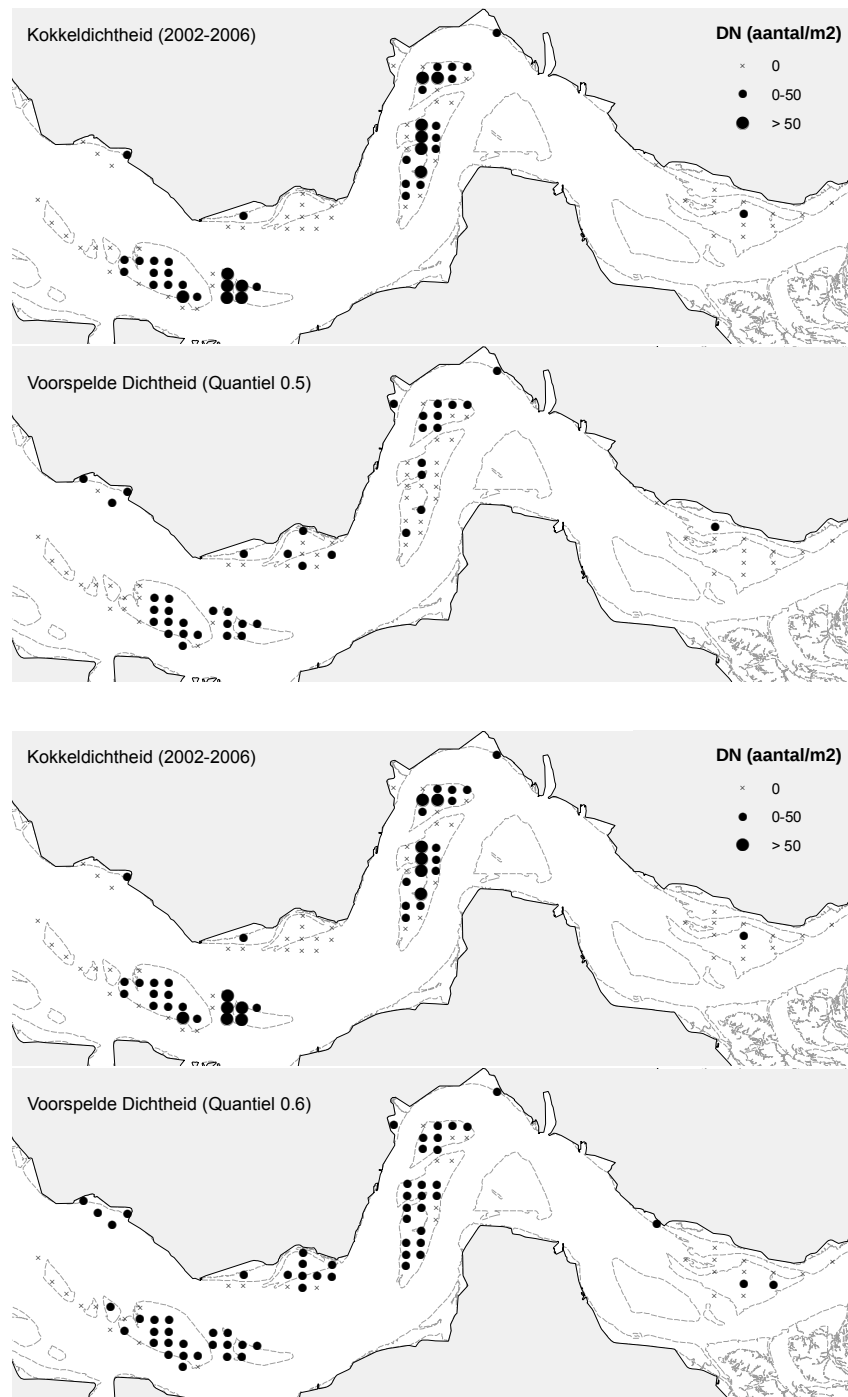
Bijlage 2 Modeluitkomsten GAM

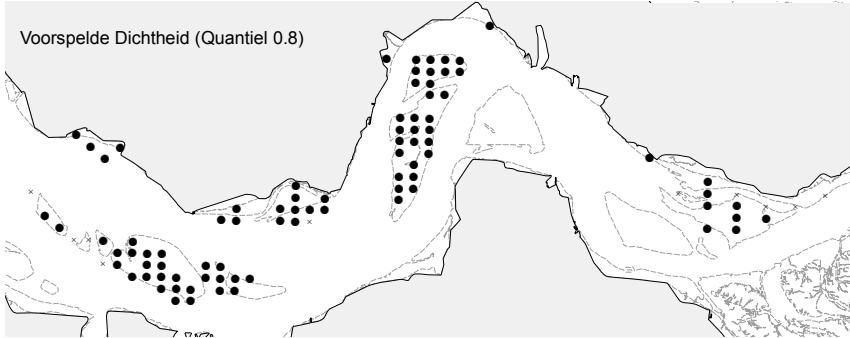
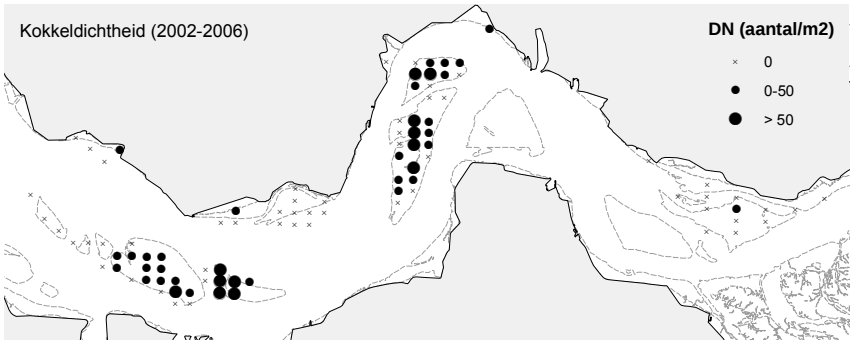
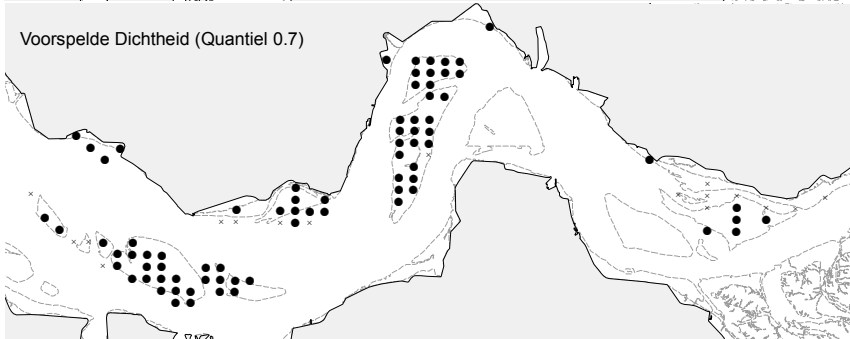
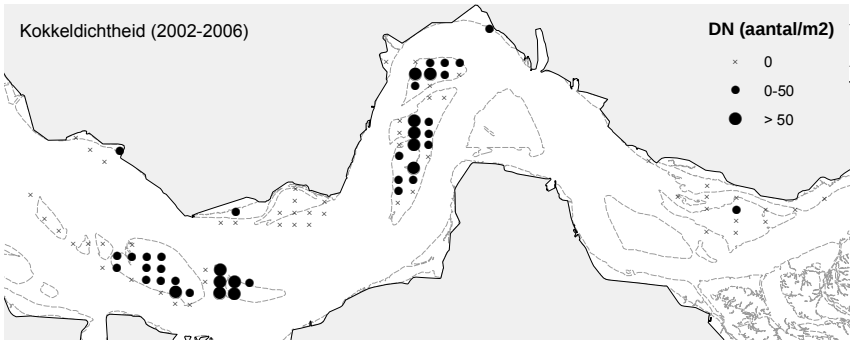
Signif. codes: 0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

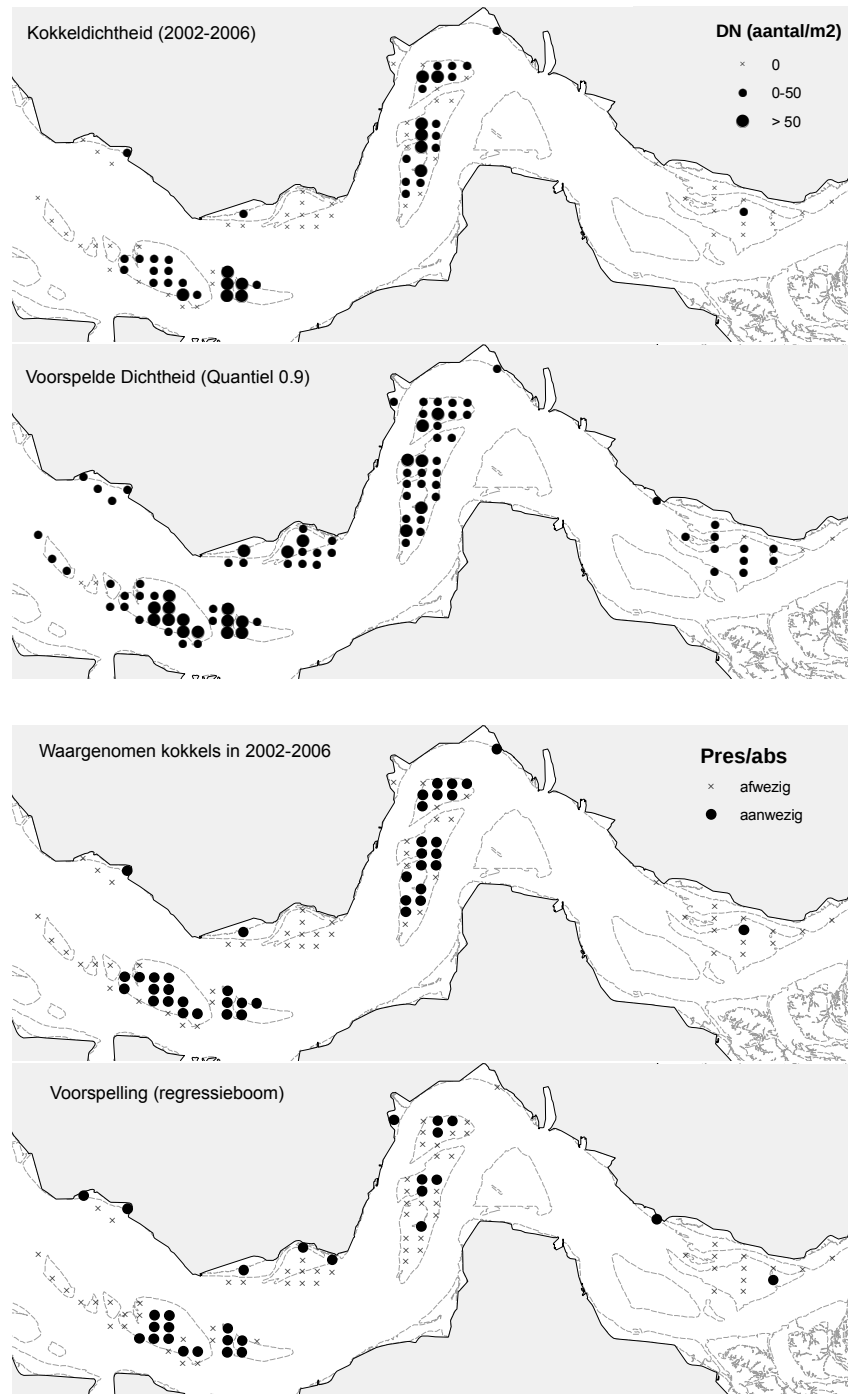
	edf	Est.ran k	F	p-value
s(SD50phi)	6.354	9	4.08	0.000275 ***
s(zoutgem)	6.023	9	3.541	0.001061 **
s(diepte)	3.576	8	4.178	0.000362 ***
s(ssh)	5.624	9	3.568	0.000991 ***

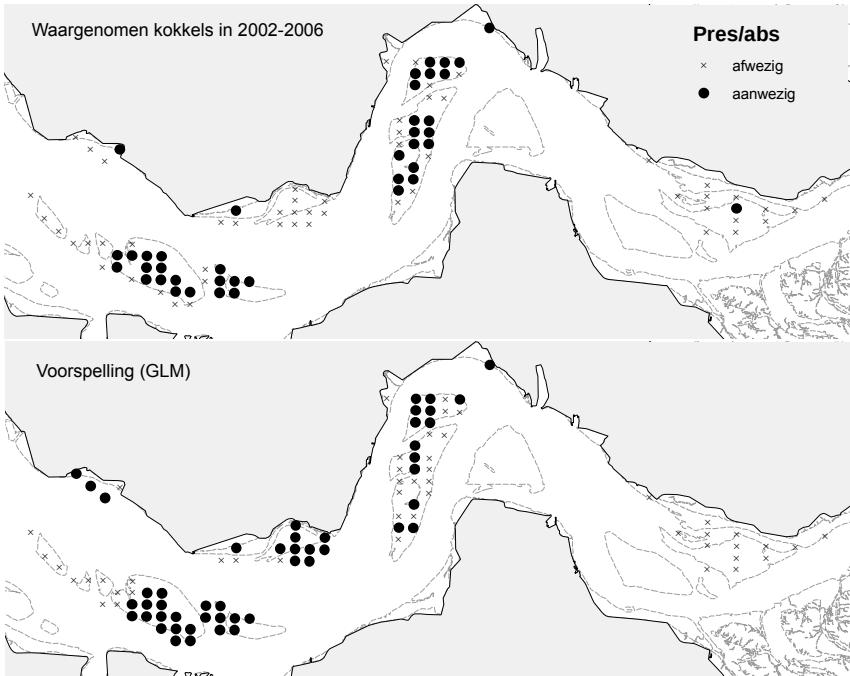
Bijlage 3 GIS-plaatjes van de waargenomen dichtheden kokkels (gemiddeld 2002-2006) en de voorspelde dichtheden.











Bijlage 4 Overzicht bestaande habitatmodellen voor kokkels in de Nederlandse kustwateren.

Ysebaert et al (2002) hebben modellen ontwikkeld voor 10 macrobentische soorten in de Westerschelde, waaronder de kokkel. Hiertoe voerden zij een logistische regressie uit op een grote dataset om de kans op voorkomen van onder andere kokkels als een respons op de abiotische factoren: saliniteit, diepte, stroomsnelheid en sediment te analyseren. De berekeningen zijn uitgevoerd op een dataset in zowel het litoraal als het sublitoraal in de gehele Westerschelde.

Bult et al (2003) hebben habitatmodellen ontwikkeld voor kokkels in de Waddenzee om eventuele effecten van extra spuimiddel in de afsluitdijk te kunnen schatten. Door middel van deviatieanalyses/Poisson regressie zijn habitatmodellen ontwikkeld voor de verspreiding van kokkels ($\# \cdot m^{-2}$ en $g \cdot m^{-2}$) in de westelijke Waddenzee. Tevens is een logistische regressie uitgevoerd om de kans op voorkomen te bepalen. De abiotische factoren die zijn gebruikt zijn: gemiddeld zoutgehalte, maximale stroomsnelheid, hoogteligging t.o.v. NAP en mediane korrelgrootte.

Kater & Baars (2002) hebben met behulp van habitatmodellen gekeken naar verschuiving van geschikte kokkelgebieden als gevolg van de bouw van de Oosterscheldekering. Hiertoe zijn modellen ontwikkeld die de biomassa van kokkels beschrijven aan de hand van abiotische factoren (droogvalduur, saliniteit, stroomsnelheid en mediane korrelgrootte) en is een vergelijking gemaakt tussen een periode voor de bouw van de Oosterscheldekering en twee perioden na de Oosterscheldekering. De modellen zijn ontwikkeld met behulp van generalized linear modeling GLM en een Poisson verdeling.

Steenbergen et al (2004) ontwikkelden een habitatmodel voor kokkels in de Westerschelde om de verspreiding van kokkels in relatie tot habitat te beschrijven. Dit model is verder uitgewerkt in onderliggende rapportage waarbij verschillende criteria zijn gebruikt en verschillende gebieden. Indertijd zijn 2 modellen gemaakt, voor aantallen en biomassa kokkels in relatie tot droogvalduur, diepte, zoutgehalte stroomsnelheid en slibgehalte.

Voor deze modellen is nagegaan of de modellen getest zijn op genericiteit en zo ja hoe generiek deze modellen zijn

Bijlage 5 Algemene discussie over bestaande habitatmodellen***Uit Meesters et al 2006 (in prep.)***

Hoewel de uitkomst van habitatanalyses vaak kans op voorkomen of mogelijke dichtheden (aantallen of biomassa) van organismen betreft, moeten effecten van veranderingen worden gelezen in termen van een veranderde habitatgeschiktheid.

Zo zijn er vele habitatmodellen gemaakt die alle als voornaamste eigenschap hebben dat ze eigenlijk alleen geldig zijn voor het geografisch gebied waar de gebruikte data uit afkomstig zijn. Dit wordt mede veroorzaakt door het feit dat er mathematische beperkingen zijn aan de modellen omdat ze vaak mindere geldigheid bezitten indien de waarden van de sturende (i.e. verklarende) variabelen buiten het interval vallen van de meetdata waarmee de habitatmodellen gemaakt zijn.

- Spisula-model voor Nederlandse kustzone (Craeymeersch 2001).
- Modellen voor laagwatersverspreiding van wadvogels (Brinkman & Ens 1998), onderzoek in het kader van de bodemdalingsstudie (NAM 1998).
- Modellen voor verspreiding van wadvogels in de Waddenzee (Smit et al 2003)
- Modellen voor de verspreiding van benthos over het gehele NCP (Brinkman et al 2001)
- Habitatgeschiktheidmodel voor meerjarige mosselbanken in de Nederlandse Waddenzee (Brinkman & Bult 2002)
- Habitatgeschiktheidmodel voor kokkelbanken in de Nederlandse Waddenzee (Kater et al 2003)
- Habitatgeschiktheidmodel voor kokkelbanken in de Westerschelde (Steenbergen et al 2004)
- Modellen voor verspreiding van wadvogels in de Westerschelde (Brinkman, Meesters, Ens & Dijkman 2005)
- Habitatmodel voor klein zeegras in de Nederlandse Waddenzee (De Jonge & De Jong 1999, Essink et al 2003)
- Habitatmodel voor verspreiding van gewone zeehonden in de Nederlandse Waddenzee (Hetmank 2004)

Bij vrijwel al deze modellen is de verspreiding van soorten, bv. van kokkel- en mosselbanken, gerelateerd aan de abiotische omstandigheden (slibgehalte, mediane korrelgrootte, droogvalduur of diepte t.o.v. NAP). Stroomsnelheid en orbitaalsnelheden zijn soms een sturende factor, maar vaak ontbreken daarvoor de nodige data.

De modellen hebben vaak als eerste doel om aan te geven wat de meest waardevolle gebieden zijn of meest kansrijke locaties zijn voor het (toekomstig) voorkomen van de betreffende soort

of habitat. Een tweede doel is een eerste schatting te geven van de verwachte veranderingen in biomassa- of aantalsdichtheid, of trefkans op bijvoorbeeld een mosselbank wanneer er veranderingen optreden in de sturende variabelen. Deze benadering is bijvoorbeeld gevolgd bij de studie naar de effecten van de aanleg van een vliegveld in zee. Een groot eiland heeft gevolgen voor de slibhuishouding, en daarmee vermoedelijk ook voor de slibgehalten van het sediment. Omdat slib in de habitatbeschrijvingen een sturende grootheid is, kan de verandering van slibgehalte gebruikt worden om te schatten wat de verandering aan habitat zal zijn indien de slibgehalten van de bodem veranderen. Omdat de habitatgeschiktheid vaak uitgedrukt wordt in (verwachte) biomassadichtheid wordt de verandering óók in die eenheid uitgedrukt. Adaptief gedrag wordt hierbij dus buiten beschouwing gelaten, evenals allerlei relaties tussen de voorkomende andere soorten.

Habitatgeschiktheidsmodellen worden steeds meer gebruikt, maar hun voorspellend vermogen is beperkt. Veelal kunnen ze alleen gebruikt worden voor de locatie waar de gegevens vandaan komen en geven ze geen inzicht in relaties tussen soorten en de gevolgen van adaptief gedrag. Ook geven habitatmodellen geen inzicht in accumulatie van effecten, synergie of neutralisatie van effecten. Habitatmodellen zijn dus vooral beschrijvend en bevatten geen of nauwelijks informatie over processen. Habitatmodellen zijn wel vaak de enige beschikbare modellen en bieden in ieder geval een kader om eventuele gevolgen van maatregelen te bestuderen. Dynamische modellen zijn relatief arbeidsintensief om te maken, maar leveren wel meer inzicht in sturende variabelen en hebben een groter voorspellend vermogen. Er bestaat een discrepantie tussen gekoppelde fysische modellen en ecologische modellen waarbij de laatste meestal een stuk minder geavanceerd zijn.

A handwritten signature in black ink, consisting of stylized, overlapping loops and a long horizontal stroke extending to the right, positioned above a horizontal line.

Handtekening:

Datum:

21 december 2006