



Trinity College Dublin
Coláiste na Tríonóide, Baile Átha Cliath
The University of Dublin

School of Computer Science & Statistics

A thesis submitted in partial fulfilment of the requirements
for Doctor of Philosophy in Statistics

Design and Analysis of Biodiversity Experiments

Laura Byrne

Supervisor:

Prof. Caroline Brophy

Head of School:

Prof. Gregory O'Hare

**Trinity College Dublin, Ireland,
2025**

Declaration

I declare that this thesis has not been submitted as an exercise for a degree at this or any other university and it is entirely my own work. I agree to deposit this thesis in the University's open access institutional repository or allow the Library to do so on my behalf, subject to Irish Copyright Legislation and Trinity College Library conditions of use and acknowledgement. I consent to the examiner retaining a copy of the thesis beyond the examining period, should they so wish (EU GDPR May 2018).

A handwritten signature in black ink that reads "Laura Byrne". The script is cursive and fluid, with the first letters of "Laura" and "Byrne" being capitalized and prominent.

Laura Byrne, author

Acknowledgements

I would like to use this section to thank the numerous people who have helped me throughout my research. Firstly, my supervisor, Prof. Caroline Brophy, has been of constant support throughout this journey, always going above and beyond to keep me on track, I couldn't have done any of this without her. In the same way, thank you to everyone in the School of Computer Science and Statistics and our ever-growing research group, but special thanks need to go to Rishabh Vishwakarma, who was on this journey with me from the start.

Thank you to John Finn for supporting me from the beginning, from providing data to organising collaborations and reviewing my writing, your help has been invaluable. Thank you to Forest Isbell and Katherine Yurkonis for hosting me on my trips to America, the provision of data, and for teaching me all about the wonderful world of ecology. Thank you to Rafael de Andrade Moral for continuing to teach me R long after you finished being my lecturer. I would like to thank Trinity College Dublin and the Science Foundation of Ireland for funding my PhD and all included conferences and research trips.

I will be forever grateful to my family, partner, and friends for keeping me sane and making sure I didn't run myself into the ground. Special thanks to my mother for always pushing me, even when I complained, believing in me and my ability, even when I had doubts, and for stocking the house with my favourite snacks and drinks when the going got tough. Thank you to the studio behind my favourite show for releasing its next season this summer, and for airing it on the weekends so that I could fully take my mind off of work for 30 minutes each week.

Lastly, to my dear friend Frank Clancy, who inspired me to pursue academia but unfortunately didn't make it to see this submission, I hope I've made you proud.

Table of contents

Declaration.....	i
Acknowledgements.....	ii
Table of contents	iii
Abstract.....	vi
Publications.....	viii
List of abbreviations	ix
1 Introduction.....	1
1.1 Background.....	1
1.2 Designing BEF experiments.....	5
1.3 Analysing BEF experiments.....	10
1.4 Development of Diversity-Interactions modelling.....	12
1.5 Research questions	17
2 Diversity in community designs may improve the estimation of complex interaction structures: A simulation study	19
2.1 Introduction	20
2.2 Materials and methods.....	23
2.3 Results	30
2.4 Conclusions	35
2.5 Discussion and future work.....	38
3 A Diversity-Interactions model for the fitting of multivariate repeated measures data	41

3.1	Introduction.....	42
3.2	Multivariate and repeated measures DI models.....	44
3.3	DImodelsMulti R package.....	49
3.4	A case study	55
3.5	Discussion.....	66
4	Diversity-Interactions models for modelling the total community response and breakdown by species proportions in biodiversity experiments	67
4.1	Introduction.....	68
4.2	Proposed model	71
4.3	Case study.....	77
4.4	Additive partitioning method.....	86
4.5	Discussion.....	90
5	Discussion	95
6	References	105
7	Appendices.....	118
A	The full table of results for each experimental design used in the simulation study	118
B	The R code used in the simulation study for testing experimental design principles	128
C	The linear transformation applied to the case study ecosystem functions	134
D	The results of each step in the multivariate repeated measures case study model selection process	135

E	Comparison between final multivariate repeated measures and ‘equivalent’ univariate repeated measures models for chapter 3 case study	137
F	The code used for the multivariate repeated measures case study	141
G	The results of each step in the realised proportions case study model selection process	147
H	The estimated optimal communities for two metrics in the realised proportions case study.....	148
I	Ternary diagrams showcasing biomass predictions from the realised proportions case study.....	149
J	Stacked barchart showcasing observed vs. predicted biomass from the realised proportions case study	151
K	The code used for the realised proportions case study	153
L	The code used for the ternary diagrams in Appendix I using the realised proportions case study model	182
M	The code used for the comparative plots in Appendix J using the realised proportions case study model	187
N	The code used for the additive partitioning method using the realised proportions case study model.....	192

Abstract

Biodiversity and ecosystem functioning (BEF) relationships define the ways in which the diversity of species in an ecosystem drive the quantity and quality of the goods and services provided. Species diversity may be defined in various ways, such as richness (species number) and/or evenness (relative proportions of species), and may or may not include species identity. Designing BEF experiments, where species diversity is manipulated across experimental units, can present unique challenges given that species diversity has multiple dimensions. In this thesis, using a simulation study, I test the best practices for designing a BEF experiment, with a focus on capturing the nature of how species interact within BEF relationships.

The provision of simultaneous functions by an ecosystem is referred to as multifunctionality. In multifunctionality BEF studies which span across multiple time points, each experimental unit has a multivariate repeated measures response; thus the data that arises from each experimental unit potentially has multiple sources of correlation which must be accounted for in analysis of the data. The development of the Diversity-Interactions (DI) modelling framework builds upon the traditional usage of ‘richness-only’ models in the BEF literature; richness-only models equate species diversity to solely the number of species, which reduces the true dimensionality of species diversity and therefore may cause loss of information and confounding within an analysis. DI modelling is a regression-based modelling approach to quantifying the BEF relationship which has the ability to separate the effects of individual species and their interactions on ecosystem functioning and is highly generalisable. This approach has previously been extended to model multivariate data. In this thesis, I extend DI models for use with BEF data containing responses with multiple sources of correlation (multivariate repeated measures data). I develop new statistical software tools to fit the new methodologies that are user-friendly and aim to promote uptake

of advanced statistical modelling techniques for multifunctionality BEF research. Finally, I develop a new modelling approach for jointly modelling continuous and proportional multivariate response measurements; this work is motivated by challenges identified in the BEF literature but has wider applicability.

Keywords: BEF relationship, Diversity-Interactions models, regression modelling, species diversity, compositional data, longitudinal data, multifunctionality, experimental design, simulation study.

Publications

Peer reviewed journal papers

Vishwakarma, R., **Byrne, L.**, Connolly, J., de Andrade Moral, R. and Brophy, C., 2023.

Estimation of the non-linear parameter in Generalised Diversity-Interactions models is unaffected by change in structure of the interaction terms. *Environmental and Ecological Statistics*, 30(3), pp.555-574.

Moral, R.A., Vishwakarma, R., Connolly, J., **Byrne, L.**, Hurley, C., Finn, J.A. and Brophy, C., 2023.

Going beyond richness: Modelling the BEF relationship using species identity, evenness, richness and species interactions via the DImodels R package. *Methods in Ecology and Evolution*, 14(9), pp.2250-2258.

CRAN published R packages

Byrne, L., Vishwakarma R., Moral, R., and Brophy, C., 2024. DImodelsMulti: Fit

Multivariate Diversity-Interactions Models with Repeated Measures. *The Comprehensive R Archive Network*, R package version 1.1.1, <https://cran.r-project.org/package=DImodelsMulti>.

Vishwakarma, R., Brophy, C., **Byrne, L.**, Hurley, C., 2024. DImodelsVis: Visualising

and Interpreting Statistical Models Fit to Compositional Data. *The Comprehensive R Archive Network*, R package version 1.0, <https://cran.r-project.org/package=DImodelsVis>.

List of abbreviations

Important abbreviations that have been used in the text are briefly explained.

BEF	Biodiversity and ecosystem function(ing)
DI	Diversity-Interactions
STR	Structural model (a form of a DI model)
ID	Identity effects model (a form of a DI model)
AV	Average pairwise interaction model (a form of a DI model)
ADD	Additive species effects model (a form of a DI model)
FG	Functional groups; or Functional group model (a form of a DI model)
FULL	Full pairwise interactions model (a form of a DI model)
UN	Unstructured/general (a type of covariance structure)
CS	Compound symmetry (a type of covariance structure)
AR(1)	Autoregressive model of order one (a type of covariance structure)
ML	Maximum likelihood (estimation method)
REML	Restricted maximum likelihood (estimation method)
MLE	Maximum likelihood estimation (estimation method)

1 Introduction

1.1 Background

The biodiversity and ecosystem function (BEF) relationship refers to the idea or theory that ecosystem functioning, i.e. the capability of an ecosystem's processes to provide useful goods or services (de Groot, 1992; Costanza *et al.*, 1997), depends directly on the levels of diversity within an ecosystem. The study of this relationship has rapidly expanded in recent decades and includes both observational and experimental study designs. In BEF experiments, species diversity is often manipulated between experimental units with some ecosystem function response, e.g. yield or soil water content, measured sometime after establishment (allowing for growing or development time, depending on the type of ecosystem being studied). The ecosystem function responses may be any data type, e.g. continuous, ordinal, count, etc., but are most commonly continuous. Species diversity in a community is often defined using richness, which is the number of unique species present (usually at establishment). Although this is a popular method of quantifying species diversity in a local scale community, it is also the simplest, often resulting in a loss of information in species dimensionality in the form of individual species identities or their proportional representation in the community, and when used within an analysis, may assume that all species contribute to ecosystem functioning in the same way. Other methods of defining species diversity tend to suffer from this same issue, for example, Simpson's index, Shannon's index, and evenness (Shannon, 1948; Simpson, 1949; Menhinick, 1964; Kirwan *et al.*, 2009; Roberts, 2019; Moore, 2023). As such, analyses of BEF relationship study data have been steadily moving away from their reliance on species richness, instead finding that individual species, functional redundancies (wherein multiple species perform the same function in an ecosystem, these redundancies are therefore thought to occur within functional groups of species), and other measures, such as the relative abundances of species employed

by the Diversity-Interactions (DI) modelling method, can more accurately capture the BEF relationship (Whittaker, 1972; Wittebolle *et al.*, 2009; Kirwan *et al.*, 2009; Stuart-Smith *et al.*, 2013; Moral *et al.*, 2023; Finn *et al.*, 2024).

The definition of species diversity utilised by an experimental BEF relationship study will affect the design of said experiment, changing the degree to which species community compositions, both in the identity of the species and their relative abundances, may need to vary in order to detect differences between species mixing effects and prevent confounding variables. Studies utilising richness need only vary the number of unique species between experimental units, disregarding their identities and relative abundance. On the other hand, an experiment which seeks to utilise the DI modelling approach for the analysis of its data may benefit from the inclusion of each species in the pool in a number of experimental units at different proportions, ensuring sufficient coverage of the simplex space (a sample space comprised of compositional data) created by the species proportions.

Many differing habitats have been studied in the BEF literature, from managed grasslands, such as those seen commonly in agriculture, wherein one might measure the influence of an increase in the presence of legumes (therefore a decrease in the relative presence of grasses) on the water content in the soil or the yield (Finn *et al.*, 2013; Binder *et al.*, 2018; Grange *et al.*, 2021), to vast areas of oceans, where the diversity of marine life may influence the biomass produced (Strong *et al.*, 2015; Maureaud *et al.*, 2019). In many ecosystems, biodiversity is thought to drive ecosystem functioning due to known symbiotic relationships between species. For example, in a grassland ecosystem, legumes can provide an extra nitrogen source from the atmosphere, through their nitrogen fixing capabilities, which directly benefits any grass species present. This grass-legume relationship in particular has such been found to be a desirable alternative to the application of excess artificial nitrogen fertiliser (Grange *et al.*, 2021). This finding raises questions about the validity of replacing potentially harmful practices, such as the excessive application of

nitrogen fertiliser, pesticide, and herbicide treatments, with the inclusion of additional beneficial interspecific interactions. However, antagonistic relationships may also exist between species in relation to ecosystem functioning, for example, in the case that two grass species are established together in a field and vie for the same resources, e.g. nutrients in the soil or root space. Following this grassland ecosystem example, a BEF experiment interested in these grass and legume relationships may include plots with varying relative abundances for a number of species within the two functional groups (grasses and legumes) and some ecosystem function response measured, see Table 1.1.

Table 1.1: A table showing an example dataset from a BEF experiment interested in examining the effect of interactions between and within groups of grasses and legumes on some ecosystem function. The four species (two grasses, two legumes) are shown to each be sown in monoculture (at a proportion of 1), along with some mixtures (Richness > 1), both equi-proportionally (plots 5, 6, 9) and with a dominating species (plots 7, 8). Two mock ecosystem function responses are also included, where Function1 is count data, e.g. resulting species richness or aphid count, and Function2 is continuous, e.g. biomass. Generally, this design would be replicated some number of times, e.g. three replicates per community design for a total of 27 plots.

Plot	Richness	Grass1	Grass2	Legume1	Legume2	Function1	Function2
1	1	1	0	0	0	5	200
2	1	0	1	0	0	3	150
3	1	0	0	1	0	8	125
4	1	0	0	0	1	12	100
5	2	0.5	0.5	0	0	5	130
6	2	0	0	0.5	0.5	11	110
7	3	0.25	0.25	0.5	0	4	150
8	3	0.5	0	0.25	0.25	5	250
9	4	0.25	0.25	0.25	0.25	4	300

The simultaneous provision and maintenance of multiple ecosystem functions is referred to as ecosystem multifunctionality (Maestre *et al.*, 2012; Byrnes *et al.*, 2014). While the majority of the BEF relationship literature focuses on the analysis of a single ecosystem function, this could cause the effects of changes to species diversity on an ecosystem's overall functioning to be underestimated. This is especially true if readings are only taken at a singular time point, whereas changes in the observed function being caused by unseen changes in the other functions of the ecosystem may be picked up, e.g., a community of plants may perform well for the observed function of yield but may devastate the unobserved water content of the soil, eventually leading to the drying of plant roots and other domino effects. Multifunctionality has been a topic of interest for a number of years, with an extensive library of literature to match, however, many multifunctionality BEF studies analyse each ecosystem function response in isolation, excluding the inherent covariance between readings from the same experimental/observational units. More recent studies have begun to investigate multiple ecosystem functions simultaneously, with results indicating that higher measures of biodiversity are required to increase multifunctionality (Emmett Duffy *et al.*, 2003; Hector & Bagchi, 2007; Dooley *et al.*, 2015; Suter *et al.*, 2021), likely due to the contributions of different species being more beneficial to certain ecosystem functions. This makes it of great interest to study a multitude of processes in tandem through multivariate methods (Hector & Bagchi, 2007; Gamfeldt *et al.*, 2008).

The longevity and stability of individual species effects, species interactions, and overall ecosystem functioning are also research topics of interest within BEF studies as many ecosystems are intended to be long-standing, e.g. grasslands intended for grazing, reconstructed/restored ecosystems, or fields intended to be systematically resown with different species, e.g. following crop-rotation (Hunter Jr, 2002; Wang *et al.*, 2019). These studies usually tackle such questions in one of two ways, (1) analysing the change in species diversity/composition (Sebastià *et al.*, 2008; Lloret *et al.*, 2009; Liu *et al.*, 2018) or (2)

analysing the legacy effect of species, management, or environmental effects (Wilsey & Polley, 2003; Cuddington, 2011; Kaisermann *et al.*, 2017; Perring *et al.*, 2018; Heinen *et al.*, 2018; Heinen *et al.*, 2020). Over time, ecosystem functioning changes and this may be due to a number of reasons, including, but not limited to, changes in the environment, such as temperature, sunlight, or water availability, changes to species compositions, e.g. the loss of species or the invasion of new species, or due to the fading/lessening of legacy effects. These changes are generally tracked in longitudinal studies, where multiple readings are taken from each subject in the study, usually plots in BEF experiments, at various different time points.

1.2 Designing BEF experiments

Designing fit-for-purpose experiments can regularly present many challenges, from which BEF relationship studies are not exempt. An experiment's design is ideally catered to the study's specific research question(s) and therefore aims to achieve unbiased estimators with low variance for any model produced using the data gathered. In the case that a study instead seeks to test a hypothesis, the experimental setup must provide enough power to detect any differences between units (Bailey, 2008). Constraints such as cost must also be considered when designing an experiment, so the ideal design would be small and cheap, but still provide enough power to draw clear conclusions. However, utilising too small of a BEF relationship experiment runs the risk of not fully capturing the true variability of the ecosystem (and wasting available resources), therefore care must be taken when generalising results, especially in the case of small-scale or localised studies. As such, the number of experimental units made available to an experiment should be carefully considered.

One popular method of lowering variance in an analysis through the experimental design is replication, which involves the repetition of a treatment, such as a particular species composition, on multiple experimental units. This method increases the number of units in a study, adding more power, but may occasionally raise the variance through the introduction

of more variable experimental units. It also increases cost and manual labour involved in the experiment (Bailey, 2008). Given that space granted in a study may also be limited, replication may require that less experimental unit (species community) designs are included in the study. In BEF experiments, this raises the question as to whether replication of species community designs, either exactly or with minute changes, are to be preferred over the inclusion of more unique species communities, see Figure 1.1.

Different BEF relationship studies are designed to manipulate diversity according to some chosen definitions, e.g. richness (number of species) or evenness (measure of relative abundance of all species), and should adjust community designs to reflect their chosen metric of interest. For instance, if species richness is of interest, then there is no requirement to fluctuate species' relative abundances within richness levels, outside of present/absent, while if dominant or rare species are of interest then these different species combinations will need to be represented. In the case that interactions between different species are of interest to the study and intend to be included in an analysis of the data, the species will need to be represented together in the same community a number of times, preferably at different relative proportions. Individual species themselves may not be of interest to the study, but rather their functional traits, in which case the species may be divided into functional groups by characteristics of interest, e.g., root depth or establishment speed, and community designs may be chosen based on these groups as opposed to the species (Hector *et al.*, 1999; Tilman, 2001). The simplex space is often utilised to conceptualise the coverage of the species communities included in a BEF experiment. As such, a species community's design may be defined not just by richness, but also by the identities and relative abundances (proportions) of the species included in each design point. There are different approaches possible when it comes to systematically varying species proportions for an experiment; for example, all units can be heavily replicated at only equi-proportional combinations (e.g., Fig. 1.1, Design 1), or some unbalanced and equi-proportional communities can be included, each moderately

replicated (Fig 1.1, Design 3), or many points may be selected across the simplex space, with just one replicate at each point (Figure 1.1, Design 2). This coverage defines the information awarded by the experiment's design for the prediction of an ecosystem function response from a community within the simplex space.

Many methods exist for the systematic variation of treatments within an effectiveness relationship study, the “gold standard” of which is considered to be randomised controlled trials (Tarnow-Mordi *et al.*, 2017; Hariton & Locascio, 2018). In BEF relationship studies, this refers to the random selection of species from a pool, for a number of different richness levels, which are then typically sown in equi-proportional communities (Tilman *et al.*, 1996). This method is thought to remove bias through the randomisation of the treatment, i.e. the species community designs, applied to each experimental unit, however, as richness levels increase, there is an increased chance that a particularly high-performing species may be included in the community, boosting ecosystem functioning through its identity rather than the increased level of richness. This phenomenon is referred to as the selection probability effect (Huston, 1997). This effect may cause confounding within an analysis of BEF data which uses species richness as a predictor due to the number of species within a community and its ecosystem functioning performance becoming correlated. As such, in smaller-scale studies utilising a manageable number of species, it is often preferred to manually assign species to experimental units, ensuring the species identities within their compositions vary within richness levels. However, for experiments which contain a large pool of species, for which it would be effectively impossible to include each in differing species compositions for every richness level due to spatial constraints, the random selection of species may be required to prevent bias. In this case, the inclusion of species identities (or functional groupings) as predictors in the resulting analysis, such as through the usage of DI models (Kirwan *et al.*, 2009; Connolly *et al.*, 2013), and the variation of species proportions from

the equi-proportional standard, the selection probability effect can be avoided as high-performing species can be identified, with their effect quantified.

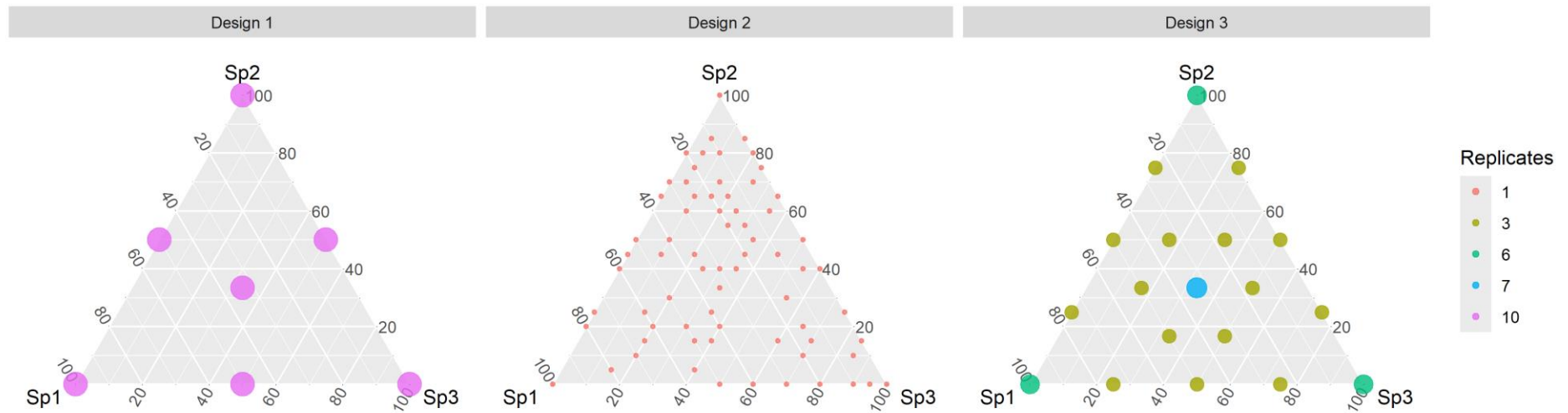


Figure 1.1: Ternary diagrams showing the spread of species proportions across the simplex space in three example datasets of equal size (70 units), for a three-species BEF experiment. Design 1 features solely equi-proportional communities, each with ten replicates, which creates a very sparse design with lots of information on particular points but any design points in between requiring interpolation. Design 2 features the antithesis of Design 1, with many points spread across the simplex, granting information for many different species community designs but with no replication, reducing confidence in the results from each point. Design 3 is presented as an example of a middle-ground alternative to the two previous extreme designs, wherein both equi-proportional and non-equi-proportional (uneven) designs are included and are replicated some number of times each, with communities of particular interest (in this case the monocultures and three-species equi-proportional mixture) replicated additional times.

1.3 Analysing BEF experiments

For univariate data from a BEF experiment, that is data with a single ecosystem function response recorded at a single time point from each experimental/observational unit, many studies utilise regression, specifically ANOVA, modelling to test the relationship between the response and some chosen measure(s) of diversity (Hooper *et al.*, 2005). For ANOVAs, the community design, represented as a categorical variable, is generally used as a predictor, along with any additional treatments, such as additional nitrogen fertiliser or simulated drought (Daepf *et al.*, 2001). This method does not provide any information on the effects of individual species within a community, simply defining the community as belonging to a particular category, and allows for no extrapolation to communities that were not included, in the original study, meaning that each community of interest must therefore be included potentially increasing the cost and labour of any experimental study which plans to utilise this method in its analysis. Other measures of species diversity may be used with linear regression models, e.g. species richness or evenness. An example of a regression model relating some ecosystem function response y to the evenness of species in some experimental unit m can be seen in Eq. 1.1, where S is the total number of species included in the study and p_i is the initial relative abundance, or proportion, of some species i .

$$y_m = \frac{2S}{S-1} \sum_{\substack{i,j=1 \\ i < j}}^S p_{im} p_{jm} + \varepsilon_m, \quad \varepsilon_m \sim N(0, \sigma^2) \quad (\text{Eq. 1.1})$$

Species diversity may also be separated into identity (ID) effects and interaction effects (Loreau, 1998) which allows for the definition of the interactions between species, such as resource partitioning (where species evolve to use resources differently to enable coexistence with minimal to no competition) and facilitation (where a species' presence may benefit another directly, e.g. nurse plants protecting germinating seeds). The additive partitioning method (Loreau & Hector, 2001; Gering *et al.*, 2003) is one such approach,

estimating an overall interaction effect of the species present in a community on the measured ecosystem function. This is referred to as the complementarity effect. However, it is difficult to design experiments for this method due to very strict data requirements, needing a quantification of the contribution of each species to the ecosystem function, even in large mixtures where it may even be physically impossible to do so. It also does not account for negative or insignificant species interactions.

Multiple measurements taken from the same experimental unit, e.g. multiple ecosystem function responses (multivariate) or readings taken at multiple time points (repeated measures), have inherent correlations between them. When modelling data with these response types, ignoring these correlations will result in pseudo-replication, affecting test results and possibly leading to incorrect inferences. One could model each response separately (Allan *et al.*, 2013), e.g. separate time points '1' and '2' from data with a single ecosystem function response and fit two univariate models, but, while this approach is valid, in the absence of quantifying the covariances between ecosystem functions and/or time points, it is not possible to test relative comparisons across time or functions.

Other current methodologies include summarising the effect of the multiple response types on some chosen ecosystem functioning metric, e.g., by standardising and averaging the responses to form an 'average multifunctionality index' (Hooper & Vitousek, 1998; Mouillot *et al.*, 2011), applying the turnover approach (Hector & Bagchi, 2007; Isbell *et al.*, 2011), or using the threshold approach (Gamfeldt *et al.*, 2008). While these methods are valid and often useful for investigating generalised research questions over relatively short time periods, they reduce the dimensionality of the responses and may also remove the ability to detect trade-offs between the different response types in addition to the aforementioned loss of testing capabilities.

1.4 Development of Diversity-Interactions modelling

One method of modelling the BEF relationship is the DI modelling framework (Kirwan *et al.*, 2009). It is a regression-based approach that assumes species' identities and their proportions to be the driving force behind changes in ecosystem functioning, meaning that the species proportions are included as predictors in the models, and they define a simplex space (where the species proportions sum to one). DI models can capture the effects of individual species, their interactions, evenness, and richness, moving beyond the reliance on species richness as a primary measure of species diversity (Kirwan *et al.*, 2009; Connolly *et al.*, 2013; Moral *et al.*, 2023). An example DI model is given in Equation 1.2,

$$y_m = \sum_{i=1}^S \beta_i p_{im} + \sum_{\substack{i,j=1 \\ i < j}}^S \delta_{ij} (p_{im} p_{jm})^\theta + \alpha \mathbf{A}_m + \varepsilon_m, \quad \text{where } \varepsilon_m \sim N(0, \sigma^2) \text{ (Eq. 1.2)}$$

where y_m is a single response variable from an experimental/observational unit m ($m = 1, 2, \dots, M$). S is the number of species included in the study and p_{im} represents the initial proportion of the i^{th} species from the m^{th} experimental/observational unit. The parameter β_i scaled by p_i is the expected contribution of the i^{th} species to the response and is referred to as the i^{th} species' identity (ID) effect; if no species interactions or treatments are required in the model, the expected y is the weighted average of the species ID effects. $\mathbf{A}_m$ is a vector which may include blocks or treatment terms (that are in addition to the species diversity manipulation) and α is a vector of the corresponding effect coefficients, as such these terms may be excluded from the model if there are no additional structures in the study. The full pairwise (FULL) interaction structure is shown here as an example, but DI interactions are highly customisable. The inclusion of the non-linear parameter θ (Connolly *et al.*, 2013) allows the models to deviate from the assumption that interactions are directly proportional to the product of each species' proportions, thus allowing the shape of the BEF relationship to change (Moral *et al.*, 2023; Vishwakarma *et al.*, 2023; Finn *et al.*, 2024). It is applied as

a power to each pairwise product of species proportions and is usually bound to the range $[0, 1.5]$.

These models can be easily fit using least squares estimation, as with any linear regression model (Aitken, 1936). The usage of individual species effects enables the investigation of the BEF relationship across communities that were not necessarily observed, given that enough information is provided in the space surrounding the chosen point in the simplex (Kirwan *et al.*, 2009). The non-linear parameter θ can be estimated using profile log-likelihood optimisation (Brent, 1973).

The number of parameters in Eq. 1.2 can be reduced by imposing constraints on species interaction terms (Kirwan *et al.*, 2009; Moral *et al.*, 2023); similarly, the model can be extended to include further interactions involving the variables in A . These proportion-based interaction effects include, but are not limited to, those shown in Table 1.2. The presented interaction structures are hierarchical in nature and can therefore be tested against one another using F-tests or metrics such as Akaike’s Information Criteria (AIC) (Kirwan *et al.*, 2009).

Univariate DI models can be easily fit using the `DImodels` R package (Moral *et al.*, 2023), which is available on CRAN. It has a selection of pre-programmed interaction structures available, an option for the inclusion and estimation of the non-linear parameter θ , and an automatic model selection procedure, making the application the DI modelling framework to any univariate single-site data a straightforward process. It also includes many S3 methods and customisable features, enabling deeper analyses after the initial model fitting.

Table 1.2: Examples of appropriate interaction structures for use with DI models, their name, algebraic expression, and general interpretation.

Structure Name	Algebraic Expression	Interpretation
Average Pairwise (AV)	$\delta_{AV} \sum_{\substack{i,j=1 \\ i < j}}^S (p_i p_j)^\theta$	All pairwise interactions are equal and defined by a single interaction term
Additive (ADD)	$\sum_{\substack{i,j=1 \\ i < j}}^S (\delta_i + \delta_j)(p_i p_j)^\theta$	The contribution of a species i to a pairwise interaction is constant, regardless of the species it is interacting with
Functional Group (FG)	$\sum_{q=1}^C \delta_{qq} \sum_{\substack{i,j \in FG_q \\ i < j}} (p_i p_j)^\theta + \sum_{1 \leq q < r \leq C} \delta_{qr} \sum_{i \in FG_q} \sum_{j \in FG_r} (p_i p_j)^\theta$	Species are grouped into C functional groups by similar traits and interact within and between these groups
Full Pairwise (FULL)	$\sum_{\substack{i,j=1 \\ i < j}}^S \delta_{ij} (p_i p_j)^\theta$	Each pair of species is assumed to interact uniquely

While these previous examples are for normally distributed continuous responses, with no within- or between-subject variance, the DI models framework has also been adapted to fit data which include a number of different complexities. This includes a multinomial baseline category logit model (Agresti & Min, 2002) in order to assess multinomial ecosystem functions (Brophy *et al.*, 2011), the inclusion of random effects to the framework (Brophy *et al.*, 2017), which allows for the addition of variability within summary metrics for interaction effects, such as the average or evenness terms, for each pair of species, and the use of phylogenetic distances between species in a community to explain the differences in ecosystem functioning between communities of different designs (Connolly *et al.*, 2011). This has been possible due to the generalisable nature of the framework.

When multiple responses are read from the same experimental/observational unit in a study, they are not independent, in other words, an inherent correlation will exist between them, for example, a high-performing study subject will likely outperform a low-performing subject, despite both being subjected to the same tests. To account for covariance between measurements of a single ecosystem function response taken from some unit m , where $m = 1, 2, \dots, M$, at multiple time points, the DI models framework from Equation 1.1 was adapted to allow for the analysis of these repeated measures in BEF data (Frankow-Lindberg *et al.*, 2009; Finn *et al.*, 2013). Equation 1.3 is an example formulation of the repeated measures DI model, where the subscript t represents the different time points from which an ecosystem function response was measured, $t = 1, 2, \dots, T$, and Σ^* is an $MT \times MT$ block diagonal matrix, with one $T \times T$ block, Σ , for each experimental unit m , where Σ_m is the variance-covariance matrix of the repeated measures responses for a subject m (Eq. 1.4).

$$y_{mt} = \sum_{i=1}^S \beta_{it} p_{im} + \sum_{\substack{i,j=1 \\ i < j}}^S \delta_{ijt} (p_{im} p_{jm})^{\theta} + \alpha_t \mathbf{A}_m + \varepsilon_{mt}, \quad \varepsilon \sim MVN(0, \Sigma^*) \quad (\text{Eq. 1.3})$$

$$\Sigma^* = \begin{bmatrix} \Sigma_1 & & 0 \\ & \ddots & \\ 0 & & \Sigma_M \end{bmatrix} \quad (\text{Eq. 1.4})$$

To account for covariance between differing ecosystem functions, the DI models framework from Eq. 1.2 was also adapted to allow for the analysis of multivariate datasets (Dooley *et al.*, 2015). Eq. 1.5 is an example formulation of the multivariate DI model, which follows the same formulation as the repeated measures DI model, where the subscript k represents the different ecosystem functions, $k = 1, 2, \dots, K$, and Σ^* is a $KM \times KM$ block diagonal matrix, with one $K \times K$ block, Σ , for each experimental unit m , where Σ is the variance-covariance matrix of the multivariate responses (Eq. 1.4).

$$y_{km} = \sum_{i=1}^S \beta_{ik} p_{im} + \sum_{\substack{i,j=1 \\ i < j}}^S \delta_{ijk} (p_{im} p_{jm})^{\theta_k} + \alpha_k \mathbf{A}_m + \varepsilon_{km}, \quad \varepsilon \sim MVN(0, \Sigma^*) \quad (\text{Eq. 1.5})$$

As no between-subject variation is included in either of the above models, all Σ_m within a model (Eq. 1.4) are considered to be equal. These models can be fit using generalised least squares methods (Aitken, 1936).

Using the above methods, one can predict from the fitted model for each ecosystem function or time point for all possible species proportions, even if not included in the original design, allowing the identification of conditions where all functions/times perform relatively well or where trade-offs between responses appear. The variance-covariance matrices estimated for the response types, Σ , may also contain useful information for a researcher regarding the relationships between the different response readings. Another strength of these DI models is the ability to compare and statistically test the relative performance of species (and their interactions) across response types (Frankow-Lindberg *et al.*, 2009; Finn *et al.*, 2013; Dooley *et al.*, 2015). These advantages come at the cost of an increase in the complexity (number of terms) of the model, when compared to previously discussed methods. One critique of this model structure is the difficulty of interpretation of the results for a large number of ecosystem functions or time points, however, a summary statistic for these responses can be produced using the response values predicted from the model to simplify interpretation, looking at the overall functioning of the ecosystem across the response types, for example, the mean log response ratio (MLRR) could be used (Suter *et al.*, 2021). This is similar to the threshold approach (Gamfeldt *et al.*, 2008) and averaging method (Hooper & Vitousek, 1998) but, as predicted response values are used to create the summary statistic rather than the raw observed response values, with the added benefits provided by the repeated measures and multivariate DI methods.

1.5 Research questions

This thesis aims to facilitate a better understanding of the design of BEF relationship studies, and to develop modelling and software tools with which to analyse BEF data, particularly those with correlated sets of ecosystem function responses.

For experimental research studies, the experiment must be planned with a certain hypothesis or research question in mind and ensure that the community designs included are sufficient for testing. Given the limited space provided for experimental units in most studies, it is of interest to test which design aspects should be prioritised, e.g. monocultures over mixtures, replicates over novel communities, or the requirement of monocultures at all (especially given that observational studies will likely never include true monocultures). In chapter 2 of this thesis, I showcase a simulated study which tests the ability of different experimental design principles to provide sufficient information around the simplex space to accurately select the true underlying species interaction structure.

While DI models have already been expanded to account for correlation between responses caused by repeated measures (Frankow-Lindberg *et al.*, 2009; Finn *et al.*, 2013) or multiple ecosystem functions (Dooley *et al.*, 2015), the framework currently cannot be used to account for covariance from both sources simultaneously. In chapter 3, I present the multivariate repeated measures DI model, which is capable of fitting DI models with multiple sources of correlation between responses taken from the same experimental or observational unit in a study.

The change in species compositions and relative abundances over time is of great interest, specifically for the facilitation of stable ecosystem functioning in long-standing ecosystems. As the structure of species communities change, the contribution of each species to ecosystem functioning will vary with it, therefore showing a need to analyse both changes in ecosystem functioning and its resulting species community simultaneously to truly form a picture of the effect of time. In chapter 4 of this thesis, I present a regression model which,

when coupled with the DI modelling framework, is capable of modelling the effect of individual species, and their interactions, on both an ecosystem function and the resulting species community composition and relative abundances simultaneously, accounting for the covariance that exists between each response and imposing a sum-to-one constraint on the species proportions for both the predictor and the response side of the model.

Univariate DI models can be easily fit using the `DImodels` R package (Moral *et al.*, 2023). It allows a user to automatically fit these models using a range of built-in interaction structures and allows for the estimation or provision of a θ value. The models produced are highly customisable, however, it cannot be used to fit any of the expansions to the framework, as such, it is not suitable for use with data containing multivariate and/or repeated measures. In chapter 3 of this thesis, I present the R package `DImodelsMulti` which enables the fitting of multivariate, repeated measures, and multivariate repeated measures DI models, following the functionality and layout of the `DImodels` package.

To summarise, the main aims of my PhD are:

1. To explore and formally test the effects of different experimental BEF relationship study designs through a simulation study.
2. To expand the DI modelling framework for usage with data containing different forms of correlated responses.
3. To develop statistical software tools to provide a user-friendly way to fit complex DI models.

In this thesis, chapter 2 addresses aim 1, chapters 3 and 4 address aim 2, and chapter 3 addresses aim 3. Please note that chapters 3 and 4 were written as stand-alone manuscripts for submission to two separate journals, thus there is some variation in their formatting and they contain some duplicate information.

2 Diversity in community designs may improve the estimation of complex interaction structures: A simulation study

Collaborators: Caroline Brophy, Kathryn Yurkonis

CB conceived the study; LB led the coding and writing for the manuscript, with input from both collaborators; KY provided real-world data and ecological expertise.

Abstract

Biodiversity and ecosystem function (BEF) relationship studies seek to quantify the relationship between species diversity and the goods and services that an ecosystem provides. The effectiveness or value of a BEF experiment greatly depends on its design's ability to provide sufficient power to address the research questions or test the hypotheses of interest. Given that most experiments have limited resources available (money, labour, and/or space), which usually imposes a limit on the number of experimental units possible, it can be necessary to compromise on some aspects of the design of a given experiment. In BEF experiments, species diversity is generally systematically varied across experimental units, with one or more ecosystem functions recorded sometime after establishment. This variation can involve changes in both species' richness (number) and their relative abundances (proportions), which form a simplex space. Ideally, the species communities included in an experimental design will cover much of the simplex space, but at minimum should have sufficient coverage around areas of interest.

In this chapter, a simulation study is used to test the ability of different experimental design principles to successfully select an underlying true species interaction structure from

a set of Diversity-Interactions models. I found evidence indicative that the inclusion of varied plot-level species community designs, spread well around the simplex space and each with one or a small number of replicates, is more beneficial than a larger number of replicates of monocultures or equi-proportional mixtures, which instead focus on narrower regions within the simplex space.

2.1 Introduction

The study of the biodiversity and ecosystem function (BEF) relationship aims to understand how changes to ecosystem functioning are influenced by species diversity. The good design of BEF relationship experiments is at the core of the study of these relationships, determining the type of statistical analysis methods available and the generalisability of the study. For instance, testing a hypothesis relies on the experiment setup providing enough power for detecting differences between units (Bailey, 2008) and smaller experiments run the risk of not fully capturing the true variability of the effects of species diversity on ecosystem functioning, meaning that conclusions on differences between experimental unit designs made may not be applicable to, for example, species communities not included in the initial design (Bruehlheide *et al.*, 2013; Klaus *et al.*, 2020). Constraints such as cost and labour must also be considered when designing an experiment, with the ideal design being logistically possible but, importantly, that also has sufficient power to detect meaningful treatment effects.

In the case of a posed research question, where statistical models will be used to estimate the effect of some predictor on an ecosystem function response, the goal of the design is to provide unbiased estimators with low variance. One popular method of lowering variance in an analysis through the experimental design is replication, which involves the repetition of a species diversity (or other) treatment, on multiple experimental units. This method increases the number of units in a study, improving estimation of the ecosystem

function response for the given community design, but may occasionally raise the variance through the introduction of more variable experimental units (between subject variation), e.g. differences in shade cover of the plots, it may also increase cost and the labour involved in the experiment, through sowing, maintenance, and the collection of observed responses (Bailey, 2008). Given that space granted in a study may also be limited, replication may prevent the inclusion of additional unique experimental unit designs in the study. When working with simplex data, i.e., where the treatment variables include a set of proportions that sum to one for each experimental unit (e.g., the manipulated species proportions in a BEF experiment), it is important to consider the coverage of the experimental design across the simplex space. It is of particular importance to ensure that areas of interest, e.g., single species (monoculture) communities, have sufficient coverage. This raises the question as to whether replication of specific communities should be preferred over the inclusion of more unique species communities, i.e., whether it is preferred to spread data points across the simplex space or to cluster points around areas of interest, see Figure 1.1.

While all BEF relationship studies seek to investigate the connection between species diversity and ecosystem functioning, each study may define diversity differently and therefore should adjust the community designs included in their experiment to reflect changes in the chosen metric of interest. If species richness is the main focus of the study, for example, then there is no requirement to fluctuate species' relative abundances within richness levels, outside of present/absent (Bell *et al.*, 2005; Weigelt *et al.*, 2010; Vogel *et al.*, 2019; Jochum *et al.*, 2020). A suitable design in this case could then be to randomly select species richness levels for each experimental unit and sow a random selection of species, at the selected richness level, equi-proportionally in each, similarly to a randomised controlled trial, which is considered the 'gold standard' for effectiveness trials (Tarnow-Mordi *et al.*, 2017; Hariton & Locascio, 2018). The randomness in these designs reduces bias in species selection and community designs, allowing for a more confident

generalisation of results to species communities not included in the original experimental design. On the other hand, if the relative rarity of a species is of interest, then different species proportion combinations other than equi-proportional will likely need to be represented. When richness levels above one are utilised in a study, species interaction effects may be of interest to include in analyses. In this case, the species should be sown together a number of times, preferably at varying relative proportions.

Diversity-Interactions (DI) modelling is a regression-based approach which can be used to analyse BEF relationship (Kirwan *et al.*, 2009; Connolly *et al.*, 2013; Moral *et al.*, 2023). This approach can simultaneously capture the effect of species richness (number), evenness (how the proportions of species are distributed), and composition (species identities) on ecosystem functioning; it does this through including species proportions and their interactions as predictors in the model (see chapter 1 of this thesis; Finn *et al.*, 2024). The specification of the interactions can vary, for example, it may be assumed that all pairs of species interact equally (average pairwise model) or that each pair of species interact in a unique way (full pairwise model), see Table 1.2 for further examples. These models assuming different interaction formats can be formally tested against one another for a given dataset. If, for example, the full pairwise model is the true underlying model, it may be that the design of the experiment could affect the power of the DI modelling framework to detect that as the best model.

When using DI models for analysis, the inclusion of plots with a single sown species (monocultures) is important for the estimation of species identity (ID) effects (the coefficients of the species proportions), however, these effects may still be estimated using solely mixtures (such would be the case in the majority of observation studies) but the associated standard error will be large in the absence of monocultures. As such, it is preferred to include monocultures in experimental studies, especially if one intends to predict a

monoculture response (to avoid extrapolation beyond the data range). Increasing the number of monoculture replicates will increase the precision of the estimates of the ID effects.

In this chapter, I test the ability of different experimental design principles to provide sufficient power for the correct detection of some true underlying interaction structure within the DI modelling framework. An assumption is made that there is a limit to the number of experimental units that can be included. The design principles of interest in this study are:

- 1) The preferential inclusion of additional monoculture replicates over mixtures;
- 2) The preferential inclusion of even (equi-proportional) mixtures, where all species sown are represented equally, over uneven mixtures, where species range from domination of an experimental unit to being rare.

2.2 Materials and methods

Overview of DI models

DI models include an identity (ID) effect for each species, interaction effects between species, and may include additional experimental treatments (captured under ‘Structure’; Eq. 2.1) (Kirwan *et al.*, 2009; Connolly *et al.*, 2013; Moral *et al.*, 2023).

$$Response = ID + Interactions + Structure + Error \quad (Eq. 2.1)$$

The ID effects of the model, for some ecosystem function response y from an experimental unit m , are generally represented by coefficients β_i , which are scaled by the initial proportion, p , of species i (Eq. 2.2). The interaction effects are usually represented as some sum of the products of pairs of species proportions, with coefficient δ , but are highly customisable, such as via the inclusion of the non-linear parameter θ (Connolly *et al.*, 2013), which allows the shape of the BEF relationship to vary (Moral *et al.*, 2023; Vishwakarma *et al.*, 2023). The example interaction structure shown in Eq. 2.2 is referred to as the average pairwise (AV) structure and assumes that each pair of species interacts in the same way and therefore may be captured through the inclusion of a single predictor term. In the case that

interaction effects are not included, referred to as the identity (ID) model, the coefficient β_i represents the expected average monoculture response for species i and the response y is simply the summation of each species' ID effect, scaled by their relative abundances. Structural effects, such as treatments or block effects, are generally represented through the predictor matrix $\mathbf{A}$. The error term of the model, ε , is assumed to follow a normal distribution, with mean 0 and variance σ^2 . These models, with various built-in interaction structures, can be automatically fit using the `DImodels` R package (Moral *et al.*, 2023).

$$y_m = \sum_{i=1}^S \beta_i p_{im} + \delta_{AV} \sum_{\substack{i,j=1 \\ i < j}}^S (p_{im} p_{jm})^\theta + \boldsymbol{\alpha} \mathbf{A}_m + \varepsilon_m, \quad \varepsilon_m \sim N(0, \sigma^2) \quad (\text{Eq. 2.2})$$

Simulation study design

I designed and ran a simulation study to test the efficacy of a number of different experimental designs for the correct selection of the interaction structure of some true underlying DI model. A number of different experimental design principles were chosen to vary and test within the study. Each combination of these experimental design principles were used to generate 10,000 experiment design matrices, between which the mixture species communities selected were varied, which were then used to generate responses from three true underlying DI models for a total of 3.36 million datasets with simulated responses. The design principles questioned in this study were the number of experimental units utilised, the size of the pool of species in the experiment, the number of monoculture replicates sown for each species, and the proportion of mixture communities which were sown equi-proportionally, see Table 2.1.

Table 2.1: A table showing the values of the experimental design principle parameters utilised in this simulation study presented here. Each experimental design simulated included at least one monoculture for each species, along with up to five additional replicates. After allotting experimental units to monocultures, the remaining units were filled with mixture designs of randomly selected richness values and species compositions (IDs). These mixtures were divided into equi-proportional mixtures and uneven mixtures based on the stated proportion. The proportions of species in the uneven mixtures were randomly selected from multiples of 0.10.

	Values
Number of experimental units	50, 100
Number of species	4, 6 (from 2, 3 functional groups respectively)
Number of monoculture replicates	1, 2, 4, 6
Proportion of all mixtures that were equi-proportional at a given level of richness	0, 0.1, 0.3, 0.5, 0.7, 0.9, 1

The 112 available combinations of the selected design variations (2 experimental units by 2 species by 4 number of monoculture replicates by 7 different categories of mixtures) were each used to generate 10,000 experiment designs, see Table 2.2 for an example simulated dataset. First, the number of monocultures required would be dictated by the number of species available in the experiment (four or six) and their number of replicates (one, two, four, or six), e.g. for an experiment with four species and two monoculture replicates each, eight experimental units would be used for monocultures. The total number of experimental units available, combined with the required number of monocultures, would therefore define the number of mixture communities to be included in the design, following the previous example of eight monocultures, for an experiment with 100 experimental units, this left 92 available units for mixtures. These mixture units are separated into equi-proportional and uneven units, the number for each being defined by the proportion of equi-

proportional communities (0.0, 0.1, 0.3, 0.5, 0.7, 0.9, or 1.0) multiplied by the number of remaining experimental units after monoculture assignment (rounded to the nearest whole number).

Table 2.2: An example dataset produced using 50 experimental units, four species, with two monoculture replicates, and where 50% of mixture communities were equi-proportional. The dataset was estimated from a true underlying model, with response y , following the functional group (FG DI) interaction structure, where species p1 and p2 are in one functional group and species p3 and p4 are in another.

p1	p2	p3	p4	y
1.00	0.00	0.00	0.00	1031.1630
1.00	0.00	0.00	0.00	1034.9056
0.00	1.00	0.00	0.00	605.1162
0.00	1.00	0.00	0.00	265.1627
0.00	0.00	1.00	0.00	560.2773
0.00	0.00	1.00	0.00	741.2800
0.00	0.00	0.00	1.00	726.8379
0.00	0.00	0.00	1.00	854.7870
0.00	0.33	0.33	0.33	888.0989
0.50	0.50	0.00	0.00	783.0426
0.50	0.00	0.00	0.50	1101.5521
0.00	0.50	0.00	0.50	956.1744
0.25	0.25	0.25	0.25	1033.5838
0.00	0.50	0.00	0.50	723.1595
0.50	0.50	0.00	0.00	876.0295
0.00	0.00	0.50	0.50	906.5609
0.50	0.00	0.00	0.50	1087.1118
0.25	0.25	0.25	0.25	701.0256
0.00	0.00	0.50	0.50	547.3065
0.33	0.33	0.33	0.00	779.3548

0.50	0.00	0.00	0.50	1176.1998
0.00	0.33	0.33	0.33	785.0651
0.50	0.00	0.50	0.00	966.5395
0.25	0.25	0.25	0.25	782.2250
0.25	0.25	0.25	0.25	957.7204
0.50	0.00	0.50	0.00	1128.3426
0.25	0.25	0.25	0.25	1133.2830
0.25	0.25	0.25	0.25	951.3567
0.00	0.33	0.33	0.33	759.7206
0.00	0.30	0.20	0.50	785.5138
0.20	0.00	0.80	0.00	783.9275
0.00	0.70	0.00	0.30	623.9295
0.60	0.40	0.00	0.00	1022.0786
0.30	0.00	0.00	0.70	1004.4088
0.50	0.20	0.00	0.30	1044.6182
0.30	0.70	0.00	0.00	816.7443
0.10	0.00	0.60	0.30	832.2852
0.40	0.00	0.00	0.60	1093.7350
0.30	0.50	0.00	0.20	968.9810
0.00	0.00	0.10	0.90	670.6440
0.10	0.50	0.00	0.40	967.9004
0.40	0.40	0.00	0.20	712.4931
0.00	0.00	0.60	0.40	647.4948
0.00	0.00	0.20	0.80	665.4705
0.30	0.20	0.10	0.40	990.6886
0.30	0.70	0.00	0.00	896.6805
0.70	0.00	0.20	0.10	1146.6438
0.40	0.00	0.00	0.60	1032.1017
0.40	0.10	0.50	0.00	1015.4691
0.10	0.90	0.00	0.00	431.6106

The assignment of the community designs to each mixture ‘unit’ was based on randomised controlled trials, which involve the random assignments of treatments to experimental units/subjects. For each of the mixture units, the richness of the community is randomly selected between two and the species pool size (four or six) and, when the richness level is lower than the species pool size, the species included in the community are randomly selected from the species pool. Therefore it is possible that some richness levels and some species within richness levels do not appear in some, or potentially (albeit unlikely) all of the generated experiment design matrices. For equi-proportional communities within the same richness level and same selection of species, there is no randomness in the relative abundances of the species (as all are sown equally), meaning that full-design replicates are likely to occur, particularly for the four-species experiments with 100 experimental units and only one monoculture replicate for each species. The relative proportions of the selected species for each non-equi-proportional mixture are assigned via the random generation of a list of non-zero positive numbers (whose size is equal to the randomly selected richness level) which sum to ten and are subsequently divided by ten to provide proportions for each species, which are multiples of 0.10. Similarly to the equi-proportional mixtures, although less likely, replicates may exist of these uneven communities within each generated design, however, care was taken to ensure that the communities were not randomly assigned all equal species proportions.

Three true underlying DI model interaction structures were selected to test for each of the 1.12 million (112 design principles by 10,000 generated design matrices) generated experiment design matrices: the identity (ID) structure, where it is assumed that there are no interactions between species; the average pairwise (AV) interaction structure, where all species are assumed to interact in the same way, as seen in Eq. 2.2; and the functional group (FG) interaction structure, see Eq. 2.3, which assumes that species, which are grouped into

C groups based on functional traits, interact within (δ_{qq} for some group q) and between (δ_{qr} for groups q and r) their functional groups. This gave a total of 3.36 million datasets.

$$y_m = \sum_{i=1}^s \beta_i p_{im} + \sum_{q=1}^c \delta_{qq} \sum_{\substack{i,j \in FG_q \\ i < j}} (p_i p_j)^\theta + \sum_{1 \leq q < r \leq C} \delta_{qr} \sum_{i \in FG_q} \sum_{j \in FG_r} (p_i p_j)^\theta + \varepsilon_m \quad (\text{Eq. 2.3})$$

The selected coefficients for the six (three interaction structures for each of the four and six species pool sizes) true underlying models from which responses were simulated can be seen in Table 2.3, and were not varied within model structures for this study as the species included in the experiments were assumed to be the same. The relation between each coefficient of the true underlying model(s) were selected based on analyses of several real-world datasets and *a priori* knowledge of the functioning of the selected functional groups for a biomass response, e.g., the synergistic relationship between grasses and legumes attributed to the nitrogen-fixing abilities of legumes and the low performance of herbs in monoculture. For the four-species experiment designs, it was assumed that two grasses and two legumes were used, while for the six-species designs it was assumed that the same grasses and legumes were used, with two additional herb species.

Each simulated experiment dataset was then analysed by fitting and comparing a series of DI models from a set of four interaction structures, where the non-linear parameter θ could be estimated to be different to one (ID, AV, FG, and FULL; see Table 1.2 of this thesis, Kirwan *et al.*, 2009, Moral *et al.*, 2023). The `auto_DI()` function from the `DImodels` R package (Moral *et al.*, 2023) was used to automatically fit and compare these models using F-tests and the number of times the true underlying interaction structure (either ID, AV, or FG) was chosen was recorded (Appendix A). The randomisation aspect of this simulation study allows me to test the general design principles, i.e. the proportion of even vs. uneven mixtures, rather than specific experimental designs or datasets.

Table 2.3: A table showing the true underlying coefficients utilised for the three DI model structures tested for this simulation study. The ID effects for each species, the non-linear parameter θ , and the variance of the error were held constant between models.

		IDs (β_i)	Interactions (δ_{ij})	Theta (θ)	Variance (σ^2)
4-species	ID	$\beta_{G1} = 1034.91$		$\theta = 1$	$\sigma^2 = 124$
	AV	$\beta_{G2} = 554.74$	$\delta_{AV} = 750.18$		
	FG	$\beta_{L1} = 654.36$ $\beta_{L2} = 713.83$	$\delta_G = 443.23,$ $\delta_L = -359.80,$ $\delta_{GL} = 827.14,$		
6-species	ID			$\theta = 1$	$\sigma^2 = 124$
	AV	$\beta_{G1} = 1034.91$ $\beta_{G2} = 554.74$	$\delta_{AV} = 750.18$		
	FG	$\beta_{L1} = 654.36$ $\beta_{L2} = 713.83$ $\beta_{H1} = 251.54$ $\beta_{H2} = 421.12$	$\delta_G = 443.23,$ $\delta_L = -359.80,$ $\delta_H = 105.37,$ $\delta_{GL} = 827.14,$ $\delta_{GH} = 1280.20,$ $\delta_{LH} = -292.24$		

2.3 Results

The findings from this simulation study, for each true underlying DI model structure, can be seen below.

ID Model (as the true underlying model)

For the DI model which only includes the ID effects of each species, with no interactions, it was found that the correct underlying DI model was selected by the `auto_DI()` function consistently at a roughly 85% rate, regardless of the underlying experimental design principles, i.e. the number of monoculture replicates and the ratio of mixtures which were

sown equi-proportionally vs those sown with some dominating species (Figure 2.1 (a)). This result was also consistent between designs which contained four or six species and those containing 50 or 100 experimental units. The only occurrences of a noticeably lower rate of correct structure selection occurred for the designs which included a single monoculture plot for each species, with no additional replicates, the lower species pool size of four, the lower experimental unit number of 50, and a high ratio of equi-proportional communities (0.9 or 1.0). However, the success rates at these two points remained relatively high ($>80\%$), see Figure 2.1 (a). In cases where the incorrect underlying model was selected, the three model structures (AV, FG, FULL) were chosen at a roughly equal rate ($AV \approx FG \approx FULL \approx 500$).

AV Model (as the true underlying model)

Similarly to the ID model, the models which contained a single interaction term (AV), performed consistently well across designs, being correctly selected by the `auto_DI()` function roughly 90% of the time (Figure 2.1 (b)). The only noticeable difference in the success rate of model selection occurred in the four-species experimental design which included one monoculture replicate and 50 experimental units, which generally performed slightly worse than the other designs, with its lowest point being for a design in which all mixtures were uneven (not equi-proportional), although this still lay above an 85% success rate, see Figure 2.1 (b). For the remaining 10-15% of cases, the incorrect model selections, the FG and FULL model structures were chosen at a roughly equal rate ($FG \approx FULL \approx 500$), while the ID model is very rarely selected.

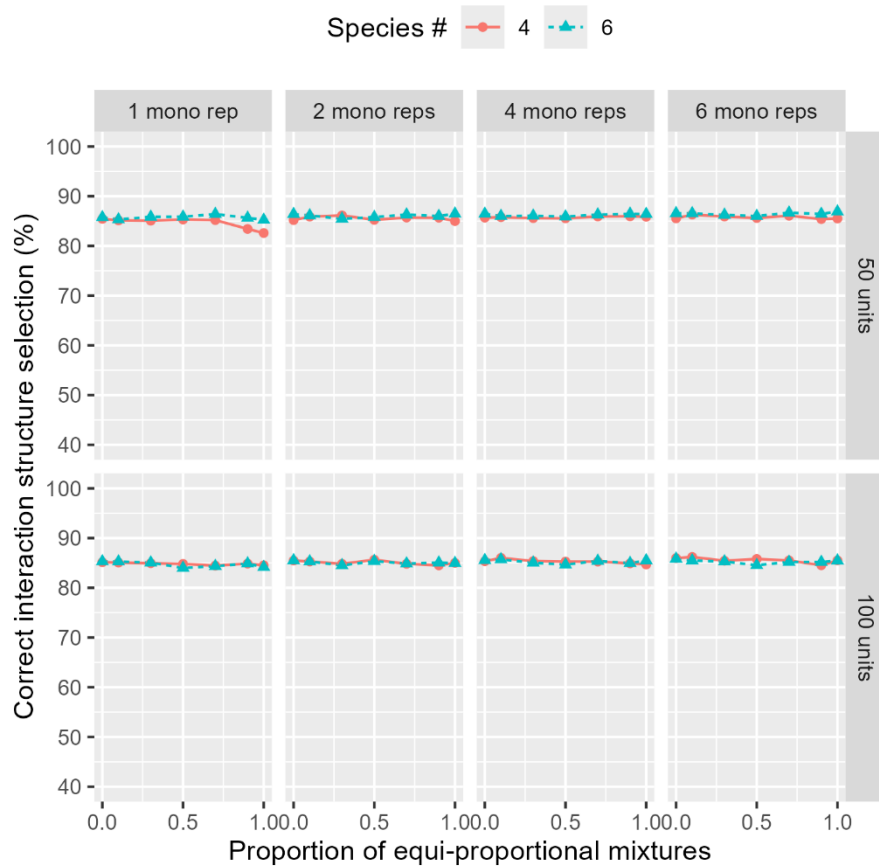
FG Model (as the true underlying model)

For the FG model, which contained either an additional three or six interaction terms compared to the ID model (for the four and six species designs respectively), the number of experimental units strongly affected the ability to detect the true underlying model as the best one (Figure 2.1 (c)). Designs including 100 experimental units consistently correctly

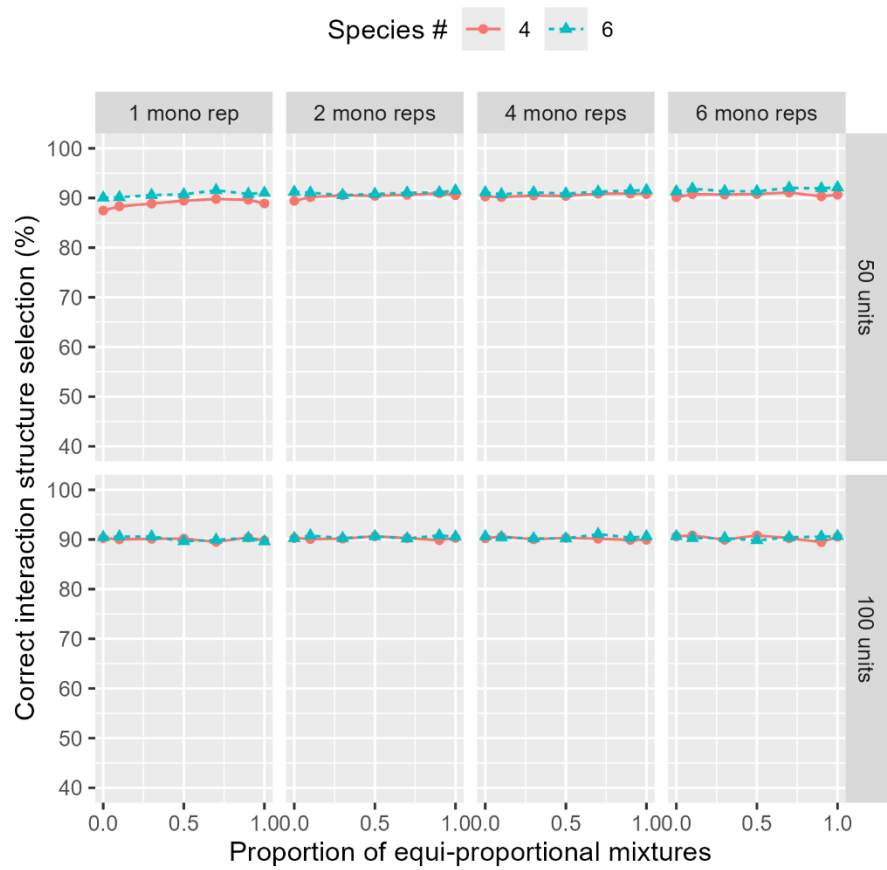
selected the true underlying model over 90% of the time while those containing 50 units ranged from roughly 45% to 90%. The results for the 100-unit designs also performed more consistently than their 50 units counterparts, although the four-species designs did experience a drop in efficacy as the proportion of sown equi-proportional mixtures increased. This drop in efficacy is consistent with that of the 50-unit designs, although much less exaggerated across monoculture replicate numbers and equi-proportional mixture proportions. For the 50-unit designs, both the four- and six-species designs experienced a fall in the rate of successful selection of the interaction structure across monoculture replicates as the proportion of equi-proportional mixtures rose, with the six-species designs out-performing the four-species designs for instances of one, two, and four monoculture replicates, but significantly dropping in efficacy, to below the four-species designs, for six monoculture replicates, see Figure 2.1 (c). For the cases of incorrect model selections, the AV model was overwhelmingly the chosen structure when only 50 experimental units were available and rarely for datasets with 100 units. The ID and FULL models were both consistently chosen across experimental unit numbers, with ID being very rarely chosen and FG being at a consistent rate of roughly 500 times.

Figure 2.1: A grid of line plots depicting the number of times, out of a possible 10,000, the `auto_DI()` function, from the `DImodels` R package (Moral *et al.*, 2023), correctly selected the true underlying (a) ID, (b) AV, or (c) FG interaction structure. The plots are divided by the number of monoculture replicates included in the design (one, two, four, or six) and the number of experimental units available (50 or 100). The solid line with circular points (red), represents designs which include four species while the dotted line with triangular points (blue) represents those with six species. As we move from left to right along the x-axis, the proportion of equi-proportional/even mixture communities increases.

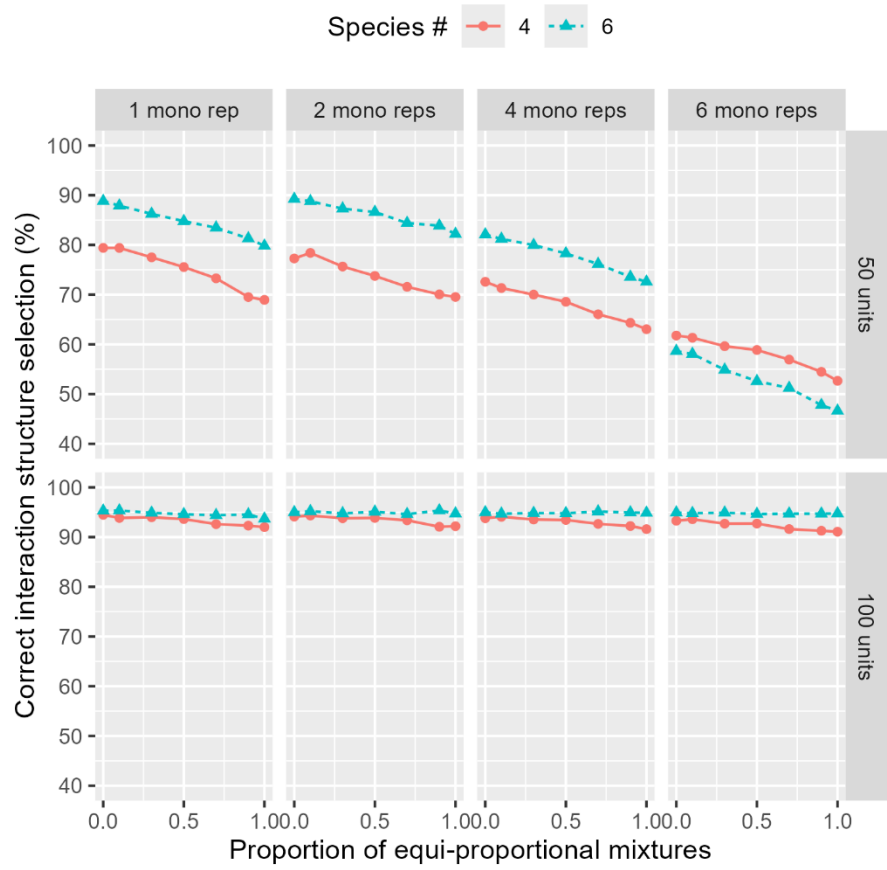
(a) - ID:



(b) - AV:



(c) - FG:



2.4 Conclusions

From the results outlined above I found that:

ID Model (as the true underlying model)

The ID model was consistently well selected as the true underlying model, seemingly unaffected by the tested variations in the experimental design principles. From this, I can conclude that, within the scope of this study, a lack of species interactions can be distinguished through the DI model framework regardless of the experimental design. However, the few points noted in the results where the success rate drops, for designs with 50 units, one monoculture replicate, and a high proportion of equi-proportional mixtures (0.9 or 1.0), could be attributed to less confidence at the monoculture points in the design, due to lack of replication. The variation between relative mixture performances and monoculture performances could therefore be assumed by the DI modelling framework to be interactions between species that weren't actually there. This highlights the importance of monoculture replicates for the estimation of species ID effects. The roughly equal selection of the three incorrect interaction structures ($AV \approx FG \approx FULL \approx 500$, Appendix A) could be due to a relatively large random error being assigned to each generated response as this could be attributed to non-existent interaction effects. This would explain the seemingly random selection of the incorrect interaction structures.

AV Model (as the true underlying model)

Similarly to the ID model, the consistently successful selection of the true underlying AV model leads me to conclude that the definition of an appropriate experimental design is robust in its ability to supply sufficient information for model estimation and selection. However, conversely to the ID model's worst designs being at high proportions of equi-proportional mixtures with 50 units and one monoculture replicate, the AV model was actually least selected at high proportions of uneven mixes, i.e. low proportions of equi-

proportional mixtures. This could be due to too high of variability in the performance of the mixtures, with little to no replicates, compared to the designs which included equi-proportional mixtures, for which there are limited options for species community designs and therefore will have some replication each. Variation in uneven mixture communities would then likely cause the DI modelling selection method to attribute differences in communities to differences in species compositions (IDs) and their resulting interactions, coupled with no additional monoculture replications per species leading to a further estimation of variable species interactions, leaving the DI modelling framework to select structures containing multiple interaction effects. The generally better performance of the AV structure compared to the ID (90% vs 85%) could be explained using the same logic as the ID model results, that is the incorrect assignment of the random error to random interactions ($FG \approx FULL \approx 500$, Appendix A). As the AV model is the correct structure, the roughly 5% difference in performance between the ID and AV models can be attributed to the error inflating the AV interaction coefficient term, but not changing its strength enough for an alternative structure to be incorrectly chosen.

FG Model (as the true underlying model)

For the most complex interaction structure tested in this study, the FG model, a considerable change was found in the success rates of model selection between experimental design principles compared to the simpler ID and AV models. This simplification of the results could immediately point to more specific information requirements from an experiment's dataset for models with larger numbers of fixed effects to be estimated. More specifically, the designs which included one or two monoculture replicates experienced very similar results across species pool size, experimental units available, and the proportion of equi-proportional mixtures, while the designs with four monoculture replicates follows the same downward trend as the proportion of equi-proportional mixes increases, but with a lower start point. The designs with six monoculture replicates suffer greatly from their lowered

availability of mixture communities, significantly more so for the six-species designs (top-right portion of Figure 2.1(c), the only instance found where the four-species outperforms the six-species designs) as the highly distinguishable herb ID effects (Table 2.3) can seemingly no longer compensate for the lack of sufficient replication of both monocultures and mixtures. In cases where the incorrect interaction structure was selected, the FULL model was consistently chosen at the same rate as the true ID and AV models (~500), regardless of experiment design, and the ID was model rarely chosen incorrectly at all across simulations. While the AV model was very rarely chosen for experiments with 100 units, it made up the vast majority of incorrectly chosen models for 50 unit experiments (Appendix A). The FULL and ID model selections follows with conclusions from the AV and ID true models, wherein it is assumed that the random error is attributed to non-existent interaction terms, therefore pushing for a more complex DI model to be selected. The AV model for 100 units also follows previous conclusions, whereas for the 50 unit designs it is more likely that the lack of experimental units available resulted in not enough information being provided to the models to accurately capture the strength of the between- and within- functional group interactions, leading to the `auto_DI()` function selecting the simpler AV structure.

Overall

These conclusions may be generalised to say that, within the scope of this study, it is generally beneficial for the correct selection of the true underlying model to include a small, but non-zero, proportion of equi-proportional mixtures and a lower number of monoculture replicates than six. The general outperformance of the six-species experiments compared to their four-species counterparts could be attributed to the additional two (herb) species having significantly lower true ID effects than grasses and legumes (as is usually seen in real-world experiments examining the total yield of an ecosystem, see section 4.3 of this thesis), see Table 2.3. The ease of differentiation between these ID effects would likely lower the requirement of information provided by an experimental design to separate these effects in

a DI model, despite their inclusion in this study (within the six-species designs) resulting in more fixed effects needing estimation and more experimental units taken up by monocultures (as at least one exists for each species in every design) than those containing only the similarly performing (in monoculture) species of grasses and legumes.

2.5 Discussion and future work

In this chapter, a simulation study was used to show that there is evidence indicating that, under the conditions tested, uneven mixtures should be prioritised over equi-proportional mixtures or monoculture replicates when designing a BEF relationship study with which the DI modelling framework (Kirwan *et al.*, 2009) will be used for analysis. This knowledge will aid in the design of future BEF relationship studies, with optimal designs ensuring that limited experimental resources are used efficiently, providing sufficient coverage of the design space to correctly estimate the true underlying structure of species interactions. The code used to run this simulation study is also highly customisable and therefore may be easily adapted for the future testing of experimental design principles, see Appendix B.

In practice, the incorrect selection of the true underlying DI model may not be detrimental to the analysis of BEF data. The interaction effects captured using DI models are not derived directly from any known biological processes, e.g. niche partitioning, and are instead estimates of net interactions. This means that rather than requiring the selected interactions structure to match the true underlying structure exactly, the general interpretation of the contribution of each species to ecosystem functioning should instead be considered at the forefront. As such the potential for over-fitting in the models, and therefore the assessment of model fits, is of particular importance when using DI models for BEF studies. The selection of a more complex DI model than the true underlying interaction structure could inflate the contribution of individual species over functional groups or species richness, resulting in often more complex and/or expensive community designs to

be incorrectly favoured for better ecosystem performance. This concern is validated by the results of this simulation study, which found that the vast majority of incorrectly selected interaction structures (save for when the AV model was selected over the FG model, which is theorised to be due to lack of information from the experimental design) were of an increased complexity to the true underlying DI model.

The generalisation of the conclusions drawn from this simulation study would require a more exhaustive study, with further variance of the assumed simulation parameters. While the randomness of the experiment designs within sets of design principle combinations helped to reduce bias in this study, many assumptions were made on the nature of the experimental designs in question and future work may involve broadening their scope. For example, the design principle parameters included in this study may be varied, including additional numbers of species (species pool sizes), monoculture replicates, or experimental units, or entirely new parameters could be included, such as treatments, blocking effects, or split-plot designs.

The fixed effects and variance of the error term from the true underlying model used to simulate the response variable may also be varied to test this study's conclusions under the influence of different species dynamics. This may vastly change the conclusions from this study as it was noted that the differences between the six- and four-species designs could be explained through the noticeable difference between the herb ID effects and the grass/legume ID effects. The inclusion of the full pairwise (FULL) DI model structure (Kirwan *et al.*, 2009; Moral *et al.*, 2023) and the non-linear parameter θ at values different to one (Connolly *et al.*, 2013; Moral *et al.*, 2023; Vishwakarma *et al.*, 2023) would also be of interest to include in future iterations of this study due to their increase in the number of parameters (thus complexity) of a DI model. From the study presented in this chapter, an increase in model complexity seemingly incurs a preference for the use of uneven species

compositions over even compositions or monoculture replicates in experimental unit designs, and future work will determine how this may not hold under varying conditions.

For this study, only the correct selection of some true underlying interaction model was tracked between experimental designs, inherently putting an emphasis on mixtures over monocultures. Future studies could instead examine the effects of changing the same parameters presented here on coefficient estimates, particularly for the ID effects of each species (as the design of the experiment can heavily influence both the bias of the ID effect estimates and their corresponding standard errors, thus why monoculture replicates are considered important for their correct estimation) and the variance of the model's error term.

This study did not take into account certain design principles which aim to mitigate problems such as the replication of community designs as 'back-ups' in the case that establishment experiences issues. The simulated results here should therefore only be followed as a guide in the context of the settings that were manipulated in my study.

3 A Diversity-Interactions model for the fitting of multivariate repeated measures data

Collaborators: Caroline Brophy, Rishabh Vishwakarma, Rafael de Andrade Moral

This chapter was formatted as a standalone paper intended for submission to Methods in Ecology and Evolution.

CB conceived the package idea; LB wrote the R package with input from all co-authors on its design; LB and RV led package testing; LB led the writing of the manuscript.

Abstract

1. Diversity-Interactions (DI) modelling is a regression-based approach to modelling the BEF relationship. This methodology enables the assessment of the effects of manipulating (initial) species richness, evenness, and composition on ecosystem function responses measured at a later point in time. The models separate the effects of each individual species and test the strength of potential species interactions.
2. Univariate DI models, fit with continuous response data gathered from a singular time point, can easily be fit using the `DImodels` R package. The `DImodels` package contains highly customisable fitting options but cannot handle data with multiple ecosystem functions (multivariate) or multiple time points (repeated measures).
3. The multivariate DI model simultaneously predicts for multiple ecosystem functions whilst accounting for covariances between them. The repeated measures DI model likewise enables the prediction of a single ecosystem function at multiple time points and incorporates any existing covariances between the time points.

4. In this chapter, I introduce the multivariate repeated measures DI model which enables the joint modelling of multiple continuous ecosystem functions, measured at multiple time points, whilst accounting for any covariances between readings from the same experimental/observational unit. I also present the R package `DImodelsMulti`, a complementary package to `DImodels`. It enables the fitting of multivariate and/or repeated measures DI models in a user-friendly way. I demonstrate its usage through a worked example.

3.1 Introduction

Biodiversity and ecosystem function (BEF) relationship studies aim to quantify how species diversity affects the quality and quantity of the goods and services that an ecosystem provides (de Groot, 1992; Costanza *et al.*, 1997). The design of these studies can include complex structures, for example, the measurement of multiple responses from the same ecosystem, multi-site experiments, or split-plot designs. The measuring and recording of multiple ecosystem function responses (multivariate) at multiple time points (repeated measures) in BEF relationship studies is of high importance, as it enables the analysis of the multifunctionality and temporal patterns (e.g. the resilience) of an ecosystem under different conditions (Hector & Bagchi, 2007; Finn *et al.*, 2013; Tilman, *et al.*, 2014; Dooley *et al.*, 2015; Oliver *et al.*, 2015; Manning *et al.*, 2018). As such, the development of statistical modelling methods that can assess the multiple dimensions of species diversity on ecosystem functioning while capturing the hierarchical complexities in data structures, such as those seen in repeated measures and multivariate data, could be beneficial to improving our understanding of the BEF relationship.

Diversity-Interactions (DI) modelling is a regression-based approach which can be used for analysis in BEF relationship studies (Kirwan *et al.*, 2009; Connolly *et al.*, 2013; Moral *et al.*, 2023). The models include an identity (ID) effect for each species, represented

by its initial relative abundance or proportion, interaction effects between species (which are customisable according to varying the biological assumptions of how species may interact), and can include additional block or treatment effects. This enables the joint assessment of how initial species richness, evenness, and composition can affect ecosystem functioning at future time points. The `DImodels` R package (Moral *et al.*, 2023) can be used to automatically fit and compare a range of univariate DI models to data from a BEF relationship study.

DI models have already been adapted to fit data with various data complexities, such as the inclusion of correlated responses from repeated measures (Frankow-Lindberg *et al.*, 2009; Finn *et al.*, 2013) and from multiple recorded ecosystem functions (Dooley *et al.*, 2015). However, they have not yet been able to model responses where there are multiple sources of correlation, i.e. for data that contains multiple ecosystem functions and multiple time points. The `DImodels` package also cannot be used to fit multivariate data nor data with responses taken at multiple time points.

The availability of R technology for fitting regression models capable of capturing correlation between responses from multiple sources is limited. Most methods only allow for a single source of correlation, e.g. the `lme4` package (Bates *et al.*, 2015) or the `nlme` package (Pinheiro *et al.*, 2024), or use a Bayesian framework, which involves the selection and use of informative priors, and many require a user to format the variance-covariance matrix themselves, e.g. the `brms` package (Bürkner, 2017) and the `MCMCglmm` package (Hadfield, 2010).

In this chapter, I will:

1. Introduce the multivariate repeated measures DI model and demonstrate its benefits.
2. Present the R package `DImodelsMulti` which enables the fitting of the DI model framework to multivariate data, repeated measures data, or data with both complexities included.

3.2 Multivariate and repeated measures DI models

DI modelling (Kirwan *et al.*, 2009) assumes that the main driver behind changes in ecosystem functioning is the initial relative abundances, or proportions, of the species in the study. An example of a univariate DI model for data collected at a single point in time is:

$$y_m = \sum_{i=1}^S \beta_i p_{im} + \sum_{\substack{i,j=1 \\ i < j}}^S \delta_{ij} (p_{im} p_{jm})^\theta + \alpha \mathbf{A}_m + \varepsilon_m, \quad \varepsilon_m \sim N(0, \sigma^2) \quad (\text{Eq. 3.1})$$

where β_i , scaled by the initial proportion of species i , is the expected contribution of the i^{th} species to the ecosystem function response y and is referred to as the i^{th} species' ID effect; if no species interactions or treatments are required in the model, the expected response y is the weighted average of the species identity effects. The subscript m refers to an experimental/observational unit, e.g. a plot, with $m = 1, 2, \dots, M$. Here I show the full pairwise interaction structure (see Table 1.2 for alternatives), where it is assumed that each pair of species interacts in a unique way. Each δ_{ij} coefficient represents the interactive potential between species i and j , and is scaled by the product of their proportions to determine the contribution to the response. S is the number of unique species present in the study, $\mathbf{A}$ may include blocks or treatment or other covariate terms, and α is a vector of the corresponding effect coefficients. The inclusion of the non-linear parameter θ as a power to the product of species proportions within the interaction structure allows for the shape of the BEF relationship to change (Connolly *et al.*, 2013; Moral *et al.*, 2023; Vishwakarma *et al.*, 2023; Finn *et al.*, 2024).

DI models have previously been used to model multiple ecosystem functions or multifunctionality (multivariate data; Dooley *et al.*, 2015) and to model ecosystem functions measured repeatedly over time (repeated measurements data; Frankow-Lindberg *et al.*, 2009, Finn *et al.*, 2013). Here, I introduce an adapted DI model suitable for modelling both multivariate and repeated measures data, i.e., a DI model for data on multiple ecosystem

functions each recorded over multiple time points sometime after the initial species proportions were measured. An example multivariate repeated measures DI model with the average pairwise species interaction structure (see Table 1.2 for alternatives), which makes a biological assumption that all pairs of species interact in the same way and can therefore be represented by a single coefficient, δ_{AV} , is given in Eq. 3.2.

$$y_{kmt} = \sum_{i=1}^S \beta_{ikt} p_{im} + \delta_{AVkt} \sum_{\substack{i,j=1 \\ i < j}}^S (p_{im} p_{jm})^{\theta_k} + \alpha_{kt} \mathbf{A}_m + \varepsilon_{kmt}, \quad \varepsilon_{kmt} \sim MVN(0, \Sigma^*) \quad (\text{Eq. 3.2})$$

Where the subscript k refers to an ecosystem function and $k = 1, 2, \dots, K$ and the subscript t refers to a time point at which the ecosystem function was recorded, with $t = 1, 2, \dots, T$. As such, y_{kmt} refers to the value of the reading of the k^{th} ecosystem function on the m^{th} experimental unit at a time point t . The initial proportions of the species, p , are generally manipulated or measured at some initial time point, $t = 0$, however, the specification shown in Eq. 3.2 can be easily adapted to instead use the species proportions manipulated/measured at a time point $t - 1$ for each y_t . The nonlinear term θ_k (Connolly *et al.*, 2013) allows for the interaction structures to vary between ecosystem function responses, and thus allows the shape of the BEF relationship to change for each ecosystem function (Moral *et al.*, 2023; Vishwakarma *et al.*, 2023; Finn *et al.*, 2024). The structural terms, defined for each experimental unit m via the term $\mathbf{A}$, have the corresponding coefficient vector α_{kt} . As the ecosystem functions, k , are read from the same experimental units, the vector $\mathbf{A}_m$ cannot depend on k . A treatment applied to a unit may vary across time, t , and therefore may influence this vector ($\mathbf{A}_{mt}$), however, in this case it may be better to describe treatments using plans rather than individual applications to avoid this complication, e.g. $\mathbf{A}$ could define the level of fertiliser applied to a unit (1 or 2), which is varied across time, or it could be defined as the fertiliser treatment plan (A, B, or C) that the unit is subjected to.

The errors of these models are assumed to follow a multivariate normal distribution, with mean zero and variance Σ^* , which is in the form of a positive-definite block diagonal

matrix of size $KMT \times KMT$, with a square matrix Σ_m in each cell along the main diagonal for each experimental unit m and zero values elsewhere (Eq. 3.3).

$$\Sigma^* = \begin{bmatrix} \Sigma_1 & & 0 \\ & \ddots & \\ 0 & & \Sigma_M \end{bmatrix} \quad (\text{Eq. 3.3})$$

Each block diagonal matrix Σ_m for plot m , is of size $KT \times KT$, contains the variance of each ecosystem function at each time point along its main diagonal and the covariances between each pair of ecosystem functions and time points in the off-diagonal cells. It is structured as the Kronecker product ($\otimes$) of the separate variance-covariance matrices for the ecosystem function responses and repeated measure responses (Galecki, 1994), see Eq. 3.4, where $K = 3$ and $T = 2$. The format of the variance-covariance matrix in Eq. 3.4 is referred to as being unstructured (UN), or following the general structure, as each element is estimated and allowed to vary from one another. The structure of these matrices can be varied (simplified) according to *a priori* information or assumptions on the relationship between variance and/or covariance values (see Table 3.1), however, this is only recommended for repeated measures. In the case of multivariate DI models, simplification of the variance-covariance matrix would enforce a strict assumption on the nature of the relationships between ecosystem functions, which generally would be unrealistic. For example, the AR(1) structure assumes that the variances of all responses are equal and that the correlations between responses lessen as the distance between recordings widen, however, the order of ecosystem function responses isn't relevant, therefore making this assumption illogical. Similarly, the compound symmetry (CS) structure assumes that the variance of each response is equal and the covariances between each response are also equal, which is an equally strict yet unrealistic assumption for multiple ecosystem functions, however, this structure may be useful for the simplification of a model which otherwise does not converge. There are many other structures available for these matrices, each with their own assumptions imposed on the responses, and therefore it is possible that some structures

exist that have significant biological interpretations with a realistic assumptions for multiple ecosystem functions.

$$\begin{aligned}
\Sigma_m &= \begin{bmatrix} \sigma_{k_1}^2 & \sigma_{k_1 k_2} & \sigma_{k_1 k_3} \\ \sigma_{k_1 k_2} & \sigma_{k_2}^2 & \sigma_{k_2 k_3} \\ \sigma_{k_1 k_3} & \sigma_{k_2 k_3} & \sigma_{k_3}^2 \end{bmatrix} \otimes \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} \\
&= \begin{bmatrix} \sigma_{k_1}^2 \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} & \sigma_{k_1 k_2} \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} & \sigma_{k_1 k_3} \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} \\ \sigma_{k_1 k_2} \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} & \sigma_{k_2}^2 \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} & \sigma_{k_2 k_3} \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} \\ \sigma_{k_1 k_3} \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} & \sigma_{k_2 k_3} \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} & \sigma_{k_3}^2 \begin{bmatrix} \sigma_{t_1}^2 & \sigma_{t_1 t_2} \\ \sigma_{t_1 t_2} & \sigma_{t_2}^2 \end{bmatrix} \end{bmatrix} \\
&= \begin{bmatrix} \begin{bmatrix} \sigma_{k_1}^2 \sigma_{t_1}^2 & \sigma_{k_1}^2 \sigma_{t_1 t_2} \\ \sigma_{k_1}^2 \sigma_{t_1 t_2} & \sigma_{k_1}^2 \sigma_{t_2}^2 \end{bmatrix} & \begin{bmatrix} \sigma_{k_1 k_2}^2 \sigma_{t_1}^2 & \sigma_{k_1 k_2}^2 \sigma_{t_1 t_2} \\ \sigma_{k_1 k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_2}^2 \sigma_{t_2}^2 \end{bmatrix} & \begin{bmatrix} \sigma_{k_1 k_3}^2 \sigma_{t_1}^2 & \sigma_{k_1 k_3}^2 \sigma_{t_1 t_2} \\ \sigma_{k_1 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_3}^2 \sigma_{t_2}^2 \end{bmatrix} \\ \begin{bmatrix} \sigma_{k_1 k_2}^2 \sigma_{t_1}^2 & \sigma_{k_1 k_2}^2 \sigma_{t_1 t_2} \\ \sigma_{k_1 k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_2}^2 \sigma_{t_2}^2 \end{bmatrix} & \begin{bmatrix} \sigma_{k_2}^2 \sigma_{t_1}^2 & \sigma_{k_2}^2 \sigma_{t_1 t_2} \\ \sigma_{k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_2}^2 \sigma_{t_2}^2 \end{bmatrix} & \begin{bmatrix} \sigma_{k_2 k_3}^2 \sigma_{t_1}^2 & \sigma_{k_2 k_3}^2 \sigma_{t_1 t_2} \\ \sigma_{k_2 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_2 k_3}^2 \sigma_{t_2}^2 \end{bmatrix} \\ \begin{bmatrix} \sigma_{k_1 k_3}^2 \sigma_{t_1}^2 & \sigma_{k_1 k_3}^2 \sigma_{t_1 t_2} \\ \sigma_{k_1 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_3}^2 \sigma_{t_2}^2 \end{bmatrix} & \begin{bmatrix} \sigma_{k_2 k_3}^2 \sigma_{t_1}^2 & \sigma_{k_2 k_3}^2 \sigma_{t_1 t_2} \\ \sigma_{k_2 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_2 k_3}^2 \sigma_{t_2}^2 \end{bmatrix} & \begin{bmatrix} \sigma_{k_3}^2 \sigma_{t_1}^2 & \sigma_{k_3}^2 \sigma_{t_1 t_2} \\ \sigma_{k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_3}^2 \sigma_{t_2}^2 \end{bmatrix} \end{bmatrix} \\
&= \begin{bmatrix} \sigma_{k_1}^2 \sigma_{t_1}^2 & \sigma_{k_1}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_2}^2 \sigma_{t_1}^2 & \sigma_{k_1 k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_3}^2 \sigma_{t_1}^2 & \sigma_{k_1 k_3}^2 \sigma_{t_1 t_2} \\ \sigma_{k_1}^2 \sigma_{t_1 t_2} & \sigma_{k_1}^2 \sigma_{t_2}^2 & \sigma_{k_1 k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_2}^2 \sigma_{t_2}^2 & \sigma_{k_1 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_3}^2 \sigma_{t_2}^2 \\ \sigma_{k_1 k_2}^2 \sigma_{t_1}^2 & \sigma_{k_1 k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_2}^2 \sigma_{t_1}^2 & \sigma_{k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_2 k_3}^2 \sigma_{t_1}^2 & \sigma_{k_2 k_3}^2 \sigma_{t_1 t_2} \\ \sigma_{k_1 k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_2}^2 \sigma_{t_2}^2 & \sigma_{k_2}^2 \sigma_{t_1 t_2} & \sigma_{k_2}^2 \sigma_{t_2}^2 & \sigma_{k_2 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_2 k_3}^2 \sigma_{t_2}^2 \\ \sigma_{k_1 k_3}^2 \sigma_{t_1}^2 & \sigma_{k_1 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_2 k_3}^2 \sigma_{t_1}^2 & \sigma_{k_2 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_3}^2 \sigma_{t_1}^2 & \sigma_{k_3}^2 \sigma_{t_1 t_2} \\ \sigma_{k_1 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_1 k_3}^2 \sigma_{t_2}^2 & \sigma_{k_2 k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_2 k_3}^2 \sigma_{t_2}^2 & \sigma_{k_3}^2 \sigma_{t_1 t_2} & \sigma_{k_3}^2 \sigma_{t_2}^2 \end{bmatrix} \quad (\text{Eq. 3.4})
\end{aligned}$$

Table 3.1: Example covariance structures available in the `DIModelsMulti` R package.

For univariate repeated measures data, three structures are available, the formats of which can be seen in (a): unstructured/general (UN), compound symmetry (CS), and autoregressive model of order one (AR(1)). For multivariate data from a single time point, it is recommended to only use the UN structure, however, the CS structure may also be used if simplification is required for the model (e.g., if the UN structure leads to convergence issues). In the multivariate repeated measures case, the structures shown in (a) can be combined using the Kronecker product of each matrix (with the multivariate matrix on the left-hand side of the product and the repeated measures matrix on the right), see for example Eq. 3.4. The current available options for the combined autocorrelation structures are $\text{UN} \otimes \text{UN}$, $\text{UN} \otimes \text{CS}$, $\text{UN} \otimes \text{AR}(1)$, $\text{CS} \otimes \text{UN}$, $\text{CS} \otimes \text{CS}$, and $\text{CS} \otimes \text{AR}(1)$, two of which are shown in (b).

(a)

Unstructured (UN)	Compound Symmetry (CS)	Autoregressive Model Order 1 (AR(1))
$K = 3, \quad T = 1$ $\begin{bmatrix} \sigma_1^2 & \sigma_{1,2} & \sigma_{1,3} \\ \sigma_{1,2} & \sigma_2^2 & \sigma_{2,3} \\ \sigma_{1,3} & \sigma_{2,3} & \sigma_3^2 \end{bmatrix}$	$K = 1, \quad T = 4$ $\sigma^2 \begin{bmatrix} 1 & \rho & \rho & \rho \\ \rho & 1 & \rho & \rho \\ \rho & \rho & 1 & \rho \\ \rho & \rho & \rho & 1 \end{bmatrix}$	$K = 1, \quad T = 4$ $\sigma^2 \begin{bmatrix} 1 & \rho & \rho^2 & \rho^3 \\ \rho & 1 & \rho & \rho^2 \\ \rho^2 & \rho & 1 & \rho \\ \rho^3 & \rho^2 & \rho & 1 \end{bmatrix}$

(b)

$\text{UN} \otimes \text{AR}(1)$	$\text{CS} \otimes \text{UN}$
$K = 2, \quad T = 3$ $\begin{bmatrix} \sigma_1^2 & \sigma_{1,2} \\ \sigma_{1,2} & \sigma_2^2 \end{bmatrix} \otimes \sigma^2 \begin{bmatrix} 1 & \rho & \rho^2 \\ \rho & 1 & \rho \\ \rho^2 & \rho & 1 \end{bmatrix}$	$K = 4, \quad T = 2$ $\sigma^2 \begin{bmatrix} 1 & \rho & \rho & \rho \\ \rho & 1 & \rho & \rho \\ \rho & \rho & 1 & \rho \\ \rho & \rho & \rho & 1 \end{bmatrix} \otimes \begin{bmatrix} \sigma_1^2 & \sigma_{1,2} \\ \sigma_{1,2} & \sigma_2^2 \end{bmatrix}$

These DI models can be estimated using generalised least squares, therefore, the standard errors of the estimated coefficients can be found by taking the square root of the values along the main diagonal of the (approximate) covariance matrix of the coefficient estimates (i.e. the inverse of the Hessian matrix associated with the model). The p-values associated with each coefficient estimate may also be produced under a t approximation, however, as there is difficulty in deciding the number of degrees of freedom to be used under this approximation for generalised least squares models, many statisticians disagree with their usage. I believe that p-values may still provide value as a guide for model selection or provide a useful interpretation for an explanatory variable of interest, however, they should not be examined in isolation. For example, if the estimated coefficient of a variable of interest has a p-value which indicates insignificance, but its inclusion appears to benefit the model when model selection techniques such as likelihood ratio tests or information criteria (AIC, BIC, etc.) are employed, then, in my opinion, the variable should be kept in the model. This is likely to occur when comparing interaction structures within the DI modelling framework, as multiple new variables may be included with each step in the model selection, e.g., the FULL DI model may outperform the AV DI model for a dataset without every pairwise interaction being significant, given that there are many other pairs which may be of significance.

3.3 DImodelsMulti R package

Suitable data

Data suitable for use with this package should be from a study, either observational or experimental, which contains a set of proportional predictors with (a) multiple ecosystem function responses, (b) a single ecosystem function response measured at multiple time points, or (c) multiple ecosystem function responses measured at multiple time points, e.g. a BEF study where species proportions are manipulated for multiple ecosystem function

responses at multiple time points. The recorded ecosystem function response(s) should be continuous. The data should contain the relative abundances (proportions) of the predictors in the study measured at some time before the earliest response time point, e.g., the initial sown species proportions at $t = 0$. If a study spans over a long time period, where the earliest proportion predictors (measured at $t = 0$) would no longer logically be representative of the study environment, e.g. the ecosystem, the proportions measured at each time $t - 1$ can instead be used for each y_t . Each experimental and/or observational unit in the study should be assigned a unique identifier, which remains constant throughout time points and response types, to enable the groupings of responses to the subject of their inherent correlation. The dataset may also include treatments or blocking variables (but not required).

Package structure and usage

I have developed the `DImodelsMulti` package (Byrne *et al.*, 2024) that acts as a sister package to the `DImodels` package (Moral *et al.*, 2023), facilitating a user-friendly way to fit Diversity-Interactions models to multivariate and/or repeated measures data. I have already published the package on CRAN and the package has over 3,800 downloads to date (August 2024). The `DImodelsMulti` R package's layout and intended workflow is in Figure 3.1.

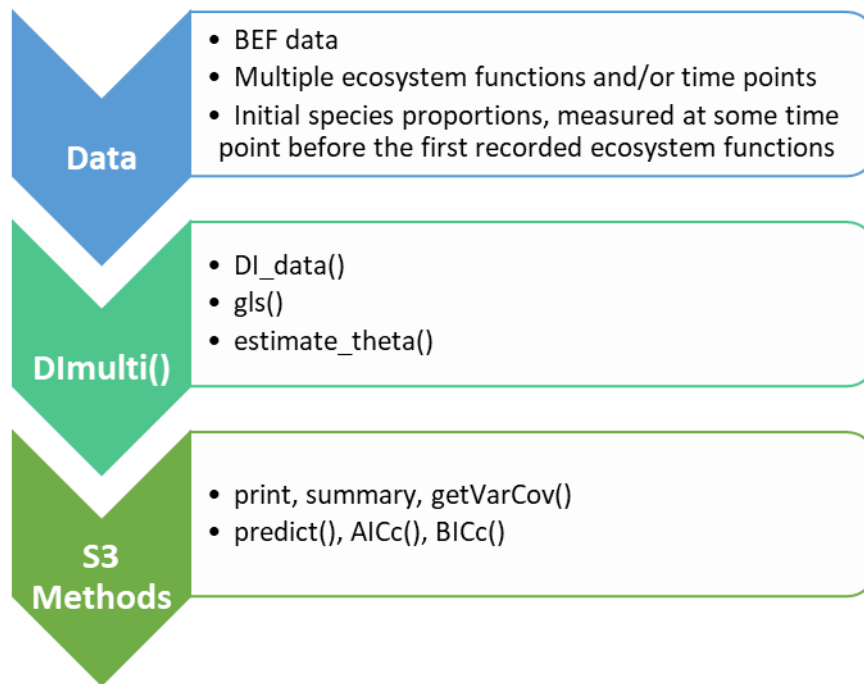


Figure 3.1: The layout and workflow of the `DImodelsMulti` R package (Byrne *et al.*, 2024). Data from a BEF study is passed to the main function of the package, which returns a model object of custom class `DImulti`. The package contains a number of S3 methods which can be used to further analyse the model output.

The main function of the package is `DImulti()`, which can fit multivariate (and univariate) repeated measures DI models. The syntax for this function can be seen below:

```
DImulti(y, eco_func = c("NA", "NA"), time = c("NA", "NA"),
        unit_IDs, prop, data, DImodel,
        FG = NULL, ID = NULL, extra_fixed = NULL,
        estimate_theta = FALSE, theta = 1,
        method = "REML")
```

The `y`, `eco_func`, and `time` arguments specify the responses of the model and `unit_IDs` is used to specify the experimental/observational unit. The `prop` argument specifies the columns in the dataset, supplied through `data`, which hold the initial species proportions, while any additional columns which hold structural effects information may be passed through `extra_fixed`. The identity effects of the species specified through `prop`

may be grouped, assuming equal identity contributions of groups of species, using the `ID` parameter. The function includes a number of species interaction structures for automatic fitting, which can be specified through the `DImodel` argument. The options included align with those previously described in DI models literature (Kirwan *et al.*, 2009; McDonnell *et al.*, 2023) and available in the `DImodels` package (Moral *et al.*, 2023) (see Table 1.2 for further details):

- **ID**: The identity effect structure, where it is assumed that species do not interact.
- **AV**: The average pairwise interaction structure, where it is assumed that all pairs of species interact in the same way.
- **ADD**: The additive species effects structure, where it is assumed that each species makes the same contribution to a pairwise interaction no matter which species it is interacting with.
- **FG**: The functional group interactions structure, where species are grouped by their functional traits and it is assumed that species interact in one way with species from the same functional group (within FG interactions) and another way with each other group (between FG interactions). The groupings are specified through the `FG` parameter.
- **FULL**: The full pairwise interactions structure, where it is assumed that each pair of species interacts in a unique way.

The θ_k parameter is applied as a power to each pairwise product of proportions in the interaction terms of the model. The current recommendation for the order of estimation of θ in the model selection procedure for multivariate DI models follows that of Vishwakarma *et al.*, 2023 for univariate DI models:

1. Simultaneously estimate each value for θ_k using the FULL interaction structure if sufficient data is available and the AV interaction structure otherwise.

2. Fit K new `DImulti` multivariate model objects, one for each θ_k value estimated, holding each other value of θ constant at a value of one.
3. Individually test the inclusion of each estimated θ_k value by comparing each new model against the FULL (or AV) model with all $\theta_k = 1$ using information criteria, selecting the value of θ for response type k from the better performing model.
4. For each value k , hold the selected value of θ constant for the remainder of the analysis.

These interaction structures can be compared in the same hierarchical fashion as their univariate counterpart (Kirwan *et al.*, 2009), using likelihood ratio tests for nested models or the standard correction to Akaike's Information Criteria (AICc) otherwise. The corrected version of AIC is recommended to prevent a bias for more complex models as the number of parameters in the models increases (Hurvich & Tsai, 1989; McQuarrie & Tsai, 1998). The ML (maximum likelihood) estimation method should be used for these comparisons as the fixed effects vary between models. The final selected model should then be refit using the REML (restricted maximum likelihood) method, in order to produce unbiased variance and covariance estimates.

The `DImulti()` function currently has three autocorrelation structures available for repeated measures (UN, CS, & AR(1)) and two structures available for multivariate data (UN & CS), see Table 3.1. A variance-covariance matrix is said to be unstructured (UN) or follow a general covariance structure when each element is allowed to differ from one another, as such, it is the most flexible structure. A simpler option is compound symmetry (CS), where it is assumed that each ecosystem function or time point has the same variance value σ^2 and each pair has the same covariance value $\sigma^2\rho$. This structure is only recommended for use with multivariate data if estimation using the UN structure is not possible, e.g. due to overparameterisation, as it makes a strong assumption about ecosystem function covariances that may not be realistic. A structure exclusive to repeated measures

data is an autoregressive model of order one (AR(1)), which assumes that, like CS, each time point has equal variance, σ^2 , and as the distance in pairs of time points increases, the covariance between them changes by a factor of ρ . It is currently recommended to fit and test the different repeated measures autocorrelation structures using the simplest model available (usually the ID model) and information criteria, such as AICc (corrected Akaike information criterion) or BICc (corrected Bayesian information criterion), holding the selected structure constant for the remainder of the analysis. In addition, once the final fixed effects models has been selected, a final re-check of the best correlation structure should be carried out for completeness. The REML estimation method should be used to fit models where autocorrelation structures are being compared and fixed effects should be held constant for these comparisons (Hox *et al.*, 2010; Snijders & Bosker, 2011).

Underlying estimation processes

The `DImulti()` function relies on the `gls()` function from the R package `nlme` (Pinheiro *et al.*, 2024) as its basis. The `gls()` function was selected as it uses generalised least squares to estimate the model coefficients, which is a method that can account for correlation between responses from a shared study subject. This function isn't typically used for fitting multivariate regression models, but following the methodology given by Bolker, 2012, this 'repeated measures' function can be used for either repeated measures or multivariate analyses. I build upon this function's usage to enable the fitting of multivariate repeated measures models (i.e., for modelling multivariate responses that are collected over time so *both* multivariate and repeated measures data). Within the `DImulti()` function, this complex variance-covariance matrix is created by first fitting models for the multivariate data and the repeated measures data separately using the `gls()` function, holding either the time or ecosystem function type constant respectively, extracting the estimated variance-covariance matrix from the model and combining these matrices using the Kronecker

product, as shown in Eq. 3.4 and Table 3.1, fixing their coefficients (so that no further estimation occurs) before passing them back through `gls()` for the final model fitting.

For the fixed effects, when both multiple time points and multiple ecosystem functions are included in the model, it is ensured that each predictor has a corresponding coefficient for each time point and ecosystem function. This is done by including the time and ecosystem function indicators as factor predictor terms in the model and crossing them with each of the other fixed effect terms in the model.

The values of θ for each ecosystem function (i.e. each θ_k) can be estimated using profile log-likelihood optimisation (Brent, 1973) or can be assigned a set value based on an *a priori* assumption/knowledge. To estimate θ_k within the `DImulti()` function, the R function `optim()` is used to perform joint profile likelihood optimisation, where each θ_k value is bound between 0.01 and 1.50. This range is chosen (following its usage in the `DImodels` R package) as a value of 0 for θ would result in the interactions between species not present contributing to the response, e.g. for $p_1 = 0$ and $p_2 = 0.5$, $(p_1 p_2)^{\theta=0} = 1$, while values over 1.50 have not yet been observed in existing DI analyses of BEF data.

3.4 A case study

The 'Belgium' dataset

Dooley *et al.*, 2015 modelled annual data from a grassland field-scale BEF experiment with one year of data and three ecosystem functions recorded; they developed the multivariate DI model for analysing data collected at a 'single' point in time, i.e. the aggregated data collected over one growing season. In that experiment, the aggregated data was the sum of biomass yield measurements over four harvests. The three ecosystem functions measured were: aboveground biomass of the sown species (sown), aboveground biomass of unsown species (weed), and the nitrogen yield in the total aboveground harvested biomass (N). Here, I model the individual harvest level data from these three ecosystem functions for the same

year of data analysed in Dooley *et al.*, 2015, i.e., I analyse data with a multivariate and repeated measures structure, with three responses (ecosystem functions) each measured across four harvests (time points). The data is from the Belgian site involved in the grassland Agrodiversity Study ('Site 1' in Kirwan *et al.*, 2014) which was set up to test whether increases in plant diversity, defined by species richness and evenness, in intensively managed grasslands can improve ecosystem functioning. The experiment includes four species, two grasses (*Lolium perenne* 'G1' and *Phleum pratense* 'G2') and two legumes (*Trifolium pratense* 'L1' and *Trifolium repens* 'L2'), which were sown in 30 plots with a range of species proportions (single species and four-species mixtures with varying proportions) for a total of fifteen community designs, each of which were sown at both high and low seeding density.

Methods

In agronomic settings, nitrogen yield and sown biomass are ecosystem functions for which large positive values are desirable while low values are desirable for weed biomass. In order for the desirability of all functions to match in direction, weed biomass was transformed into a weed suppression index, for which a large positive value is desirable. All functions were also linearly transformed to lie on a percentage scale, where each response value is divided by the mean of the top 10% of values for the corresponding ecosystem function, then multiplied by 100, such that higher values were more desirable, to allow for ease of comparison between functions (see Appendix C).

Four DI models were fit to the multivariate repeated measures data that varied in the interactions included in the fixed effects linear predictor: using the ID, AV, FG (with species being grouped by grasses and legumes), and FULL interaction structures available in the `DImodelsMulti` R package (see Table 1.2) (Kirwan *et al.*, 2009; Moral *et al.*, 2023; Byrne *et al.*, 2024). Prior to fitting the varying interactions structures, joint profile likelihood was used to estimate the θ_k values using the full pairwise interaction structure; each θ_k value

was then tested one by one, by testing each model with a θ_k value estimated against the same model using all θ_k values set equal to one. Once the estimate of θ was determined for each ecosystem function, that value of θ (the estimated value or one, depending on the test result) was used (i.e. assumed known) for the function for the remainder of the model selection process (Vishwakarma *et al.*, 2023). The general (UN) covariance structure was assumed for the ecosystem function responses while all three available structures (UN, CS, and AR(1)) were fit and tested for the harvests, using REML estimation. The model selection process followed was:

1. Fit ID model using each of the covariance structures and the REML method, compare using AICc, taking the structure with the lowest value.
2. Fit the ID and AV models using the ML method and the selected covariance structure and compare using a likelihood ratio test.
3. Estimate values of θ using the FULL DI model, fit using the ML method, testing each individually against the FULL model with all $\theta_k = 1$ using AICc, taking the θ_k value from each model with the lower value.
4. Fit the AV, FG, and FULL models using the new values of θ and the ML method and test hierarchically (Kirwan *et al.*, 2009) using likelihood ratio tests.
5. Refit the chosen model, including the seeding density treatment as an additive factor, and test its inclusion using a likelihood ratio test.
6. Refit the final selected model using REML

See Appendix D for further details and results of each step. In the case that the likelihood ratio test is insignificant (i.e. that the more complex model is not preferred), the model selection steps should continue, e.g. if ID is selected over AV in step 2, the FG and FULL models (with θ estimated) should still be tested. One case where this may be necessary is in a three-species system, where the AV interaction term does not significantly differ from 0 but the full pairwise interaction terms contain one insignificant, one negative, and one

positive coefficient, where the positive and negative coefficients are of equal magnitude, therefore cancelling one another out when averaged.

Results

The structure of the final selected model can be seen in Eq. 3.5, where the chosen interaction structure was the functional group terms (δ) with a seeding density treatment (where $A = 1$ for high density and $A = 0$ for low density) and θ not significantly different from one for sown biomass and nitrogen yield and $\hat{\theta} = 0.433$ for weed suppression. The selected autocorrelation structure was $UN \otimes UN$. The parameters $p_1 - p_4$ represent the initial proportions of species G1, G2, L1, and L2, respectively.

$$y_{kmt} = \sum_{i=1}^{S=4} \beta_{ikt} p_{im} + \delta_{Gkt} (p_{1m} p_{2m})^{\theta_k} + \delta_{Lkt} (p_{3m} p_{4m})^{\theta_k} + \delta_{GLkt} \sum_{\substack{i \in \{1,2\} \\ j \in \{3,4\}}} (p_{im} p_{jm})^{\theta_k} + \alpha_{kt} A_m + \varepsilon_{kmt} \quad (\text{Eq. 3.5})$$

Examination of the residuals from this model, via residual plots, indicated that the multivariate normality assumption for the error term was reasonable.

From the estimated coefficients for the fixed effect terms (Table 3.2), we can see that high seeding density has a positive effect on sown biomass functioning at harvest 1, but generally minimal effects on the three ecosystem functions across harvests. For nitrogen yield, the presence of legume species has a strong positive effect, as expected thanks to their unique ability to fix nitrogen from the atmosphere, and later harvests also generally perform better than earlier ones, as we can see from the increasing coefficient values for most predictors. Sown biomass is improved by the presence of multiple different species, with grass-legume mixtures generally interacting positively across harvests and within functional group interactions being positive, negative, or insignificant at different harvests, this could be due to seasonal changes. Weed suppression performance is heavily dictated by the community design in harvest one, with the presence of grasses greatly increasing this

function's performance, but it performs well regardless of species present in other harvests, with the ID effect coefficient for each species being greater than 90.

From the estimated variance-covariance matrices, Σ , in Table 3.3, we can see that the error term for weed suppression at harvest 1 is the most variable of the ecosystem functions, at a value of 204.11, but that variability lessens over time, while the sown biomass error remains rather variable throughout the first three harvests. We can also note that trade-offs will likely occur between nitrogen yield and the remaining functions due to their negative covariance, in other words, if a community were to be designed with the intention of maximising sown biomass and weed suppression, it would likely see a decrease in nitrogen yield functioning. As to be expected, it appears that the value of the covariances between harvests fall as the time points grow further apart, for example, we see that nitrogen yield at harvest 1 and harvest 2 have a covariance value of 15.64 while the same ecosystem function taken at harvest 1 and harvest 4 have a covariance value of only 0.73. From this variance-covariance matrix, we can not only see that readings of the same ecosystem function covary less as time passes, but that the covariance between ecosystem functions falls as time passes. As we are working with one year of data, these changes may be due to seasonality rather than the passing of time in general. This could be further investigated using the additional years of data available in the study (Kirwan *et al.*, 2014).

We can also visualise the predicted responses for each ecosystem function across harvests for select communities, as shown in Figure 3.2, where we can see that, as we move from earlier to later harvests, sown biomass decreases while weed suppression increases, seemingly irrespective of community composition. The nitrogen yield response appears more variable through harvests and more dependent on the initial community design. In general, mixtures also appear to outperform most monocultures throughout time, across functions.

Table 3.2: The estimated coefficients $\pm$ the standard errors of the selected functional group model for each ecosystem function response at each time point, see Eq. 3.5. Interaction effects (δ) and seeding density effects (α) that are significant at $\alpha = 0.05$ are highlighted in bold and at $\alpha = 0.1$ are underlined.

	Harvest 1			Harvest 2			Harvest 3			Harvest 4		
	N yield (%)	Sown biomass (%)	Weed suppress (%)	N yield (%)	Sown biomass (%)	Weed suppress (%)	N yield (%)	Sown biomass (%)	Weed suppress (%)	N yield (%)	Sown biomass (%)	Weed suppress (%)
β_{G1}	33.6 $\pm$ 4.418	91.5 $\pm$ 6.605	74.8 $\pm$ 9.424	40.3 $\pm$ 2.780	56.9 $\pm$ 6.582	93.6 $\pm$ 3.135	63.2 $\pm$ 1.884	23.4 $\pm$ 7.500	101.6 $\pm$ 3.294	67.9 $\pm$ 3.116	17.0 $\pm$ 1.776	99.7 $\pm$ 0.121
β_{G2}	32.0 $\pm$ 4.418	59.1 $\pm$ 6.605	84.9 $\pm$ 9.424	38.9 $\pm$ 2.780	45.6 $\pm$ 6.582	99.5 $\pm$ 3.135	59.5 $\pm$ 1.884	10.4 $\pm$ 7.500	99.5 $\pm$ 3.294	58.5 $\pm$ 3.116	20.9 $\pm$ 1.776	99.9 $\pm$ 0.121
β_{L1}	79.4 $\pm$ 4.418	46.8 $\pm$ 6.605	30.2 $\pm$ 9.424	73.1 $\pm$ 2.780	81.9 $\pm$ 6.582	99.0 $\pm$ 3.135	68.6 $\pm$ 1.884	66.4 $\pm$ 7.500	99.3 $\pm$ 3.294	104.6 $\pm$ 3.116	22.5 $\pm$ 1.776	100.0 $\pm$ 0.121
β_{L2}	82.5 $\pm$ 4.418	37.6 $\pm$ 6.605	18.8 $\pm$ 9.424	78.3 $\pm$ 2.780	49.1 $\pm$ 6.582	94.5 $\pm$ 3.135	78.4 $\pm$ 1.884	34.5 $\pm$ 7.500	95.2 $\pm$ 3.294	105.0 $\pm$ 3.116	24.2 $\pm$ 1.776	99.8 $\pm$ 0.121
δ_G	27.4 $\pm$ 39.477	7.7 $\pm$ 58.855	-59.7 $\pm$ 66.831	95.1 $\pm$ 24.842	91.2 $\pm$ 58.646	5.9 $\pm$ 22.235	35.9 $\pm$ 16.837	141.0 $\pm$ 66.831	-8.6 $\pm$ 23.360	104.5 $\pm$ 27.838	37.0 $\pm$ 15.821	0.3 $\pm$ 0.859
δ_L	-97.3 $\pm$ 39.477	122.9 $\pm$ 58.855	36.3 $\pm$ 66.831	-48.9 $\pm$ 24.842	67.6 $\pm$ 58.646	6.9 $\pm$ 22.235	-18.0 $\pm$ 16.837	15.1 $\pm$ 66.831	12.7 $\pm$ 23.360	<u>-52.6 $\pm$ 27.838</u>	-10.2 $\pm$ 15.821	0.0 $\pm$ 0.859
δ_{GL}	-43.0 $\pm$ 17.719	117.8 $\pm$ 26.507	31.6 $\pm$ 31.887	29.8 $\pm$ 11.150	36.6 $\pm$ 26.413	0.9 $\pm$ 10.609	5.0 $\pm$ 7.557	69.6 $\pm$ 30.099	-0.3 $\pm$ 11.146	31.2 $\pm$ 12.495	18.7 $\pm$ 7.125	0.1 $\pm$ 0.410
α	-2.8 $\pm$ 2.472	<u>6.2 $\pm$ 3.704</u>	5.5 $\pm$ 5.217	<u>-2.9 $\pm$ 1.555</u>	-1.9 $\pm$ 3.691	-2.8 $\pm$ 1.736	-1.0 $\pm$ 1.054	2.3 $\pm$ 4.206	-1.4 $\pm$ 1.823	-0.1 $\pm$ 1.743	-0.1 $\pm$ 0.996	-0.1 $\pm$ 0.067

Table 3.3: The estimated variance-covariance matrices of the ecosystem function and harvest level responses. (a) shows the Kronecker product formulation of the matrix, while (b) shows the full, multiplied-out version of the matrix. In this matrix, we can see the covariances between different response types at different times. The matrix is symmetrical as $cov(a, b) = cov(b, a)$. The shaded groups of cells represents the within-ecosystem function variability, with the cells along the main diagonal representing the variance of each ecosystem function at each time point. The off-diagonal shaded cells representing the covariance between the same ecosystem function measurements taken at different time points. The non-shaded cells represent the covariances between ecosystem functions at the same or different harvests, for example, we see that sown biomass and nitrogen yield have negative covariance values at each time point, meaning that as one ecosystem function increases, the other tends to decrease.

(a)

	Sown	N	Weed
Sown	88.82	-13.95	82.33
N	-13.95	34.59	-22.04
Weed	82.33	-22.04	155.82

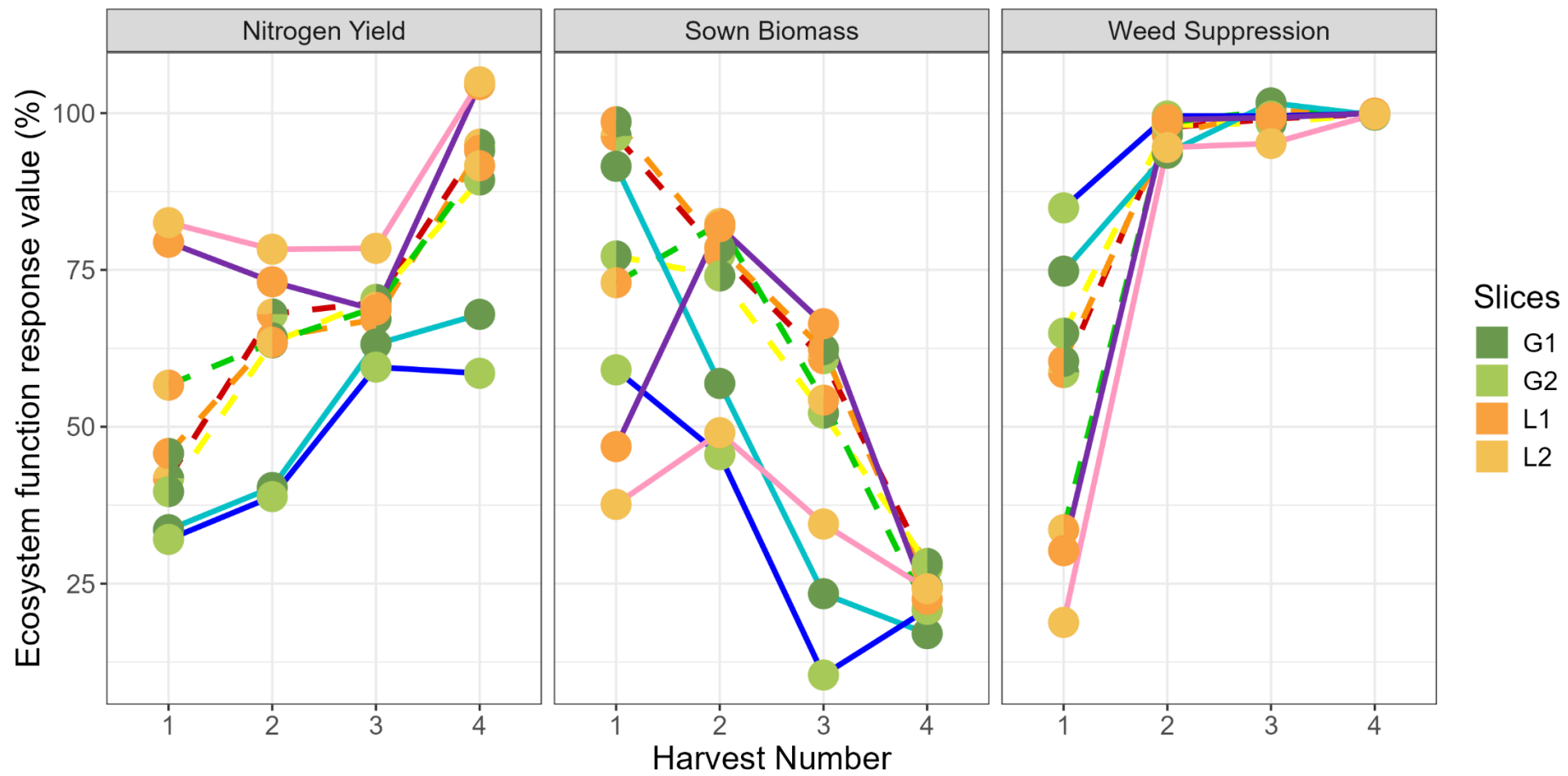
 $\otimes$

Harvest	1	2	3	4
1	89.03	43.67	20.98	0.46
2	43.67	72.78	25.93	4.12
3	20.98	25.93	56.80	3.43
4	0.46	4.12	3.43	4.61

(b)

	Sown:1	Sown:2	Sown:3	Sown:4	N:1	N:2	N:3	N:4	Weed:1	Weed:2	Weed:3	Weed:4
Sown:1	102.90	55.63	34.47	0.63	−17.28	−5.90	−2.17	−0.28	101.42	18.31	10.46	0.03
Sown:2	55.63	102.17	46.96	6.19	−9.34	−10.84	−2.96	−2.73	54.83	33.62	14.25	0.29
Sown:3	34.47	46.96	132.68	6.65	−5.79	−4.98	−8.37	−2.93	33.97	15.45	40.25	0.31
Sown:4	0.63	6.19	6.65	7.44	−0.11	−0.66	−0.42	−3.28	0.62	2.04	2.02	0.35
N:1	−17.28	−9.34	−5.79	−0.11	45.82	15.64	5.76	0.73	−29.03	−5.24	−2.99	−0.01
N:2	−5.90	−10.84	−4.98	−0.66	15.64	18.14	4.96	4.57	−9.91	−6.08	−2.58	−0.05
N:3	−2.17	−2.96	−8.37	−0.42	5.76	4.96	8.33	2.92	−3.65	−1.66	−4.33	−0.03
N:4	−0.28	−2.73	−2.93	−3.28	0.73	4.57	2.92	22.78	−0.47	−1.53	−1.51	−0.26
Weed:1	101.42	54.83	33.97	0.62	−29.03	−9.91	−3.65	−0.47	204.11	36.84	21.05	0.06
Weed:2	18.31	33.62	15.45	2.04	−5.24	−6.08	−1.66	−1.53	36.84	22.59	9.57	0.20
Weed:3	10.46	14.25	40.25	2.02	−2.99	−2.58	−4.33	−1.51	21.05	9.57	24.94	0.19
Weed:4	0.03	0.29	0.31	0.35	−0.01	−0.05	−0.03	−0.26	0.06	0.20	0.19	0.03

(a)



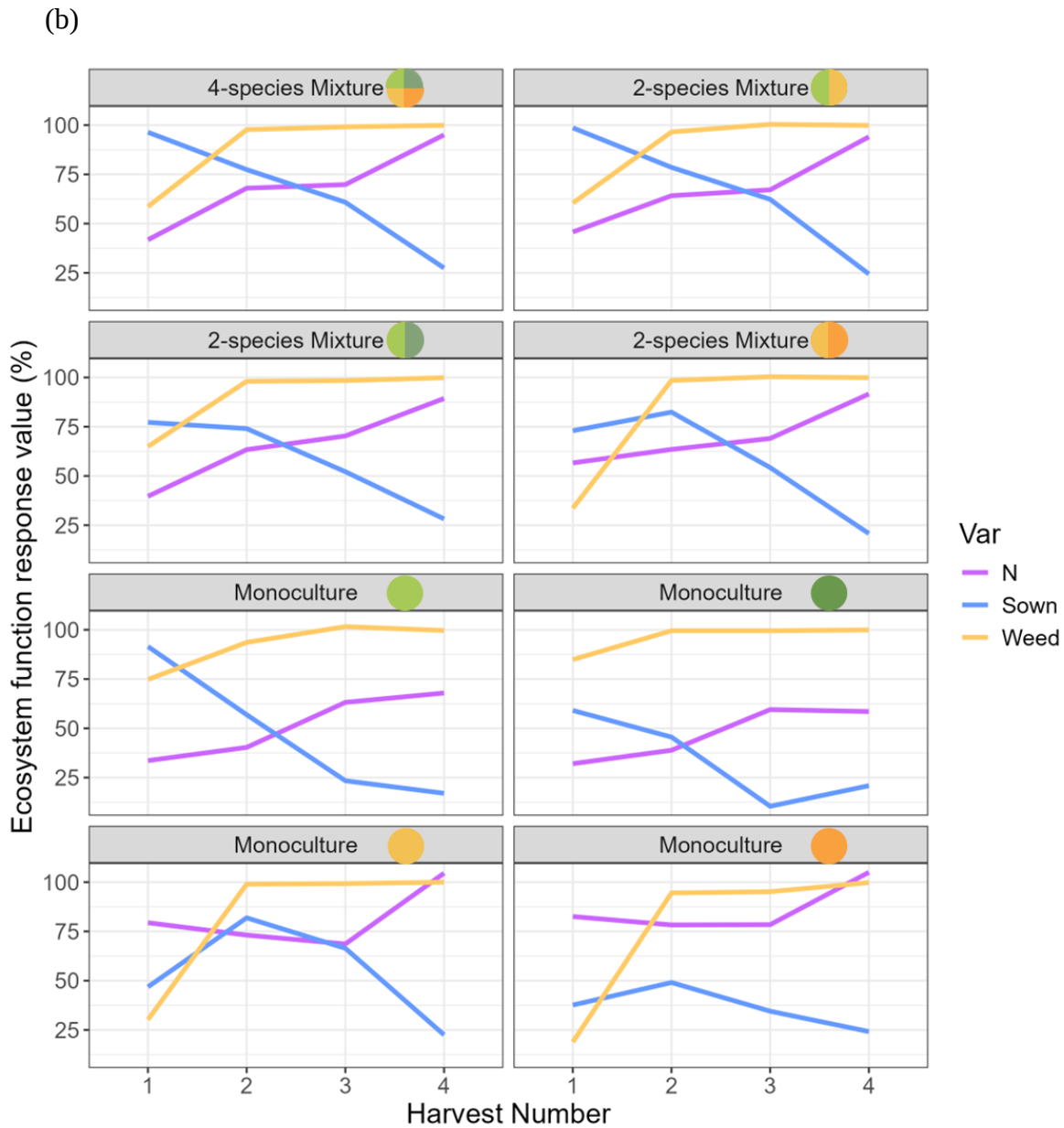


Figure 3.2: Line plots depicting the change in predicted ecosystem function responses for select communities seeded at low density over the four harvests. Figure (a) depicts each ecosystem function response for eight communities in a separate plot, mixtures and monocultures are represented by dashed and solid lines respectively, whilst (b) presents each selected community, and their corresponding ecosystem function responses, in separated plots. Pie-glyphs (Vishwakarma *et al.*, 2024a) are used to represent the proportion of each species sown in the selected community. See Appendix E for figures from the equivalent univariate repeated measures models for comparison.

Figure 3.2 (a) is simplified by producing a summary statistic for the multivariate responses, in this case, the average multifunctionality was used (Grange *et al.*, 2024; Finn *et al.*, 2024). This was done by using the predictions from the model shown in Eq. 3.5 to form the multifunctionality index, averaging across ecosystem function responses for each community design input (see Figure 3.3). Using such multifunctionality indices on the prediction values rather than the raw data can simplify the interpretation process whilst maintaining the complexity, and thus the benefits, of the multivariate repeated measures DI model. From Figure 3.3, we can see that grass monocultures, while stable, are relatively low performing across harvests for this metric, while legume monocultures, which were among the worst performing communities at harvest one thrive at later harvests. We can also see that mixtures appear to outperform the majority of monocultures at most times, regardless of their composition.

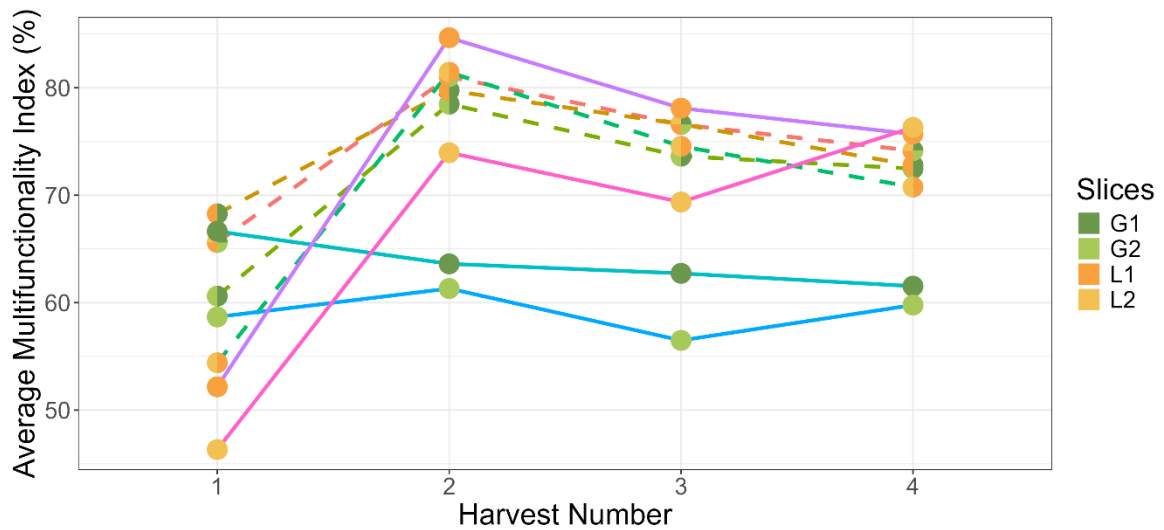


Figure 3.3: A line plot depicting the change in the ‘average’ multifunctionality index produced using the predictions from the multivariate repeated measures model shown in Eq. 3.5 for some selected communities, where pie glyphs are used to represent the proportion of each species sown in each community (Vishwakarma *et al.*, 2024a). Mixtures and monocultures are represented by dashed and solid lines respectively. A low seeding density is used for all predictions.

3.5 Discussion

In this chapter, I have presented a Diversity-Interactions model capable of modelling multivariate repeated measures responses against a gradient of initial species proportions, representing species diversity, whilst accounting for the inherent correlation that exists between readings from the same experimental/observational unit in a study. I have also presented the `DImodelsMulti` R package, which provides a user-friendly way to fit multivariate and/or repeated measures DI models. The package also contains a number of S3 functions, to enable the further analysis of models produced, and example datasets, both real-world and simulated, to aid a user in understanding the package usage.

Future work in this space could involve the expansion of the functionality of the `DImodelsMulti` R package, such as the creation of additional functions for the automatic calculation of multifunctionality indices for multivariate DI model predictions, and/or the inclusion of additional options for the existing `DImulti()` function, e.g., variance-covariance structures. The multivariate repeated measures DI model could also be further expanded to include the estimation of random effects, such as those that arise from multi-site BEF studies wherein species ID effects may differ by site due to variation from unmeasured variables like soil conditions, sunlight availability, or pests. It could also be of interest to investigate the fitting of different multi-response DI models, e.g., for non-normally distributed multivariate responses, as the models presented in this section are currently restricted by the normality assumption of the error term.

The DI model selection process, including the best practice for θ estimation, have been thoroughly studied in univariate models (Kirwan *et al.*, 2009; Connolly *et al.*, 2013; Vishwakarma *et al.*, 2023) and were generalised for the multivariate repeated measures models presented here. However, these generalisations have not yet been formally tested and suggestions made in this chapter are subject to change, dependent on future work.

4 Diversity-Interactions models for modelling the total community response and breakdown by species proportions in biodiversity experiments

Collaborators: Caroline Brophy, Forest Isbell, Rafael A Moral, John Finn, Guylain Grange, Hannah Comiskey

This chapter is formatted as a standalone paper intended for submission to the Journal of Agricultural, Biological and Environmental Statistics (JABES).

CB conceived the model idea; LB led the creation of the model specification and its implementation, supported by CB, HC and RAM; FI conceptualised the usage of the final model with the additive partitioning method; LB led the writing of the manuscript, supported by CB. GG and JF provided the data used in the case study and provided ecological insights on aspects of the model development. All authors will contribute to the writing of the manuscript prior to its submission to JABES.

Abstract

In biodiversity and ecosystem function (BEF) experiments studying grassland systems, typically the number of species and the proportions in which they are sown (species diversity) are manipulated across field-scale plots at establishment and the yield (and/or other responses) of each plot is measured after a period of growth. Models of the BEF relationship usually focus on the relationship between species diversity and the response (or ecosystem function, for example, the total yield biomass of a grassland plot). One such

approach, called Diversity-Interactions (DI) modelling, can directly relate species diversity, quantified using the sown proportion for each species in the pool, to the total yield. However, in a grassland system, the "realised" proportions of each species in the yield may change over time compared to the sown proportions. Understanding how both the total yield and its proportional breakdown by species are affected by sown species diversity (and/or other manipulated predictors), may provide ecological insight in grassland systems, particularly, for example, in agricultural grasslands where the proportional contributions of species or the application of different levels of nitrogen fertiliser may affect the grazing forage quality. In this chapter, I extend the DI modelling approach to jointly model the continuous yield response and the vector of proportions measuring the yield attributed to each species and weeds as a function of the initial sown species proportions.

4.1 Introduction

The biodiversity and ecosystem function (BEF) relationship defines how changes in species diversity affect the goods and services provided by an ecosystem, called ecosystem functions (de Groot, 1992; Costanza *et al.*, 1997). Grassland experiments studying the BEF relationship typically involve a manipulation of the initial sown species diversity across field-scale plots, i.e., the identity, number and/or sowing proportions of species are varied across plots. After a period of establishment and growth, an ecosystem function response is measured on each plot, for example, the plot-level total biomass yield may be harvested and measured. Current literature has shown evidence that grassland ecosystems containing multiple species (mixtures) produce higher yields than monocultures of their comprising species (Cardinale *et al.*, 2011; Finn *et al.*, 2013; Grange *et al.*, 2021; Tian *et al.*, 2021; Andrade *et al.*, 2022). However, the proportions of grassland species that comprise the harvested total yield (the realised proportions) are liable to vary from the sown species proportions, due to environmental effects, such as temperature or resource availability,

and/or between-species effects, such as interspecific competition, niche partitioning, or the invasion of undesired species/weeds into the ecosystem (Teasdale, 1998; Loreau & Hector, 2001; Tracy & Sanderson, 2004; Verma *et al.*, 2014; Liu *et al.*, 2018; Qianwen *et al.*, 2022; Surault *et al.*, 2024). These effects could result in the loss of important species or a general decline in species diversity over time, potentially vitiating the benefits of establishing a mixture over a monoculture.

While the effect of sown species diversity on the total yield is almost always of interest in these types of studies (e.g., Kirwan *et al.*, 2009; Finn *et al.*, 2013; Dooley *et al.*, 2015; Grange *et al.*, 2021; Suter *et al.*, 2021), there are situations where the proportional breakdown of the total yield by species may also be of interest. For example, in productive agronomic grasslands, maintaining some species above a certain threshold may benefit some agronomic outcomes and understanding what sown proportion of that species is required to achieve the threshold would be beneficial. Therefore, a model which can be used to plan the construction of an ecosystem, based on both its functional output and its realised species proportions is potentially of crucial importance for practical management (Hunter Jr, 2002). In this chapter, I develop a regression-based modelling approach with a multivariate response that includes the total yield (or other ecosystem function) and the proportions of the total yield for each component species and weeds that is modelled as a function of the sown species proportions and/or other experimental treatments.

The model developed in this chapter is motivated by a grassland biodiversity field experiment that investigated the impact of manipulating the diversity of six sown species on total yield (Grange *et al.*, 2021). The six species could be categorised according to three functional groups, grasses, legumes, and herbs that each have different ‘functions’ within productive grasslands. Grange *et al.*, (2021) showed that increasing sown species diversity improved the total yield over two years. Here, I apply the newly developed approach to the first year of data from Grange *et al.*, (2021) to address the question: how does sown species

diversity jointly affect the total yield *and* the proportional breakdown of the yield into grasses, legumes, herbs, and weeds?

There are other analytical approaches that have modelled species proportion dynamics in BEF studies. These methods generally involve the computation of some metric to describe species composition or proportion changes rather than directly modelling the resultant species proportions (*relative growth rates*: Brophy *et al.*, 2017; *CAFE approach*: Bannar-Martin *et al.*, 2018; Ladouceur *et al.*, 2022; *Price equation*: Ladouceur *et al.*, 2022). However, while these metrics are highly valuable for assessing overall changes to ecosystem functioning, they do not allow for the direct modelling of the realised ecosystem species proportions and therefore incur a loss of dimensionality. This loss may cause conflating results, where two sown communities are predicted to have the same ecosystem function response but achieved this result through different underlying biological mechanisms, which could be separated and inferred from the realised/achieved species composition. For example, two plots with similar predicted yield responses could have realised proportions of all species approximately equal or could have one species highly dominating the yield; modelling only yield (and not its breakdown by species proportions) may not uncover the differences driving those two similar yield scenarios.

In this chapter, I will describe the DI modelling framework and how it can be used with the new model proposed. While I describe the approach in the context of grassland BEF experiments and will demonstrate this new DI model in a case study using the data from the grassland biodiversity field experiment previously mentioned (Grange *et al.*, 2021), the approach is applicable to many other BEF studies, such as observational studies or effectiveness tests for treatments (wherein species diversity is not manipulated across units). Lastly, I will demonstrate how the use of the new model with the DI framework can be used in conjunction with the additive partitioning method (Loreau & Hector, 2001; Gering *et al.*, 2003).

4.2 Proposed model

In this chapter, I propose a novel modelling method for the joint prediction of the total yield produced (y_D) and its breakdown into the realised species proportions ($y_1, y_2, \dots, y_C$) from an ecosystem. Data collected from a grassland field plot experiment is suitable for this approach if the total biomass yield (y_D) has been collected on each plot some length of time after the establishment of the plots, *and* if the total yield breakdown by the realised species proportions ($y_1, y_2, \dots, y_C$), has also been measured at that time. The total yield response y_D is continuous and assumed normally distributed, while $y_1, y_2, \dots, y_C$ are proportions. The approach developed models this multivariate response vector ($y_1, y_2, \dots, y_C, y_D$) as a function of the experimental manipulations in a multivariate regression approach. Suitable data can arise from either an experiment that did or did not include species diversity manipulations in the design. For example, suppose an experiment where three species are sown in equal proportions in all plots, but the level of nitrogen fertiliser application has been manipulated across the plots. In this case, the response vector would include y_1, y_2, y_3, y_4, y_D (where y_1, y_2, y_3 are the proportions of sown species responses, y_4 is for the proportion of weeds response and y_D is the total yield response and $D = 5$), and only the treatment nitrogen fertiliser would appear in the linear predictor of the regression model. Alternatively, if sown species proportions have been manipulated across the plots (solely or in addition to other experimental manipulations), the approach can be embedded in the DI approach to model the relationship between sown species proportions and the multivariate vector of yield and realised species proportions. This methodology will therefore allow for the estimation of the effect of any set of predictors on both a total community response, such as biomass, and the set of realised species proportions.

In the remainder of this section, I will first give a brief overview of DI models and then describe the new modelling approach in the context of DI models, i.e., assuming that there have been experimental manipulations of species diversity. p is used to denote a sown

species proportion and y to denote a realised species proportion (i.e., the proportion of the total biomass yield attributed to a given species).

Overview of DI models

The DI modelling framework (Kirwan *et al.*, 2009; Connolly *et al.*, 2013; Moral *et al.*, 2023) is a regression-based approach to quantifying the BEF relationship. It models an ecosystem function such as yield as a function of the sown (or initial) proportions of species and their interactions. It has been commonly applied in experimental grassland studies, particularly for those studying the effects of species diversity on biomass and weed suppression (Kirwan *et al.*, 2009; Finn *et al.*, 2013; Dooley *et al.*, 2015; Grange *et al.*, 2021; Suter *et al.*, 2021; McDonnell *et al.*, 2023; Argens *et al.*, 2023). An example of a univariate DI model, which could be used to predict the total biomass, can be seen in Eq. 4.1 where the univariate response y_m represents the recorded value of an ecosystem function from an experimental unit m , numbered $1, 2, \dots, M$. This response is modelled as the summation of the species' identity (ID) effects, which, for some species i from a pool of S total species, is expressed as the coefficient β_i scaled by the species' initial proportion p_i , and the species' interactions effects, which can be formulated in a number of different ways, see Table 1.2, in this example the full pairwise (FULL) structure is given, where it is assumed that each pair of species (i and j) interact in a unique way with strength δ , scaled by the product of the proportions of the species. The non-linear parameter θ may be applied as a power to the pairwise products of species in order to change the shape of the BEF relationship. Additional structural effects, such as treatments, can be easily included in the model through the term A with coefficient α . The error term of the model, ε , is assumed to follow a normal distribution with mean zero and variance σ^2 .

$$y_m = \sum_{i=1}^S \beta_i p_{im} + \sum_{\substack{i,j=1 \\ i < j}}^S \delta_{ij} (p_{im} p_{jm})^\theta + \alpha A_m + \varepsilon_m, \quad \varepsilon_m \sim N(0, \sigma^2) \quad (\text{Eq. 4.1})$$

At present, DI models have already been expanded to allow for the modelling of multiple responses, taken from the same experimental unit, simultaneously whilst accounting for the inherent correlations between them (Frankow-Lindberg *et al.*, 2009; Finn *et al.*, 2013; Dooley *et al.*, 2015). However, these approaches assume that all responses are continuous/unbound and are drawn from a multivariate normal distribution.

New approach for modelling the multivariate response vector $(y_1, y_2, \dots, y_C, y_D)$

The proposed model relates the set of responses, $(y_1, y_2, \dots, y_C, y_D)$ for each experimental unit m (for $m = 1, \dots, M$), to a function of the sown species proportions and, where applicable, their interaction effects along with any additional experimental treatments. In the example model given by Eq. 4.2, β_{di} represents the ID effect of species i for a response d (for $d = 1, \dots, D$), which is scaled by its initial proportion, p_i , from experimental unit m . The interaction term δ_{AVd} represents the interaction effect of all species sown in experimental unit m for the d^{th} response and is included as an example interaction structure, see Table 1.2 for alternatives. Interaction effects are only explicitly included for the total biomass (y_D) and the proportion of weed responses (which is $y_C = y_{D-1}$ if weeds have been recorded as a single proportion in the breakdown of the total yield). Interaction effects are not included for modelling the realised proportions from species sown in the experiment as those interaction effects are already captured by the identity effects; the ID effect of each species represents the effect of the sown species' presence on the response species, i.e. their interaction. For responses where interactions are included, the non-linear parameter θ may vary in value between responses, allowing the shape of the BEF relationship to change for different responses. Any additional treatment effects may also be included through A .

$$y_{dm} = f(\omega_{dm}) =$$

$$f \left(\begin{cases} \sum_{i=1}^S \beta_{di} p_{im} + \alpha_d A_m, & \text{if } d \text{ represents a species sown in the experiment} \\ \sum_{i=1}^S \beta_{di} p_{im} + \delta_{AVd} \sum_{\substack{i,j=1 \\ i < j}}^S (p_{im} p_{jm})^{\theta_d} + \alpha_d A_m, & \text{otherwise} \end{cases} \right)$$

$$+ \varepsilon_{dm}, \quad \varepsilon_m \sim MVN(0, \Sigma^*) \quad (\text{Eq. 4.2})$$

The residual term ε_{dm} is assumed to follow a multivariate normal distribution with mean 0 and variances Σ^* , where Σ^* is a $DM \times DM$ block diagonal matrix (Eq. 4.3) with $D \times D$ matrix Σ along the main diagonal and zeroes elsewhere.

$$\Sigma^* = \begin{bmatrix} \Sigma_1 & & 0 \\ & \ddots & \\ 0 & & \Sigma_M \end{bmatrix} \quad (\text{Eq. 4.3})$$

The dimensions of Σ_m will depend on the number of proportions that the total yield has been divided into. For the example where we have the responses: proportions of three species, a weeds proportion and a total yield, Σ_m follows the form shown in Eq. 4.4, where σ_1^2, σ_2^2 and σ_3^2 represent the variance of the proportions y_1, y_2 and y_3 respectively, σ_4^2 represents the variance of proportion of weeds (y_4), and σ_5^2 represents the variance of the total yield (y_D); while covariances are on the off-diagonal.

$$\Sigma = \begin{bmatrix} \sigma_1^2 & \sigma_{12} & \sigma_{13} & \sigma_{14} & \sigma_{15} \\ \sigma_{12} & \sigma_2^2 & \sigma_{23} & \sigma_{24} & \sigma_{25} \\ \sigma_{13} & \sigma_{23} & \sigma_3^2 & \sigma_{34} & \sigma_{35} \\ \sigma_{14} & \sigma_{24} & \sigma_{34} & \sigma_4^2 & \sigma_{45} \\ \sigma_{15} & \sigma_{25} & \sigma_{35} & \sigma_{45} & \sigma_5^2 \end{bmatrix} \quad (\text{Eq. 4.4})$$

The regression parameters in the linear predictor are estimated by maximising the log-likelihood (Eq. 4.5). The first line of this log-likelihood equation comes from the multivariate normal distribution, where M represents the number of experimental units, D represents the number of responses taken from each experimental unit m , and μ is the vector of expected values for the responses, $y = (y_{1m}, y_{2m}, \dots, y_{Dm})$. This line accounts for the covariance between the responses and imposes a multivariate normal assumption on the

error. The subsequent lines come from zero-adjusted regression models (Tsagris & Stewart, 2018), which extends the multivariate beta (Dirichlet) distribution to allow for zero-valued responses. The dataset is divided into two, rows with any zero-valued compositional response denoted by 0 and those with no zeroes denoted by * , therefore M^* would refer to the number of experimental units which had no zero-value responses, e.g. a plot m_1 which observed 0% of any species would be sorted into the 0 matrix, while a plot m_2 which observed some % greater than zero for all species would be sorted into the * matrix. The second line of the equation deals solely with the rows of data which contain no zero responses and is the log-likelihood for the Dirichlet distribution. The third line adapts the usual Dirichlet log-likelihood to allow for zero proportions, trading the problematic $\log(\widehat{y}_n)$ term, which is undefined at zero, for a function L . The final line in the equation is a constant term, added for the marginal probability that an observation would fall into each of the separated matrices (denoted using * and 0). The nonlinear parameters θ_d are estimated using a joint profile likelihood optimisation (Brent, 1973; Moral *et al.*, 2023).

$$\begin{aligned}
\ell = & -\frac{MD}{2}\log(2\pi) - \frac{M}{2}\log(\det(\Sigma)) - \frac{1}{2}\sum_{m=1}^M (y_m - \mu_m)^T \Sigma^{-1} (y_m - \mu_m) + \\
& (M^*)\log \Gamma(W^*) - (M^*)\sum_{n=1}^C \log \Gamma(\omega_n^*) + (M^*)\sum_{n=1}^C (\omega_n^* - 1)\log(\widehat{y}_n^*) + \\
& (M^0)\log \Gamma(W^0) - (M^0)\sum_{\substack{n=1 \\ \omega_n \neq 0}}^C \log \Gamma(\omega_n^0) + (M^0)\sum_{n=1}^C (\omega_n^0 - 1)L + \\
& (M^*)\log\left(\frac{M^*}{M}\right) + M^0\log\left(\frac{M^0}{M}\right)
\end{aligned} \tag{Eq. 4.5}$$

Where $L = \log(\widehat{y}_n^0)$ if $\widehat{y}_n^0 \neq 0$ and $L = 0$ otherwise.

The standard errors of the fixed effect coefficients may be calculated using the inverse of the Hessian matrix produced at the maximum likelihood estimation (MLE) point,

which may also be referred to as the observed Fisher's information matrix, as it acts as an approximation of the fixed effects variance-covariance matrix (Tsagris & Stewart, 2018). However, as the likelihood is disjointed, i.e., not continuous due to the splitting of the dataset into rows with zero-valued responses and those without, it is possible that the Hessian matrix may not be computable or may not follow the correct format for a variance-covariance matrix (symmetrical and positive definite) at the point of MLE.

The untransformed response of the model, ω_{dm} (Eq. 4.2), is referred to as the d^{th} response's shape parameter from experimental unit m . This parameter, after estimation, can be transformed to the realised proportions (y_c , for $c = 1, \dots, C$) and the total ecosystem function (y_D), using a function $f()$, shown in Eq. 4.6.

$$\hat{y}_{dm} = f(\hat{\omega}_{dm}) = \begin{cases} \frac{e^{\hat{\omega}_{dm}}}{\hat{W}}, & \text{if } d \in \{1, \dots, C\}, \\ \hat{\omega}_{dm}, & \text{otherwise} \end{cases}, \quad \hat{W} = \sum_{c=1}^C e^{\hat{\omega}_c} \quad (\text{Eq. 4.6})$$

A set of compositional responses following a Dirichlet distribution is said to be neutral (Connor & Mosimann, 1969), therefore the relationships between any subset of responses are unaffected by the removal of any number of elements. In the case that it is known that a species did not appear in the realised biomass, such as in the event that the species was not sown in a plot and therefore would be categorised as a weed if it did appear in the plot, the estimated realised proportion of said species can be set to zero, and the proportions of the remaining responses may be recalculated/proportionally redistributed.

Models that vary in their linear predictors may be compared using information criteria such as Akaike's Information Criteria (AIC), although the corrected version is preferred for models with large numbers of parameters (Hurvich & Tsai, 1989; McQuarrie & Tsai, 1998). Likelihood ratio tests are also applicable in the case that the models to be compared are nested, such is the case with DI models (Kirwan *et al.*, 2009; Moral *et al.*, 2023)

4.3 Case study

Data

The data used in this section originates from a grassland field experiment carried out in Co. Wexford, Ireland (Grange *et al.*, 2021). There were 43 plots established in April of 2017, seeded using a total 19 distinct plant community designs. There were six species used in the experiment, which could be grouped into three functional groups; ‘Grasses’: *Lolium perenne* and *Phleum pratense*, ‘Legumes’: *Trifolium pratense* and *Trifolium repens*, and ‘Herbs’: *Cichorium intybus* and *Plantago lanceolata*. There were monocultures for each of the six species and 13 unique mixtures, with richness levels varying from two to six species. The sown species proportions followed a simplex design using seeding weight, based on the advised weight for each species in monoculture. The species were sown both at equi-proportional and asymmetric levels, with monocultures at 100%; mixtures of two species from the same functional group at 50% of each; four-species mixes with 25% of each, where species came from two of the three functional groups; a single dominant species at 60% with four other species, each at 10%, from the remaining two functional groups; and a six species equi-proportional mixture of all species. These communities were randomly assigned to the 43 plots in the field, with the number of replicates for each community pre-determined. See Table S1 from the supplementary information of Grange *et al.*, 2021 for an overview of the design and replicates per community type; (note that the study manipulated water supply, here only the rainfed/control plots are studied, omitting the drought plots). At seven harvests over the course of one growing season in 2018, the dry matter yield of each plot (t/ha) was measured, along with a visual assessment of the proportional species breakdown in the plot, which could be summarised into the three sown functional groups and ‘weeds’. If a species was not sown in the plot it was classified as weeds (even if one of the six species in the experimental design). The individual harvest measurements were combined to give annual

values for each plot for analysis here, therefore the set of responses to be used here includes the total annual yield (the summation of each harvest's yield) and its breakdown into the three sown functional groups (grasses, legumes, and herbs) and weeds.

Methods

The multivariate vector of responses for the m^{th} plot is: $(y_{1m}, y_{2m}, y_{3m}, y_{4m}, y_{Dm})$, where y_1, y_2, y_3 and y_4 are the realised proportions of total yield (y_D) attributed to grasses, legumes, herbs and weeds respectively. A series of DI models were fit to the multivariate response data, using the ID, AV, and FG interaction structures (see Table 1.2 for further information on these structures) and compared across the model fits, and use p_1 to p_6 to denote the sown proportion of the six species. The nonlinear parameters θ_d were estimated and tested using the FG interaction structure, as the most complex interaction structure fit, then held constant throughout the analysis (Vishwakarma *et al.*, 2023). A likelihood ratio test was performed to select between models with differing interaction structures, while the corrected version of Akaike's Information Criteria (AICc) was used to select between models with differing values of the non-linear parameter θ , the results for each step are in Appendix G. All models were fit using the R programming language (R Core Team, 2022).

Results

The final selected model (Eq. 4.7) used the functional group (FG) interaction structure for the total dry matter yield and the proportion of weeds responses, and only identity effects for proportions of grasses, legumes, and herbs responses. This structure assumes that species within the same functional group will interact in one way, while pairs of species from different functional groups interact in another. The interaction term δ_{qq} is referred to as the within-FG interaction for some functional group q , while δ_{qr} is the between-FG interaction term for some functional groups q and r (I assign the functional groups grass, legume, and herb the values 1, 2, and 3).

$$\begin{aligned}
y_{dm} &= f(\omega_{dm}) = \\
f &\left(\begin{cases} \sum_{i=1}^6 \beta_{di} p_{im}, & \text{if } d \text{ represents a species sown in the experiment} \\ \sum_{i=1}^6 \beta_{di} p_{im} + \sum_{q=1}^3 \delta_{qq} \sum_{\substack{i,j \in FG_q \\ i < j}} (p_{im} p_{jm})^{\theta_d} + \\ \sum_{1 \leq q < r \leq 3} \delta_{qr} \sum_{i \in FG_q} \sum_{j \in FG_r} (p_{im} p_{jm})^{\theta_d}, & \text{otherwise} \end{cases} \right) \\
&\quad + \varepsilon_{dm}
\end{aligned} \tag{Eq. 4.7}$$

The nonlinear parameter was estimated for the interactions, at values 0.01 and 1.50 for the yield and weeds respectively, however, these values are at the bounds of the estimation method and likely indicates lack of fit. After model selection, it was found that the inclusion of θ different to 1 for the total yield was not beneficial to the model and so was set equal to one ($\theta_D = 1$), while the estimated value of θ for the realised proportions of weeds was found to benefit the model's estimation ($\theta_4 = 1.5$). As lack of fit with these θ values is likely, it would probably be appropriate to exclude the non-linear parameter from the model, i.e. set all $\theta = 1$. It may instead be of interest to extend the range of θ past 1.5 to test if this border value actually relates to the maximum log-likelihood or if the increasing trend continues past this value.

The coefficients for the total yield (y_D) are interpreted in the same way as in a univariate DI model, as there is no transformation applied to the linear predictors for this response, see Eq. 4.6. From the estimated coefficients of this model (Table 4.1), it was found that between-FG interactions had strong positive effects on the total yield (biomass), y_D , which supports the theory of niche partitioning, by which, rather than directly competing with one another, species instead adapt to fill niches, having different resource requirements for growth, and therefore non-similar species (such as those from differing functional groups), may thrive together in the same ecosystem. Further supporting this interpretation is the negative interactions within-grasses and within-herbs, where these species, which share a functional group, may compete with one another for the same resources, potentially hurting

both of their growths. However, this is not the case for the within-legumes interaction, which is relatively strong and positive, at a value of 2.43.

The parameter estimates for the proportion responses in Table 4.1 are not directly interpretable (due to the back-transformation required for the proportion scale), however, we can see that the ID effect of the species from the sown functional groups are generally higher for their own corresponding proportional response (highlighted in light grey cells table); it would be expected that sowing a species in a plot would increase the proportion of realised yield for its corresponding functional group. On the flip side, the ID effect of a species on the realised proportion of a functional group to which it doesn't belong is generally negative, therefore suppressing growth, or relatively small, therefore not contributing greatly to the growth of that functional group's species. The strongest of these interactions is the effect of the herb *Ci* on the realised proportion of grasses, at a value of -5.10. Exceptions to this general case, within this study, appear to be the relatively strong positive effect of the legume *Tp* on grasses, likely due to their well-known synergistic relationship wherein the nitrogen-fixing capabilities of legumes creates additional nitrogen within the plot, which facilitates the growth of grasses, and the effect of the legume *Tr* on herbs. The groups of unsown species, *Weed*, was found to be greatly suppressed by between-FG interactions and marginally so by within-grass and within-herb interactions. Interestingly, the presence of legume species appears to facilitate the growth of unsown species, with both legume species (*Tp* & *Tr*) having relatively large ID effects and a positive, albeit smaller, within-functional group effect. This may be due to their nitrogen-fixing capabilities, which we credit with their facilitation of boosted grass realised proportions, making additional nutrients available for other species, therefore, when no other functional group is sown in mixture with legumes, this invites the invasion of weeds. This may also explain the positive within-legume effect on the total yield, with the additional biomass gathered actually being attributed to undesirable weed species which could not be captured by the univariate yield model.

The standard errors of the model's fixed effect coefficients could not be estimated at this point of MLE, which, as previously discussed, is likely due to the disjoint nature of the likelihood function, with this dataset having many zero-valued responses. Profiling the likelihood of the model or using the bootstrapping estimation method may be a suitable next step in this analysis to estimate these errors.

The variance-covariance matrix of the responses, Table 4.2, shows that the weed proportion is the most variable of the compositional responses, at a value of 3.25, this could be attributed to the fact that this 'functional group' is possibly comprised of numerous different species, each of which may interact very differently with the sown species. The proportion of weeds is also most highly correlated with the legume proportion response, with a covariance value of 0.21, pointing to a relationship between the two wherein, as the realised proportion of legumes increases, generally driven by an increase in the proportion of sown legumes, Table 4.1, the realised proportion of weeds will increase at a relative rate of 0.21. The realised proportions of herbs and grasses have an antagonistic relationship, with a negative covariance value of -0.34, therefore, as the realised proportion of one species increases, the other will decrease. This is an important relationship to note in conjunction with the positive covariance between the total yield and the realised herb proportion (0.60) as, despite the positive interaction effect between sown herbs and grasses on the total yield response (4.17, Table 4.1), the negative effect of sown herbs on realised grass proportions (Ci : -5.10, Pl : -0.82, Table 4.1) and this negative covariance indicate that high total yields will be comprised of higher realised proportions of herbs than grasses, possibly even in the case that grasses are initially sown at a higher proportion than herbs (as the negative effect of sown grass proportions on realised herb proportions are weaker than those of sown herb on realised grasses, Table 4.1).

Table 4.1: The table of fixed effect coefficients for the final selected model shown in Eq. 4.7.

The coefficients for the total yield are directly interpretable, as with any univariate DI model, however, for the proportional responses, a back-transformation is required for scaling (Eq. 4.6). In the case of the untransformed proportional-response coefficients, they can be inferred from in relation to one another, e.g. a relatively large positive value will indicate a relatively large positive effect of the predictor while a small or negative value will have minimal or negative (suppression) effect. The ID effect of a species on a functional group that it does not belong to can be interpreted as their interaction effect, e.g. the ID effect of the legume *Tr* on Weed proportion has a relatively large and positive value of 1.60, indicating that the presence of this species has a positive interaction with the unsown species and may facilitate their growth, whilst the herb *Pl* has a negative effect, indicating a negative interaction with weeds and therefore may suppress weed invasion through its presence.

$\theta_D = 1.0$ $\theta_4 = 1.5$		Responses				
		Grass proportion	Legume proportion	Herb proportion	Weed proportion	Total yield (t/ha)
Predictors	$\widehat{\beta}_1$ (Lp)	2.76	-1.11	-0.33	-0.31	9.06
	$\widehat{\beta}_2$ (Pp)	2.70	-0.23	-0.54	1.25	10.91
	$\widehat{\beta}_3$ (Tp)	1.39	2.70	0.69	2.02	11.38
	$\widehat{\beta}_4$ (Tr)	-0.51	2.22	1.11	1.60	10.26
	$\widehat{\beta}_5$ (Ci)	-5.10	0.26	2.42	1.10	8.40
	$\widehat{\beta}_6$ (Pl)	-0.82	-3.75	2.35	-1.14	11.07
	$\widehat{\delta}_{11}$ (G)				-1.80	-1.23
	$\widehat{\delta}_{22}$ (L)				0.86	2.43
	$\widehat{\delta}_{33}$ (H)				-5.14	-6.43
	$\widehat{\delta}_{12}$ (GL)				-31.38	6.16
	$\widehat{\delta}_{13}$ (GH)				-29.07	4.17
	$\widehat{\delta}_{23}$ (LH)				-36.02	5.34

Table 4.2: The variance-covariance matrix for the responses of the final selected model shown in Eq. 4.7. The variance of each response lies along the main diagonal of the matrix (see shaded cells) whilst the covariance between pairs of responses lie on the off-diagonal.

	Grass proportion	Legume proportion	Herb proportion	Weed proportion	Total yield
Grass proportion	1.65	-0.01	-0.34	0.01	0.25
Legume proportion	-0.01	0.80	-0.01	0.21	0.10
Herb proportion	-0.34	-0.01	0.91	0.04	0.60
Weed proportion	0.01	0.21	0.04	3.25	0.05
Total yield	0.25	0.10	0.60	0.05	0.80

This model can be used to predict for specific sown species communities, both those included in the original experimental design and those outside of it, given there was sufficient coverage around the design space. From Figure 4.1, we can see that mixtures with higher numbers of functional groups tend to outperform communities of lower functional group richness levels, regardless of composition. We can also see that the inclusion of herbs with grasses reduces the presence of weeds while legumes facilitate weed invasion, although this is mitigated in mixtures. The overall functioning of an ecosystem can be defined in a number of different ways, greatly depending on the goal of the stakeholders. The optimal performing community for two measures of overall ecosystem functioning are presented in Figure 4.1: (1) the least change in species composition (between sown and realised), which was a four-species mix, although heavily dominated by two species and (2) the highest sown species biomass produced, i.e. the total yield not counting weeds, which was a three-species mix, see Appendix H for further details. We can see that both estimated ‘optimal’ communities contain all three functional groups, suggesting that species interactions, specifically between functional groups, are valuable for not only the stability of an ecosystem’s species composition and their proportions, but also for the maximisation of its

biomass production. Interestingly, the community which experienced the least shift in species proportions (between sown and realised) contains one grass species (50.9% *Lp*) and one herb species (47.2% *Pl*), in addition to a very small proportion of legumes (1.1% *Tp*, 0.8% *Tr*) despite the previously noted antagonistic relationship between herbs and grasses. This is likely due to the fact that the two species actually have very similarly negative relationships with one another, which effectively balance one another out when they are sown in nearly equal proportions. The presence of all three functional groups also promotes a resistance to weeds, as previously noted, with the ‘weed boosting’ effect of legumes minimised via their rare (non-dominating) inclusion.

From Figure 4.2, we can clearly see that as the sown proportion of sown legumes increases, so does the presence of weeds in the resulting biomass. It is also worthy to note that despite being sown in equal proportions, herbs appear in larger proportions in the realised communities than grasses. However, as the proportions of sown legumes increases, the realised proportions of herbs decreases at a faster rate than the grasses, insinuating that either legumes and herbs have a stronger antagonistic relationship than legumes and grasses, that the benefits of sowing legumes on the realised proportions of grasses partially outweighs the negative effects of decreasing their sown proportion, or possibly both.

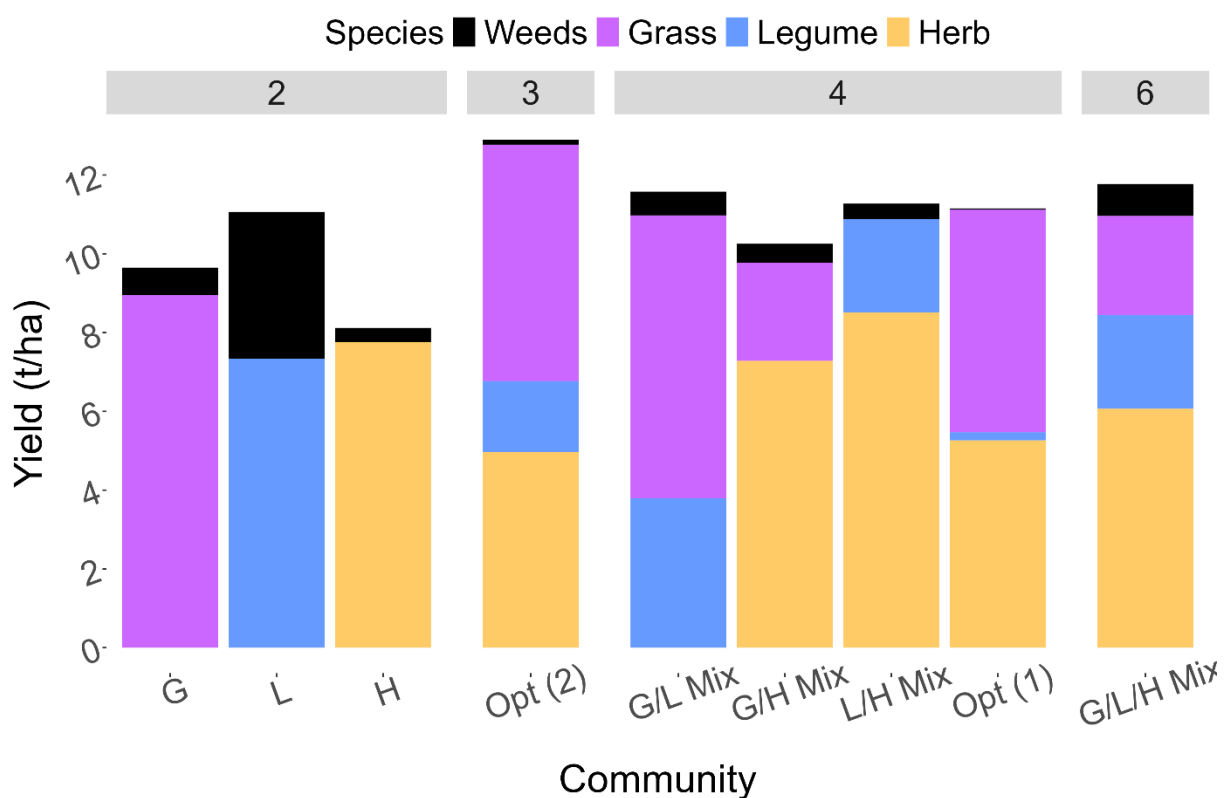


Figure 4.1: A stacked bar chart showcasing the predicted yield (biomass) for a variety of initial sown species proportions, using the model given in Eq. 4.7. The total yield and its proportional breakdown were predicted using the DI model, which were then multiplied to find the predicted yield of each functional group, the breakdown of which is shown here via the different coloured segments of each bar. The sown communities represented are: 50/50 mix of each species within functional groups (G, L, and H), 25/25/25/25 mix of each species between pairs of functional groups (G/L, G/H and L/H mix), an equal mix of all six species (G/L/H mix), and the best performing community for two measures, ‘Opt (1)’, which represents the least change in species proportions (50.9% *Lp*, 0.0% *Pp*, 1.1% *Tp*, 0.8% *Tr*, 0.0% *Ci*, and 47.2% *Pl*) and ‘Opt (2)’, which produced the highest sown species yield (i.e. excluding weeds) produced (sown proportions: 0.0% *Lp*, 27.9% *Pp*, 40.6% *Tp*, 0.0% *Tr*, 0.0% *Ci*, and 31.5% *Pl*).

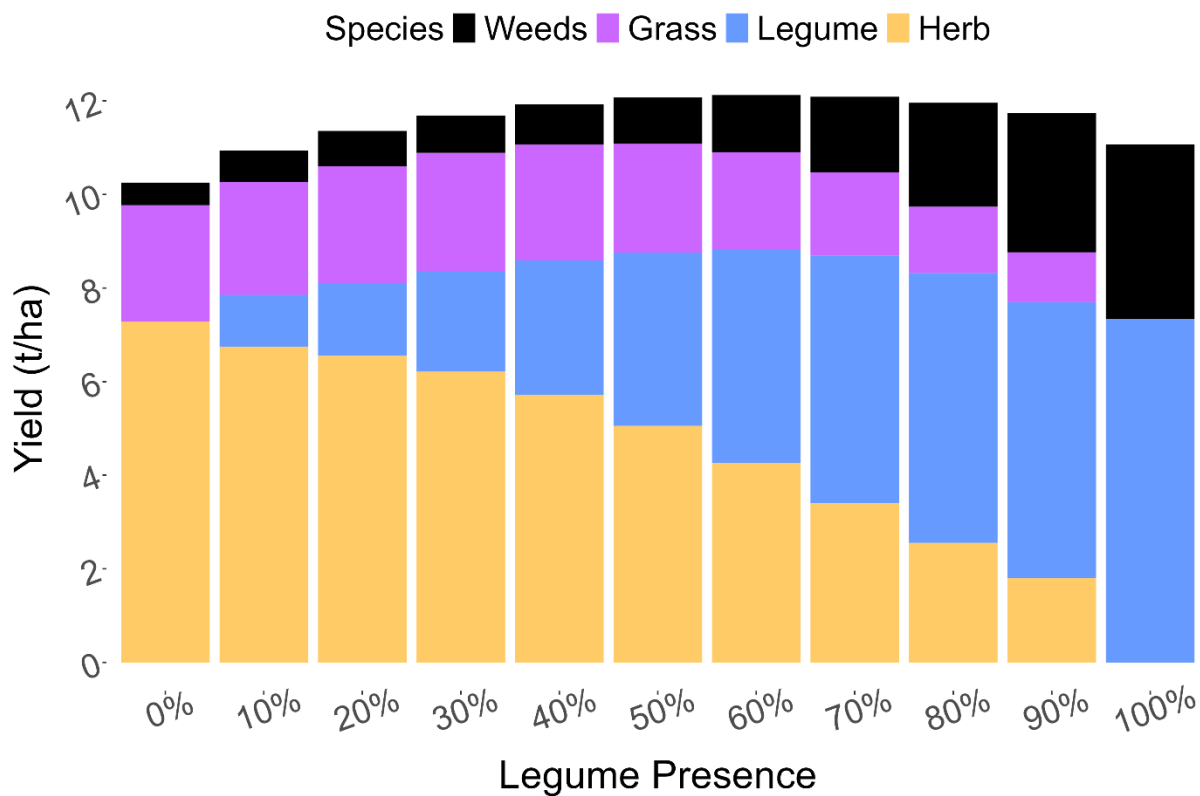


Figure 4.2: A stacked barchart showing the change in the biomass, and its breakdown into the functional groups, as the percentage of the two legume species increase equally. All other species remain equi-proportional, e.g. each sown at 10% when legumes represent 60% (30% Tp and 30% Tr).

4.4 Additive partitioning method

An interest in the effects of the change in species proportions is not novel in BEF studies. The popular additive partitioning method (Loreau & Hector, 2001) utilises this change in relative abundances over time as a lens through which to better understand the mechanisms underlying the relationship. Using additive partitioning, an analysis can break down the net biodiversity effect into two components:

- the complementarity effect, theorised to be responsible for increased total resource use through resource partitioning or positive interactions between species, resulting in better quality and quantity of goods and services provided by an ecosystem, and

- the selection effect which refers to how dominance in coverage of an ecosystem by species with certain traits affects ecosystem functioning.

Whilst this method can theoretically be applied to any ecosystem function, it is frequently used in the context of yield from a grassland biodiversity study, which is also the focus of the realised proportion DI modelling presented in this chapter.

From Loreau and Hector (2001), the net biodiversity effect is calculated as the difference between the observed total yield (Y_{OBS}) from an experimental unit and its expected total yield (ETY), see Eq. 4.8. The ETY is assumed to be the summation of the products of each species' planted/initial proportion and their mean output in monoculture ($\widehat{Y_{MONO}}$), see Eq. 4.9. This is similar to the use of a DI model with no included interaction effects where only monoculture plots are supplied in the dataset.

$$NBE_m = Y_{OBSm} - ETY_m \quad (\text{Eq. 4.8})$$

$$ETY_m = \sum_{i=1}^S \hat{Y}_{MONO_i} p_{im} \quad (\text{Eq. 4.9})$$

The complementarity effect is obtained by taking the product of the richness (the number of unique species present in an ecosystem), the mean change in relative yield (proportions) of the species present, and the mean monoculture yields. Finally, the selection effect is defined as the product of the richness and the covariance between the change in species proportions of the mixtures and the change in the relative yields of the monocultures.

In short, the additive partitioning method assumes that without species interactions (defined by the net biodiversity effect), the achieved yield of each species is equal to the mean monoculture yield for the species, scaled by the initial proportion of the species in the ecosystem. The net biodiversity effect, along with the complementarity and selection effects, can then be inferred by comparing these assumed values with the true observed values for the sorted yields. The effectiveness of this method is therefore reliant on the estimation of

the expected relative yield of each species present. The ID DI model (where it is assumed that species do not interact), could be utilised for this estimation, improving upon the original by allowing both monoculture and mixture plots to contribute to the estimation of the ETY, increasing the data available for estimation while maintaining the assumption of a null interaction effect (Eq. 4.10).

$$NBE_m = Y_{OBSm} - \sum_{i=1}^S \hat{\beta}_i p_{im} \quad (\text{Eq. 4.10})$$

Since the model developed in this chapter (Section 4.2) can predict both the ETY (y_D) and the resulting/achieved relative species abundances ($y_1, \dots, y_C$), my approach can also be used to estimate the selection and complementarity effects which can be valuable summarising metrics for the model (which may contain an otherwise large number of coefficients). For example, using the additive partitioning method in conjunction with my model, Eq. 4.7, and ignoring weeds, the net biodiversity effect, along with its breakdown into the complementarity and selection effects, can be estimated for any mixture of the component species, regardless of whether the selected community was sown in the original experimental design or not (given there was sufficient information provided around this space), see Figure 4.2 for an example. From this figure, we can see that the communities which had the highest predicted sown species yield (i.e. ignoring the proportion of weeds), according to Figure 4.1, also appear to have the highest net biodiversity effect. The majority of the net biodiversity effects presented can be attributed to complementarity effects, which indicates resource partitioning or positive species interactions for the total yield, however, selection effects can be seen in all selected sown communities with the exception of the equi-proportional grass and herb mixture (G/H). Selection effects indicate that in the community, the realised dominance of a species is influencing the total yield (either positively or negatively), for the grass and herb mixture noted, selection effects likely do not take place,

despite a seeming realised dominance of herb (see Figure 4.1) as the ID effects of grasses and herbs for the total yield response are roughly equal (see Table 4.1).

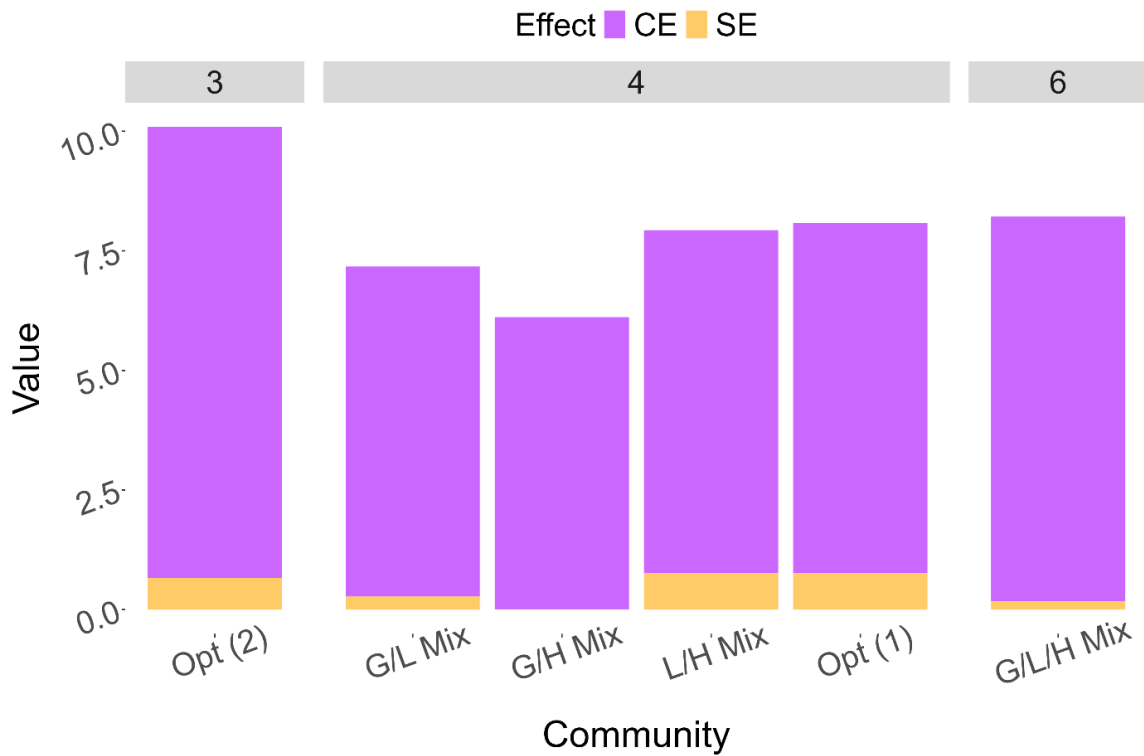


Figure 4.3: A stacked bar chart showcasing the predicted net biodiversity effect, and its breakdown into the complementarity (CE) and selection (SE) effects, of a set of selected sown communities, using the additive partitioning method (Loreau & Hector, 2001) and the DI model given in Eq. 4.7 with the data from Grange *et al.*, 2021. The sown communities represented match those of Figure 4.1 which contained multiple functional groups: a 25/25/25/25 mix of each species between pairs of functional groups, an equal mix of all six species, and the best performing community for two optimal ecosystem functioning measures, ‘Opt (1)’, which is the community which experienced the least change in species proportions (50.9% Lp, 0.0% Pp, 1.1% Tp, 0.8% Tr, 0.0% Ci, and 47.2% Pl) and ‘Opt (2)’ which produced the highest sown species yield (i.e. excluding weeds) (sown proportions: 0.0% Lp, 27.9% Pp, 40.6% Tp, 0.0% Tr, 0.0% Ci, and 31.5% Pl).

4.5 Discussion

The newly developed DI model for jointly modelling a total response and its proportional breakdown by species presented in this chapter can allow researchers to gain a deeper insight into the mechanisms underlying changes in the BEF relationship. The prediction of changes in species proportions over time using the DI modelling framework in conjunction with the additive partitioning method (Loreau & Hector, 2001) also allows for the estimation of the selection and complementarity effects for communities of interest without said communities necessarily needing to be sown in the original experiment, given that there is sufficient coverage around the design space.

Alternative methods for analysing the change in species proportions typically include a transformation, or less commonly a regression, between an initial set of species proportions and a resulting set of proportions, therefore not allowing for species interactions to be examined (Aitchison, 1982; Tsagris *et al.*, 2011; Alenazi, 2019). Basic Dirichlet regression was also unsuitable, due to its inability to handle responses with zero values, generally, if zero values for any compositional response are required/desired, zero values are transformed to near-zero values and then back-transformed in the prediction process (Smithson & Verkuilen, 2006). In BEF relationship studies, ‘true’ or ‘structural’ zeroes are common and important to distinguish from near-zero values, for example, in the case that a species was never sown in a field experiment, its observed realised proportion should remain at zero, unless the species invaded the plot. The model presented in this chapter has built upon these older methods, not only introducing novelty in the prediction of a continuous response alongside a set of compositional responses within a singular model but also, more generally, in the modelling of multiple responses which follow different distributions. Previous work in this area has not been applied to BEF relationship data, limits the distributions available to work with, and utilises either a Bayesian framework or includes random effects, which

may not be suitable for every study analysis (Hadfield *et al.*, 2007; Goldstein, *et al.*, 2009; Bürkner, 2017)

For these models, further work could also consider the inclusion of higher order species interactions, i.e. allowing the interactive effect of two sown species to affect the realised proportion of some third species, or the thorough testing of the effects/implications and best estimation practices of the non-linear parameter θ on multivariate and/or bound responses, similarly to the study done by Vishwakarma *et al.*, 2023 for univariate DI models. Further testing of the normality assumption of the error term in the model could also be done. The total ecosystem function is usually assumed to follow a normal distribution when studied using linear regression models, such as with DI models (Kirwan *et al.*, 2009). The proportional responses, if assumed to follow a Dirichlet distribution, have asymptotical normal errors (Boyles, 1997; Hijazi, 2006; Ouimet, 2022; Ouimet & Tolosana-Delgado, 2022), therefore a transformation, like Pearson standardised residuals, will certainly follow some normal distribution and can be used as pseudo residuals in place of traditional residuals in the likelihood function. Other pseudo residuals are also suggested as a method of model fit tests, e.g. compositional residuals (Aitchison, 1982). Any transformation on the compositional response residuals may affect estimation and so should only be used in the case that the assumption would be broken.

As the methodology presented here relies upon maximum likelihood estimation, it has an inherent weakness in that the process is susceptible to getting stuck in a local maxima/minima. This vulnerability may be mitigated by choosing a good set of starting values for the coefficients to be estimated and utilising appropriate optimisation methods. In the case study of this chapter, the starting values for estimation were selected by first using the `DI()` method from the `DImodels` R package (Moral *et al.*, 2023) to select a starting value for each non-linear parameter θ by fitting a univariate linear regression model (ignoring the 0-1 proportional restriction of the realised weed proportion response) and

allowing the θ to be estimated using profile likelihood optimisation (Brent, 1973). The `optim()` function was then used to jointly estimate the two θ values using joint profile likelihood optimisation. The starting values for the fixed effect coefficients of the total yield response were selected from a univariate regression model, fit using the `DI()` function (Moral *et al.*, 2023), while the coefficients for the proportional responses were estimated using univariate quasi binomial regression models, using the function `prop.reg()` from the `Rfast` R package (Papadakis *et al.*, 2023). Finally, the initial values used for the response variance-covariance matrix Σ were selected from the variance-covariance matrix of the observed responses, using the function `vcov()`. From Appendix J, we can see that generally the model fitted well, although there was some underprediction of the yield response that may benefit from further investigation.

In this chapter, I presented the model as being used to predict a single ecosystem function (yield) along with the realised species proportions, however, as it is built upon the multivariate normal distribution, it is easily expandable to allow for the investigation into multiple ecosystem functions simultaneously (leading to $y_{D1}, y_{D2}, \dots$ etc). This would allow for multifunctionality (Maestre *et al.*, 2012; Dooley *et al.*, 2015; Argens *et al.*, 2023; Grange *et al.*, 2024) and changes in species proportions to be investigated within a single model. Likewise, the model could be further expanded to account for repeated measures, however, this may be more complex as this would involve repeatedly measuring not only the normally distributed ecosystem function response(s), but also the compositional responses, which follow a Dirichlet distribution with true zeroes.

For the case study of this chapter, I modelled the realised proportions of the three sown functional groups, plus an additional functional group of weeds, as the original study from which the data came (Grange *et al.*, 2021) questioned the effectiveness of functional groups rather than individual species. This reduced the number of responses by three whilst still allowing for the investigation into inter-functional group interactions. However, the set

of realised proportions of each sown species could also be included and would allow for investigation into the interactions between pairs of individual species.

In the case study of this chapter, the responses of seven individual harvests from a single experiment (Grange *et al.*, 2021) were annualised in order to investigate the effect of species ID and interaction effects on the biomass collected over the course of the growing year (2018). However, if the change in species proportions between harvests was of interest, a series of these single-time point models could instead be fit, one per harvest for a total of seven models. The same set of predictors could be used in each of these models (i.e. the sown species proportions from the initial time point) or, as species compositions may drastically change over time, the realised proportions from each harvest could instead be used as the predictor terms for each subsequent harvest, i.e. the realised species proportions from harvest one could be used to inform the responses from harvest two, those from harvest two used for the responses of harvest three, and so on. This option will be especially important for longstanding studies, which generally last for several years. It may be of interest to test the length of time that a set of predictor species proportions, e.g. sown proportions, is viable for use in DI models as opposed to the realised species proportions from some time point closer to that of the response. For example, in a four-year study, the sown species proportions may have a strong influence on the biomass production of the first and possibly second years, however, within that time, some species may be lost or new species may have invaded the plot, leaving the sown species proportions non-viable for use with the biomass responses of years three and four. Using the model presented in this chapter, the difference/change between the sown proportions of a plot and its predicted realised proportions could be used to inform or dictate when researchers may need to change the set of predictors used, opting for the realised proportions (observed or predicted) from a later time point, or include a new predictor which scales the predicted responses by the expected change in species proportions over long periods of time.

This novel methodology could also be used to improve upon the estimation of the threshold for transgressive overyielding (Trenbath, 1974; Grange *et al.*, 2021), which refers to a mixture outperforming the best-performing monoculture. The ‘best’ monoculture is usually defined by the monoculture with the highest estimated ecosystem function response based on the set of model predictors, usually initial species proportions for DI models (Kirwan *et al.*, 2009). However, if a monoculture’s species proportions shifts significantly before the ecosystem function response is measured, then this response may be more closely related to the invasion of additional species (weeds), rather than the single species initially sown. The novel methodology presented here may be useful in the redefinition of transgressive overyielding in the case of these large shifts, as one could attribute monoculture performance to the realised proportions instead.

I acknowledge use of the experimental design and elements of the sampling protocol developed by LegacyNet (<https://legacynet.scss.tcd.ie/>).

5 Discussion

For this thesis, I aimed to test the best practices for the design of a biodiversity ecosystem function (BEF) relationship experimental study and to address some outstanding gaps in statistical modelling tools for use with BEF relationship studies, specifically with Diversity-Interactions (DI) models for use with multiple correlated responses (temporal and/or multivariate responses taken on individual experimental units). The three main aims of my thesis are listed in Section 1.5, and I have fulfilled these aims in chapters 2-4 of my thesis. I will first summarise how I have achieved these aims and will then provide more detail chapter by chapter.

*Aim 1: To explore and formally test the effects of different experimental
BEF relationship study designs through a simulation study.*

To achieve aim 1, I have conducted a simulation study to test the effects of varying an BEF relationship study's experimental design on the ability of the DI modelling framework to correctly detect the true underlying interaction structure (chapter 2).

*Aim 2: To expand the DI modelling framework for usage with data
containing different forms of correlated responses.*

To achieve aim 2, I have developed two new models. The first enables the simultaneous modelling of repeated measures and multivariate data using the DI modelling framework (chapter 3). The second allows for the joint modelling of some total ecosystem function response and the breakdown of the realised species proportions (chapter 4).

*Aim 3: To develop statistical tools to provide a user-friendly way to fit
complex DI models.*

To achieve aim 3, I have developed the R package `DImodelsMulti` and published it on CRAN (chapter 3 and Byrne *et al.*, 2024). This package allows users to easily fit repeated

measures DI models, multivariate DI models, and multivariate repeated measures DI models.

Chapter 2

In chapter 2 of this thesis, I used a simulation study to test the number of times that some true underlying DI model's interaction structure is correctly estimated using the `auto_DI()` function from the `DImodels` package (Moral *et al.*, 2023). This function automatically fits a range of DI models, with differing interaction structures, using least squares estimation and compares them using F-tests to find the best fitting model structure. The identity (ID), average pairwise (AV), and functional group (FG) interaction structures were chosen as the underlying true models to test and a total of 112 experiment design principles were tested for each of the three interaction structures, which were used to simulate 10,000 datasets each, totalling 3.36 million simulated datasets. The experimental designs were created following typical methodology seen in BEF research, particularly randomised controlled trials. The randomness of this simulation study allowed for a reduction in bias when generating experiment design matrices and for the testing of general design principles, i.e. the proportion of even vs. uneven mixtures, rather than specific experimental designs or datasets.

Through this simulation study, I have provided evidence that the general best practice when selecting community designs for a BEF relationship experiment, with a focus on estimating the true underlying species interactions structure, may be to include a wider diversity of community design/species compositions, particularly non-equi-proportional communities, over replicating monoculture designs. This advice should only be applied given that the experiment in question is within the scope of the parameters tested in this simulation study and that standard procedures are implemented with priority, such as the replication of designs as 'fail-safes' to ensure community establishment. This particularly

applies when it is hypothesised that the species included in the study will have invoke a relatively large number of interaction terms, e.g. the functional group (FG) DI structure.

I have also produced highly customisable code for the comparison between experimental designs, which future researchers may find useful, see Appendix B.

Chapter 3

The multivariate repeated measures DI model, that I developed in chapter 3, allows for the simultaneous modelling of multiple ecosystem function responses at multiple different time points. Using this model will enable researchers to investigate the effects of species diversity on multifunctionality at different time points whilst accounting the inherent covariance existing between response readings from the same experimental unit. The coefficients produced by the model may also be formally tested across responses using this methodology, in other words, differences between the effects of individual species (ID) or their interactions across functions and time can be detected. This enables a researcher to track any significant increases or decreases in the effectiveness of a species across time or any significant differences in the contribution of a species between ecosystem functions.

I developed the new R package, `DImodelsMulti`, which follows on from the release of the `DImodels` package (Moral *et al.*, 2023). The original package allowed for the automatic fitting of univariate DI models (Kirwan *et al.*, 2009), including the estimation of the non-linear parameter θ (Connolly *et al.*, 2013). The new package presented provides researchers with the coding tools to fit the multivariate (Dooley *et al.*, 2015), repeated measures (Frankow-Lindberg *et al.*, 2009; Finn *et al.*, 2013), and the newly developed multivariate *and* repeated measures DI models. As of writing, August 2024, it currently has over 3.8k downloads from the CRAN repository (Byrne *et al.*, 2024).

Chapter 4

In Chapter 4, I developed a model which enables the estimation of the effect of a set of predictors on both a total biomass response and the set of realised species proportions simultaneously. To accomplish this, I combined the multivariate normal and multivariate-beta (Dirichlet) log-likelihoods to account for both the correlation between the responses, as they are drawn from the same subject/plot, and the sum-to-one constraint on the proportional responses. The zero-adjusted Dirichlet regression method (Tsagris & Stewart, 2018) was included in this adaption in order to allow for ‘true zero’ values, as they have important interpretations which differ to ‘near-zero’ values, especially when relating to the total loss of a species from an ecosystem versus it being a ‘rare’ species. The modelling structure presented in this chapter therefore provides a statistical framework to allow BEF researchers to investigate the relationship between species diversity and the realised yield with its relative breakdown by species, whilst accounting for the correlation between readings. This allows for deeper investigation into the interspecific mechanisms underlying changes in ecosystem functioning by implicitly modelling the changes in species diversity between ecosystem function readings and explicitly modelling how the realised species proportions covaries with some ecosystem function.

Implications and limitations for BEF research

The work in this thesis contributes in many ways to the ongoing effort to quantify the strengths and shapes of the effects of species diversity on ecosystem functioning. It does this namely through the expansion of the DI modelling framework to further its usage in the analyses of complex BEF relationship studies and the development of statistical tools through which these models can become more accessible to the wider scientific community.

The simulation study that I performed in chapter 2 of this thesis grants an insight into the data requirements of DI models, testing the benefit of the inclusion of mixtures over the

replication of existing monocultures and the preferred composition of said mixtures, between equi- and non-equi-proportional. This study will enable the better planning of BEF relationship experiments for optimal use with the DI modelling framework, allowing for the efficient use of resources, such as space and labour. The provision of experimental design guides to coincide with statistical methodology can greatly improve the use of said methodology, allowing for additional confidence in both the establishment of BEF experiments, knowing that it will provide sufficient power for use with the methodology desired, and may help to remove doubt in its usage by those running the experiment. Conversely, it can inspire confidence in the usage of an analysis tool with already collected data sourced from an experiment whose design appears sufficient. Setting a precedent for developers of statistical methodology to provide formally tested advice or guides on experimental design creation compatible for use with their analysis tool will therefore undoubtedly benefit the research community. However, the use of simulation studies may only provide insight within the bounds of their testing parameters/criteria and therefore their results must be interpreted, and followed, with care.

The development in chapter 3 of a DI model capable of modelling responses which have multiple sources of correlation allows for the combined use of the repeated measures DI model (Frankow-Lindberg *et al.*, 2009; Finn *et al.*, 2013) and the multivariate DI model (Dooley *et al.*, 2015) within a singular model structure. The variance-covariance matrix of the model's error term can capture the correlation between differing ecosystem function responses at different time points. This allows for further investigation into the effect of species diversity on multifunctionality across time without the loss of dimensionality associated with previous methods, e.g., the use of summarising metrics prior to model fitting. The methodology employed for fitting these models also does not require the use of random effects or the selection of informative priors (as is required for Bayesian approaches) and is therefore quite approachable for non-statisticians. The main limitation of this model is in its

large number of parameters, which may make interpretation difficult for extensive studies. However, this issue can be mitigated, without losing the ability to detect differences between individual coefficients, via the application of summarising metrics, such as multifunctionality indexes, to predictions from the multivariate repeated measures model.

The creation and release of the `DImodelsMulti` R package (Byrne *et al.*, 2024) enables the fitting of repeated measures, multivariate, and multivariate repeated measures DI models through a user-friendly medium, following the syntax and usage of the `DImodels` R package (Moral *et al.*, 2023). This package has already been utilised in published BEF studies (Golińska *et al.*, 2023; Finn *et al.*, 2024), showing how the expansion of the available statistical tools for the analysis of BEF data will unquestionably benefit the research field, allowing for the use of potentially complex methodology without the need to create the code from scratch each time. However, tools should not be used without an understanding of how the underlying technology works, as this may cause a user to misinterpret the results obtained. This promotes a need for highly documented tools, with plenty of examples, not only of usage, but also interpretation.

The development of the realised proportions model presented in chapter 4 of this thesis will enable researchers to gain a deeper insight into interspecific relationships and how they affect not only ecosystem functioning but future species diversity. This methodology may also be used in the planning for ecosystem reconstruction, i.e. to best estimate the required initial species proportions to achieve some desired species composition and ecosystem function response value at a set future time point. In agricultural contexts, this could be used to maximise the biomass achieved from a grassland whilst maintaining a high proportion of some desirable species, such as a high protein-grass which would be fed to beef cattle. In natural ecosystems, this model could instead be used to maximise some soil-health metric, such as water-absorption rate, whilst preventing species loss. The novelty of this method may extend beyond use within the BEF research space as it is a regression-based

model which allows for the estimation of multivariate responses which follow different distribution families, in this case, the normal and multivariate-beta (Dirichlet) distribution. This methodology could also be applied to any data where a continuous variable can be broken down into its proportional components or is hypothesised to be related to some set of proportional responses, for example, the cell count in a blood sample and their cell type. This MLE approach may also be applicable to different groups of distributions.

Future work

The work presented in this thesis focuses on expanding the use-case and accessibility of the DI modelling framework in BEF relationship studies, through the development of novel modelling methodologies and the creation of user-friendly statistical tools, such as software. The methodologies and principles followed within this thesis will also be of use to other researchers conducting statistical analyses on BEF data and potentially to disciplines outside of the study of this relationship. For example, to the best of my knowledge, chapters 3 and 4 of this thesis are the first instances of the usage of the non-linear parameter θ (Connolly *et al.*, 2013) for multiple-response DI models and therefore their methods of estimation and their suggested order in which to best perform model selection may guide future usage of this parameter in multivariate settings or inform a future study on this topic. Additionally, the method by which the likelihoods of multiple distributions were combined for the realised proportions models, presented in chapter 4 of this thesis, could be utilised for different groups of responses, hailing from different distribution families.

The simulation study presented here can be easily expanded to vary the experimental design principles tested, i.e. the number of monoculture replicates and the ratio of equi-proportional to uneven mixtures. Different attributes could also be varied, such as the number of plots available, the number of species included in the study, the number of monoculture replicates, or the inclusion of monocultures at all, and the inclusion of

treatments or structural effects, such as blocking. The formulation of the experimental designs for this study was based on randomised controlled trials, which is just one method by which the richness levels and species (from a set pool) are selected for inclusion, therefore, this study could be repeated using any alternative method, such as a more curated selection of richness levels and species, which ensures that each richness level appears and that each species appears within each richness level. The true underlying models could also be varied, through variation in the interaction structure used (e.g. the ADD or FULL structures), the coefficients of the fixed effects, and the error variance of the DI models already presented in chapter 2 or a different modelling methodology could be utilised. These variations would ensure that conclusions drawn from this study are applicable to a number of different systems. The expansion of this simulation to more high-stakes experiments, which include large numbers of species, span across a large period of time, and/or include a very limited number of experimental units, would also be desirable and allow for interpolation to medium-sized designs (as the ones presented here contain a relatively small number of species). Rather than focusing on interaction structures, this study could also be repeated with a focus on estimating the model coefficients or predictions to be within some interval of the true underlying model's.

Currently, for DI models, several R packages have been released to date; `DImodels` (Moral *et al.*, 2023), for the fitting of univariate DI models, `DImodelsMulti` (chapter 3; Byrne *et al.*, 2024), for the fitting of repeated measures, multivariate, or multivariate repeated measures DI models, and `DImodelsVis` (Vishwakarma *et al.*, 2024b), for the plotting of results obtained from fitting DI models and the representation of experimental designs which utilise the simplex space, e.g. those using species proportions. New functionality may be added to these packages, or new packages could be added to the family, to increase the reach and accessibility of the DI modelling framework and other compatible functions, such as the estimation of multifunctionality indexes, the net biodiversity effect

(and its breakdown into the selection and complementarity effects), or additional comparative or plotting functions.

The DI modelling framework can be expanded further beyond those presented in this thesis and those presented in past published works. This could be through the combination of existing methodologies within the framework, similarly to the multivariate repeated measures DI model presented in chapter 3, or via the application of the framework to a totally novel response, such as shown with the set of realised proportions utilised in the model shown in chapter 4. For instance, while the multivariate, repeated measures, and realised proportions models shown here examine within-subject variation through the covariance structure of their errors, there is a growing interest in multi-site analyses, where the between-subject variation, examined through the use of random effects, is of great importance for inclusion in any model (Frankow-Lindberg et al., 2009; Finn et al 2013; Tilman, et al., 2014; Oliver et al., 2015).

The data presented in the case study of chapter 4 of this thesis came from a single site of a wider, international multi-site experiment, LegacyNet, where, due to potentially drastic changes in site conditions between countries, it is expected that the between-site variation will affect the behaviour of both the included fixed effects and the unexplained variance. The expansion of DI models, or other methodology, to be capable of accounting for not only this between-subject variation, but also within-subject variation, responses of different distributions, and/or bound responses is a logical next step to the research presented within this thesis and for the statistical modelling and BEF research fields alike.

The presented model in chapter 4 was shown through the lens of the DI modelling framework, however, it may be used in conjunction with any set of predictors. This work could also be further expanded to allow for the prediction of multiple ecosystem functions simultaneously rather easily and so may be the subject of future work. The models could also be further expanded upon to enable the prediction of these response sets at multiple

different time points (repeated measures), enabling the study of changes to species compositions, in conjunction with ecosystem functioning, across time.

The inclusion, interpretation, and best practice for estimation of the non-linear parameter θ has been thoroughly tested for univariate DI models (Connolly *et al.*, 2013; Moral *et al.*, 2023; Vishwakarma *et al.*, 2023), the results of which were generalised for use with multivariate and/or repeated measures DI models, as presented in chapters 3 and 4 of this thesis. Future work will likely involve the testing of these generalisations, and alternatives, to gain a better understanding of the effects of allowing fixed effects to vary between correlated responses.

All novel models presented within this thesis currently have no recommendations for dealing with missing data, despite it being a common problem in BEF experiments, as harvests may not be able to be collected due to weather, destruction of experimental units, the damage of samples, etc. Future work in this space could examine the suitability of different methodologies for imputing and/or deleting missing values for use with different DI models. Similarly, the time-points of ecosystem function measurements may vary irregularly, the effect of which on DI modelling has not been fully discussed nor tested. For large-scale studies this is commonly due to differences in harvesting practices between countries and/or climates, it may also be due to limited availability for any equipment required to measure certain ecosystem function responses. Some variance-covariance structures have been created for use with irregularly spaced repeated measures, e.g. the continuous autoregressive model of order one (CAR(1)), however these structures impose their own assumptions onto the data, which may not be valid. As such, other methodology for handling such data should also be the topic of future study.

6 References

- Agresti, A. and Min, Y., 2002. Unconditional small-sample confidence intervals for the odds ratio. *Biostatistics*, 3(3), pp.379-386.
- Aitken, A.C., 1936. IV. On least squares and linear combination of observations. *Proceedings of the Royal Society of Edinburgh*, 55, pp.42-48.
- Allan, E., Weisser, W.W., Fischer, M., Schulze, E.D., Weigelt, A., Roscher, C., Baade, J., Barnard, R.L., Beßler, H., Buchmann, N. and Ebeling, A., 2013. A comparison of the strength of biodiversity effects across multiple functions. *Oecologia*, 173, pp.223-237.
- Argens, L., Brophy, C., Weisser, W.W. and Meyer, S., 2023. Functional group richness increases multifunctionality in intensively managed grasslands. *Grassland Research*, 2(3), pp.225-240.
- Arif, S. and Massey, M.D.B., 2023. Reducing bias in experimental ecology through directed acyclic graphs. *Ecology and Evolution*, 13(3), p.e9947.
- Bailey, R.A., 2008. Design of comparative experiments.
- Bannar-Martin, K.H., Kremer, C.T., Ernest, S.M., Leibold, M.A., Auge, H., Chase, J., Declerck, S.A., Eisenhauer, N., Harpole, S., Hillebrand, H. and Isbell, F., 2018. Integrating community assembly and biodiversity to better understand ecosystem function: the Community Assembly and the Functioning of Ecosystems (CAFE) approach. *Ecology Letters*, 21(2), pp.167-180.
- Bates, D., Mächler, M., Bolker, B. and Walker, S., 2015. Fitting linear mixed-effects models using lme4. *Journal of Statistical Software* 67 (1), 48. [https://doi.org/10, 18637](https://doi.org/10.18637).

- Bell, T., Newman, J.A., Silverman, B.W., Turner, S.L. and Lilley, A.K., 2005. The contribution of species richness and composition to bacterial services. *Nature*, 436(7054), pp.1157-1160.
- Binder, S., Isbell, F., Polasky, S., Catford, J.A. and Tilman, D., 2018. Grassland biodiversity can pay. *Proceedings of the National Academy of Sciences*, 115(15), pp.3876-3881.
- Bolker, B., 2012. Multivariate analysis with mixed modeling tools in R. *RPubs.com/bbolker/3336*.
- Boyles, R.A., 1997. Using the chi-square statistic to monitor compositional process data. *Journal of Applied Statistics*, 24(5), pp.589-602.
- Brent, R.P., 1973. Some efficient algorithms for solving systems of nonlinear equations. *SIAM Journal on Numerical Analysis*, 10(2), pp.327-344.
- Brophy, C., Connolly, J., Fagerli, I.L., Duodu, S. and Svenning, M.M., 2011. A baseline category logit model for assessing competing strains of rhizobium bacteria. *Journal of agricultural, biological, and environmental statistics*, 16, pp.409-421.
- Brophy, C., Dooley, Á., Kirwan, L., Finn, J.A., McDonnell, J., Bell, T., Cadotte, M.W. and Connolly, J., 2017. Biodiversity and ecosystem function: making sense of numerous species interactions in multi-species communities. *Ecology*, 98(7), pp.1771-1778.
- Bruehlheide, H., Nadrowski, K., Assmann, T., Bauhus, J., Both, S., Buscot, F., Chen, X.Y., Ding, B., Durka, W., Erfmeier, A. and Gutknecht, J.L., 2014. Designing forest biodiversity experiments: general considerations illustrated by a new large experiment in subtropical China. *Methods in Ecology and Evolution*, 5(1), pp.74-89.
- Bürkner, P.C., 2017. brms: An R package for Bayesian multilevel models using Stan. *Journal of statistical software*, 80, pp.1-28.
- Byrne, L., Vishwakarma R., Moral, R., and Brophy, C., 2024. DImodelsMulti: Fit Multivariate Diversity-Interactions Models with Repeated Measures. *The*

Comprehensive R Archive Network, R package version 1.1.1, <https://cran.r-project.org/package=DImodelsMulti>.

- Byrnes, J.E., Gamfeldt, L., Isbell, F., Lefcheck, J.S., Griffin, J.N., Hector, A., Cardinale, B.J., Hooper, D.U., Dee, L.E. and Emmett Duffy, J., 2014. Investigating the relationship between biodiversity and ecosystem multifunctionality: challenges and solutions. *Methods in Ecology and Evolution*, 5(2), pp.111-124.
- Connolly, J., Bell, T., Bolger, T., Brophy, C., Carnus, T., Finn, J.A., Kirwan, L., Isbell, F., Levine, J., Lüscher, A. and Picasso, V., 2013. An improved model to predict the effects of changing biodiversity levels on ecosystem function. *Journal of Ecology*, 101(2), pp.344-355.
- Connolly, J., Cadotte, M.W., Brophy, C., Dooley, Á., Finn, J., Kirwan, L., Roscher, C. and Weigelt, A., 2011. Phylogenetically diverse grasslands are associated with pairwise interspecific processes that increase biomass. *Ecology*, 92(7), pp.1385-1392.
- Costanza, R., d'Arge, R., De Groot, R., Farber, S., Grasso, M., Hannon, B., Limburg, K., Naeem, S., O'Neill, R.V., Paruelo, J. and Raskin, R.G., 1997. The value of the world's ecosystem services and natural capital. *Nature*, 387(6630), pp.253-260.
- Cuddington, K., 2011. Legacy effects: the persistent impact of ecological interactions. *Biological Theory*, 6, pp.203-210.
- Daepp, M., Nösberger, J. and Lüscher, A., 2001. Nitrogen fertilization and developmental stage alter the response of *Lolium perenne* to elevated CO₂. *New Phytologist*, 150(2), pp.347-358.
- De Groot, R.S., 1992. Functions of nature: evaluation of nature in environmental planning, management and decision making. *Functions of nature: evaluation of nature in environmental planning, management and decision making*.

- Dooley, A., Isbell, F., Kirwan, L., Connolly, J., Finn, J.A. and Brophy, C., 2015. Testing the effects of diversity on ecosystem multifunctionality using a multivariate model. *Ecology Letters*, 18(11), pp.1242-1251.
- Emmett Duffy, J., Paul Richardson, J. and Canuel, E.A., 2003. Grazer diversity effects on ecosystem functioning in seagrass beds. *Ecology letters*, 6(7), pp.637-645.
- Finn, J.A., Kirwan, L., Connolly, J., Sebastià, M.T., Helgadottir, A., Baadshaug, O.H., Bélanger, G., Black, A., Brophy, C., Collins, R.P. and Čop, J., 2013. Ecosystem function enhanced by combining four functional types of plant species in intensively managed grassland mixtures: a 3-year continental-scale field experiment. *Journal of Applied Ecology*, 50(2), pp.365-375.
- Finn, J.A., Suter, M., Vishwakarma, R., Oram, N.J., Lüscher, A. and Brophy, C., 2024. Design principles for multi-species productive grasslands: Quantifying effects of diversity beyond richness. *Journal of Ecology*.
- Frankow-Lindberg, B.E., Brophy, C., Collins, R.P. and Connolly, J., 2009. Biodiversity effects on yield and unsown species invasion in a temperate forage ecosystem. *Annals of botany*, 103(6), pp.913-921.
- Galecki, A.T., 1994. General class of covariance structures for two or more repeated factors in longitudinal data analysis. *Communications in Statistics-Theory and Methods*, 23(11), pp.3105-3119.
- Gamfeldt, L., Hillebrand, H. and Jonsson, P.R., 2008. Multiple functions increase the importance of biodiversity for overall ecosystem functioning. *Ecology*, 89(5), pp.1223-1231.
- Gering, J.C., Crist, T.O. and Veech, J.A., 2003. Additive partitioning of species diversity across multiple spatial scales: implications for regional conservation of biodiversity. *Conservation biology*, 17(2), pp.488-499.

- Goldstein, H., Carpenter, J., Kenward, M.G. and Levin, K.A., 2009. Multilevel models with multivariate mixed response types. *Statistical modelling*, 9(3), pp.173-197.
- Golińska, B., Vishwakarma, R., Brophy, C. and Goliński, P., 2023. Positive effects of plant diversity on dry matter yield while maintaining a high level of forage digestibility in intensively managed grasslands across two contrasting environments. *Grass and Forage Science*, 78(4), pp.438-461.
- Grange, G., Brophy, C., Vishwakarma, R. and Finn, J.A., 2024. Effects of experimental drought and plant diversity on multifunctionality of a model system for crop rotation. *Scientific Reports*, 14(1), p.10265.
- Grange, G., Finn, J.A. and Brophy, C., 2021. Plant diversity enhanced yield and mitigated drought impacts in intensively managed grassland communities. *Journal of Applied Ecology*, 58(9), pp.1864-1875.
- Hadfield, J.D., 2010. MCMC methods for multi-response generalized linear mixed models: the MCMCglmm R package. *Journal of statistical software*, 33, pp.1-22.
- Hadfield, J.D., Nutall, A., Osorio, D. and Owens, I.P.F., 2007. Testing the phenotypic gambit: phenotypic, genetic and environmental correlations of colour. *Journal of evolutionary biology*, 20(2), pp.549-557.
- Hariton, E. and Locascio, J.J., 2018. Randomised controlled trials—the gold standard for effectiveness research. *BJOG: an international journal of obstetrics and gynaecology*, 125(13), p.1716.
- Hector, A. and Bagchi, R., 2007. Biodiversity and ecosystem multifunctionality. *Nature*, 448(7150), pp.188-190.
- Hector, A., Schmid, B., Beierkuhnlein, C., Caldeira, M.C., Diemer, M., Dimitrakopoulos, P.G., Finn, J.A., Freitas, H., Giller, P.S., Good, J. and Harris, R., 1999. Plant diversity and productivity experiments in European grasslands. *Science*, 286(5442), pp.1123-1127.

- Heinen, R., Hannula, S.E., De Long, J.R., Huberty, M., Jongen, R., Kielak, A., Steinauer, K., Zhu, F. and Bezemer, T.M., 2020. Plant community composition steers grassland vegetation via soil legacy effects. *Ecology Letters*, 23(6), pp.973-982.
- Heinen, R., van der Sluijs, M., Biere, A., Harvey, J.A. and Bezemer, T.M., 2018. Plant community composition but not plant traits determine the outcome of soil legacy effects on plants and insects. *Journal of Ecology*, 106(3), pp.1217-1229.
- Hijazi, R.H., 2006. Residuals and diagnostics in dirichlet regression. *ASA Proceedings of the General Methodology Section*, pp.1190-1196.
- Hooper, D.U. and Vitousek, P.M., 1998. Effects of plant composition and diversity on nutrient cycling. *Ecological monographs*, 68(1), pp.121-149.
- Hooper, D.U., Chapin III, F.S., Ewel, J.J., Hector, A., Inchausti, P., Lavorel, S., Lawton, J.H., Lodge, D.M., Loreau, M., Naeem, S. and Schmid, B., 2005. Effects of biodiversity on ecosystem functioning: a consensus of current knowledge. *Ecological monographs*, 75(1), pp.3-35.
- Hox, J.J., Maas, C.J. and Brinkhuis, M.J., 2010. The effect of estimation method and sample size in multilevel structural equation modeling. *Statistica neerlandica*, 64(2), pp.157-170.
- Hunter Jr, M.L., 2002. Fundamentals of Conservation Biology. 547 pp.
- Hurvich, C.M. and Tsai, C.L., 1989. Regression and time series model selection in small samples. *Biometrika*, 76(2), pp.297-307.
- Huston, M.A., 1997. Hidden treatments in ecological experiments: re-evaluating the ecosystem function of biodiversity. *Oecologia*, 110, pp.449-460.
- Isbell, F., Calcagno, V., Hector, A., Connolly, J., Harpole, W.S., Reich, P.B., Scherer-Lorenzen, M., Schmid, B., Tilman, D., Van Ruijven, J. and Weigelt, A., 2011. High plant diversity is needed to maintain ecosystem services. *Nature*, 477(7363), pp.199-202.

- Jochum, M., Fischer, M., Isbell, F., Roscher, C., van der Plas, F., Boch, S., Boenisch, G., Buchmann, N., Catford, J.A., Cavender-Bares, J. and Ebeling, A., 2020. The results of biodiversity–ecosystem functioning experiments are realistic. *Nature Ecology & Evolution*, 4(11), pp.1485-1494.
- Kaisermann, A., de Vries, F.T., Griffiths, R.I. and Bardgett, R.D., 2017. Legacy effects of drought on plant–soil feedbacks and plant–plant interactions. *New Phytologist*, 215(4), pp.1413-1424.
- Kirwan, L., Connolly, J., Brophy, C., Baadshaug, O.H., Belanger, G., Black, A., Carnus, T., Collins, R.P., Čop, J., Delgado, I. and De Vlieghe, A., 2014. The Agrodiversity Experiment: three years of data from a multisite study in intensively managed grasslands.
- Kirwan, L., Connolly, J., Finn, J.A., Brophy, C., Lüscher, A., Nyfeler, D. and Sebastià, M.T., 2009. Diversity–interaction modeling: estimating contributions of species identities and interactions to ecosystem function. *Ecology*, 90(8), pp.2032-2038.
- Klaus, V.H., Whittingham, M.J., Báldi, A., Eggers, S., Francksen, R.M., Hiron, M., Lellei-Kovács, E., Rhymer, C.M. and Buchmann, N., 2020. Do biodiversity-ecosystem functioning experiments inform stakeholders how to simultaneously conserve biodiversity and increase ecosystem service provisioning in grasslands?. *Biological Conservation*, 245, p.108552.
- Ladouceur, E., Blowes, S.A., Chase, J.M., Clark, A.T., Garbowski, M., Alberti, J., Arnillas, C.A., Bakker, J.D., Barrio, I.C., Bharath, S. and Borer, E.T., 2022. Linking changes in species composition and biomass in a globally distributed grassland experiment. *Ecology letters*, 25(12), pp.2699-2712.
- Liu, H., Mi, Z., Lin, L.I., Wang, Y., Zhang, Z., Zhang, F., Wang, H., Liu, L., Zhu, B., Cao, G. and Zhao, X., 2018. Shifting plant species composition in response to climate

- change stabilizes grassland primary production. *Proceedings of the National Academy of Sciences*, 115(16), pp.4051-4056.
- Lloret, F., Peñuelas, J., Prieto, P., Llorens, L. and Estiarte, M., 2009. Plant community changes induced by experimental climate change: seedling and adult species composition. *Perspectives in Plant Ecology, Evolution and Systematics*, 11(1), pp.53-63.
- Loreau, M. and De Mazancourt, C., 2013. Biodiversity and ecosystem stability: a synthesis of underlying mechanisms. *Ecology letters*, 16, pp.106-115.
- Loreau, M. and Hector, A., 2001. Partitioning selection and complementarity in biodiversity experiments. *Nature*, 412(6842), pp.72-76.
- Loreau, M., 1998. Separating sampling and other effects in biodiversity experiments. *Oikos*, 82(3), pp.600-602.
- Lüscher, A., Fuhrer, J. and Newton, P.C., 2005. Grassland: a global resource.
- Maestre, F.T., Quero, J.L., Gotelli, N.J., Escudero, A., Ochoa, V., Delgado-Baquerizo, M., García-Gómez, M., Bowker, M.A., Soliveres, S., Escolar, C. and García-Palacios, P., 2012. Plant species richness and ecosystem multifunctionality in global drylands. *Science*, 335(6065), pp.214-218.
- Manning, P., Van Der Plas, F., Soliveres, S., Allan, E., Maestre, F.T., Mace, G., Whittingham, M.J. and Fischer, M., 2018. Redefining ecosystem multifunctionality. *Nature ecology & evolution*, 2(3), pp.427-436.
- Maureaud, A., Hodapp, D., Van Denderen, P.D., Hillebrand, H., Gislason, H., Spaanheden Dencker, T., Beukhof, E. and Lindegren, M., 2019. Biodiversity–ecosystem functioning relationships in fish communities: biomass is related to evenness and the environment, not to species richness. *Proceedings of the Royal Society B*, 286(1906), p.20191189.

- McDonnell, J., McKenna, T., Yurkonis, K.A., Hennessy, D., de Andrade Moral, R. and Brophy, C., 2023. A mixed model for assessing the effect of numerous plant species interactions on grassland biodiversity and ecosystem function relationships. *Journal of Agricultural, Biological and Environmental Statistics*, 28(1), pp.1-19.
- McQuarrie, A.D. and Tsai, C.L., 1998. Regression and time series model selection. World Scientific.
- Menhinick, E.F., 1964. A comparison of some species-individuals diversity indices applied to samples of field insects. *Ecology*, 45(4), pp.859-861.
- Moore, J.C. and Brodie, J.F., 2023. Diversity, taxonomic versus functional.
- Moral, R.A., Vishwakarma, R., Connolly, J., Byrne, L., Hurley, C., Finn, J.A. and Brophy, C., 2023. Going beyond richness: Modelling the BEF relationship using species identity, evenness, richness and species interactions via the DImodels R package. *Methods in Ecology and Evolution*, 14(9), pp.2250-2258.
- Mouillot, D., Villéger, S., Scherer-Lorenzen, M. and Mason, N.W., 2011. Functional structure of biological communities predicts ecosystem multifunctionality. *PloS one*, 6(3), p.e17476.
- Oliver, T.H., Heard, M.S., Isaac, N.J., Roy, D.B., Procter, D., Eigenbrod, F., Freckleton, R., Hector, A., Orme, C.D.L., Petchey, O.L. and Proença, V., 2015. Biodiversity and resilience of ecosystem functions. *Trends in ecology & evolution*, 30(11), pp.673-684.
- Ouimet, F. and Tolosana-Delgado, R., 2022. Asymptotic properties of Dirichlet kernel density estimators. *Journal of Multivariate Analysis*, 187, p.104832.
- Ouimet, F., 2022. A multivariate normal approximation for the Dirichlet density and some applications. *Stat*, 11(1), p.e410.
- Papadakis, M., Tsagris, M., Dimitriadis, M., Fafalios, S., Tsamardinos, I., Fasiolo, M., Borboudakis, G., Burkardt, J., Zou, C., Lakiotaki, K., Chatzipantsiou, C., 2023.

Rfast: A collection of Efficient and Extremely Fast R Functions. *The Comprehensive R Archive Network*, R package version 2.1.0, <https://cran.r-project.org/package=Rfast>.

Perring, M.P., Bernhardt-Römermann, M., Baeten, L., Midolo, G., Blondeel, H., Depauw, L., Landuyt, D., Maes, S.L., De Lombaerde, E., Carón, M.M. and Vellend, M., 2018. Global environmental change effects on plant community composition trajectories depend upon management legacies. *Global Change Biology*, 24(4), pp.1722-1740.

Pinheiro, J., Bates, D., R Core Team, 2024. nlme: Linear and Nonlinear Mixed Effects Models. *The Comprehensive R Archive Network*, R package version 3.1-165, <https://CRAN.R-project.org/package=nlme>.

R Core Team, 2022. R: A language and environment for statistical computing. *R Foundation for Statistical Computing, Vienna, Austria*. URL <https://www.R-project.org/>.

Roberts, F.S., 2019. Measurement of biodiversity: richness and evenness. *Mathematics of Planet Earth: Protecting Our Planet, Learning from the Past, Safeguarding for the Future*, pp.203-224.

Sebastià, M.T., Kirwan, L. and Connolly, J., 2008. Strong shifts in plant diversity and vegetation composition in grassland shortly after climatic change. *Journal of Vegetation Science*, 19(3), pp.299-306.

Shannon, C.E., 1948. A mathematical theory of communication. *The Bell system technical journal*, 27(3), pp.379-423.

Simpson, E.H., 1949. Measurement of diversity. *nature*, 163(4148), pp.688-688.

Snijders, T.A. and Bosker, R., 2011. Multilevel analysis: An introduction to basic and advanced multilevel modeling.

Strong, J.A., Andonegi, E., Bizsel, K.C., Danovaro, R., Elliott, M., Franco, A., Garces, E., Little, S., Mazik, K., Moncheva, S. and Papadopoulou, N., 2015. Marine biodiversity

- and ecosystem function relationships: the potential for practical monitoring applications. *Estuarine, Coastal and Shelf Science*, 161, pp.46-64.
- Stuart-Smith, R.D., Bates, A.E., Lefcheck, J.S., Duffy, J.E., Baker, S.C., Thomson, R.J., Stuart-Smith, J.F., Hill, N.A., Kininmonth, S.J., Airolidi, L. and Becerro, M.A., 2013. Integrating abundance and functional traits reveals new global hotspots of fish diversity. *Nature*, 501(7468), pp.539-542.
- Suter, M., Huguenin-Elie, O. and Lüscher, A., 2021. Multispecies for multifunctions: combining four complementary species enhances multifunctionality of sown grassland. *Scientific Reports*, 11(1), p.3835.
- Tarnow-Mordi, W., Cruz, M., Morris, J.M. and Mol, B.W., 2017. RCT evidence should drive clinical practice: a day without randomisation is a day without progress. *BJOG: An International Journal of Obstetrics & Gynaecology*, 124(4), pp.613-613.
- Teasdale, J.R., 1998. Cover crops, smother plants, and weed management 247 270 Hatfield JL, Buhler DD & Stewart BA Integrated weed and soil management. *Ann Arbor Press Chelsea, MI*.
- Tilman, D., Isbell, F. and Cowles, J.M., 2014. Biodiversity and ecosystem functioning. *Annual review of ecology, evolution, and systematics*, 45, pp.471-493.
- Tilman, D., Reich, P.B., Knops, J., Wedin, D., Mielke, T. and Lehman, C., 2001. Diversity and productivity in a long-term grassland experiment. *Science*, 294(5543), pp.843-845.
- Tilman, D., Wedin, D. and Knops, J., 1996. Productivity and sustainability influenced by biodiversity in grassland ecosystems. *Nature*, 379(6567), pp.718-720.
- Tracy, B.F. and Sanderson, M.A., 2004. Forage productivity, species evenness and weed invasion in pasture communities. *Agriculture, ecosystems & environment*, 102(2), pp.175-183.

- Trenbath, B.R., 1974. Biomass productivity of mixtures. *Advances in agronomy*, 26, pp.177-210.
- Vishwakarma, R., Brophy, C., Hurley, C., 2024a. PieGlyph: Axis Invariant Scatter Pie Plots. *The Comprehensive R Archive Network*, R package version 1.0, <https://cran.r-project.org/package=PieGlyph>.
- Vishwakarma, R., Brophy, C., Hurley, C., 2024b. DImodelsVis: Visualising and Interpreting Statistical Models Fit to Compositional Data. *The Comprehensive R Archive Network*, R package version 1.0.1, <https://cran.r-project.org/package=DImodelsVis>.
- Vishwakarma, R., Byrne, L., Connolly, J., de Andrade Moral, R. and Brophy, C., 2023. Estimation of the non-linear parameter in Generalised Diversity-Interactions models is unaffected by change in structure of the interaction terms. *Environmental and Ecological Statistics*, 30(3), pp.555-574.
- Vogel, A., Ebeling, A., Gleixner, G., Roscher, C., Scheu, S., Ciobanu, M., Koller-France, E., Lange, M., Lochner, A., Meyer, S.T. and Oelmann, Y., 2019. A new experimental approach to test why biodiversity effects strengthen as ecosystems age. In *Advances in ecological research* (Vol. 61, pp. 221-264). Academic Press.
- Wang, Y., Cadotte, M.W., Chen, Y., Fraser, L.H., Zhang, Y., Huang, F., Luo, S., Shi, N. and Loreau, M., 2019. Global evidence of positive biodiversity effects on spatial ecosystem stability in natural grasslands. *Nature communications*, 10(1), p.3207.
- Weigelt, A., Marquard, E., Temperton, V.M., Roscher, C., Scherber, C., Mwangi, P.N., Von Felten, S., Buchmann, N., Schmid, B., Schulze, E.D. and Weisser, W.W., 2010. The Jena Experiment: six years of data from a grassland biodiversity experiment. *Ecology*, 91(3), pp.930-931.
- Whittaker, R.H., 1972. Evolution and measurement of species diversity. *Taxon*, 21(2-3), pp.213-251.

- Wilsey, B.J. and Polley, H.W., 2003. Effects of seed additions and grazing history on diversity and productivity of subhumid grasslands. *Ecology*, 84(4), pp.920-931.
- Wittebolle, L., Marzorati, M., Clement, L., Balloi, A., Daffonchio, D., Heylen, K., De Vos, P., Verstraete, W. and Boon, N., 2009. Initial community evenness favours functionality under selective stress. *Nature*, 458(7238), pp.623-626.

7 Appendices

A The full table of results for each experimental design used in the simulation study

A table of each experimental design principle combination (version) and the number of times, out of the 10,000 runs, it provided enough information for the `auto_DI()` function (Moral *et al.*, 2023) to correctly select the true underlying interaction structure. The number of times each interaction structure was chosen (correct or not) can also be seen. Here, ‘Monos’ refers to the number of monocultures for each species, i.e. the number of monoculture replicates, ‘Even Mixes’ refers to the proportion of the remaining experimental units, after assigning monocultures, that were equi-proportional mixtures, rather than uneven.

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
1	4	FG	50	2	1	6954	60	2491	6954	495
2	4	FG	50	4	1	6305	7	3187	6305	501
3	6	FG	50	2	1	8220	15	1303	8220	462
4	6	FG	50	4	1	7260	5	2228	7260	507
5	4	FG	50	2	0.9	7005	50	2461	7005	484
6	4	FG	50	4	0.9	6434	8	3107	6434	451
7	6	FG	50	2	0.9	8387	16	1127	8387	470
8	6	FG	50	4	0.9	7357	6	2153	7357	484
9	4	FG	50	2	0.7	7157	53	2284	7157	505
10	4	FG	50	4	0.7	6605	11	2902	6605	482
11	6	FG	50	2	0.7	8441	20	1026	8441	513
12	6	FG	50	4	0.7	7616	11	1886	7616	487
13	4	FG	50	2	0.5	7377	73	2061	7377	489
14	4	FG	50	4	0.5	6857	20	2625	6857	498
15	6	FG	50	2	0.5	8662	20	823	8662	495
16	6	FG	50	4	0.5	7829	20	1668	7829	482
17	4	FG	50	2	0.3	7566	84	1890	7566	460
18	4	FG	50	4	0.3	7001	26	2452	7001	521
19	6	FG	50	2	0.3	8730	28	715	8730	527
20	6	FG	50	4	0.3	7997	34	1457	7997	512

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
21	4	FG	50	2	0.1	7840	103	1570	7840	487
22	4	FG	50	4	0.1	7134	40	2291	7134	535
23	6	FG	50	2	0.1	8879	51	603	8879	467
24	6	FG	50	4	0.1	8122	48	1301	8122	527
25	4	FG	50	2	0	7728	159	1586	7728	527
26	4	FG	50	4	0	7259	67	2179	7259	495
27	6	FG	50	2	0	8925	53	576	8925	446
28	6	FG	50	4	0	8210	79	1236	8210	472
29	4	FG	100	2	1	9220	1	259	9220	520
30	4	FG	100	4	1	9160	0	321	9160	519
31	6	FG	100	2	1	9474	0	22	9474	504
32	6	FG	100	4	1	9488	0	11	9488	501
33	4	FG	100	2	0.9	9208	3	248	9208	541
34	4	FG	100	4	0.9	9223	0	259	9223	518
35	6	FG	100	2	0.9	9535	0	12	9535	453
36	6	FG	100	4	0.9	9494	0	12	9494	494
37	4	FG	100	2	0.7	9338	2	167	9338	493
38	4	FG	100	4	0.7	9266	0	220	9266	514
39	6	FG	100	2	0.7	9461	0	8	9461	531
40	6	FG	100	4	0.7	9514	0	6	9514	480
41	4	FG	100	2	0.5	9386	0	126	9386	488
42	4	FG	100	4	0.5	9345	0	174	9345	481
43	6	FG	100	2	0.5	9509	0	3	9509	488
44	6	FG	100	4	0.5	9480	0	5	9480	515
45	4	FG	100	2	0.3	9380	0	118	9380	502
46	4	FG	100	4	0.3	9354	1	139	9354	506
47	6	FG	100	2	0.3	9477	0	4	9477	519
48	6	FG	100	4	0.3	9483	0	2	9483	515
49	4	FG	100	2	0.1	9433	0	71	9433	496
50	4	FG	100	4	0.1	9408	0	109	9408	483
51	6	FG	100	2	0.1	9520	0	1	9520	479
52	6	FG	100	4	0.1	9470	0	2	9470	528
53	4	FG	100	2	0	9413	0	80	9413	507
54	4	FG	100	4	0	9380	0	107	9380	513
55	6	FG	100	2	0	9499	0	0	9499	501
56	6	FG	100	4	0	9498	0	0	9498	502
57	4	FG	50	6	1	5267	4	4250	5267	479
58	6	FG	50	6	1	4664	70	4795	4664	452

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
59	4	FG	50	6	0.9	5448	4	4044	5448	504
60	6	FG	50	6	0.9	4779	76	4665	4779	467
61	4	FG	50	6	0.7	5694	7	3857	5694	442
62	6	FG	50	6	0.7	5123	108	4339	5123	429
63	4	FG	50	6	0.5	5887	21	3615	5887	477
64	6	FG	50	6	0.5	5260	211	4053	5260	471
65	4	FG	50	6	0.3	5964	40	3502	5964	494
66	6	FG	50	6	0.3	5489	318	3684	5489	499
67	4	FG	50	6	0.1	6134	53	3379	6134	434
68	6	FG	50	6	0.1	5807	439	3307	5807	427
69	4	FG	50	6	0	6178	86	3241	6178	495
70	6	FG	50	6	0	5865	492	3135	5865	472
71	4	FG	100	6	1	9109	0	402	9109	489
72	6	FG	100	6	1	9475	0	27	9475	498
73	4	FG	100	6	0.9	9127	0	332	9127	541
74	6	FG	100	6	0.9	9471	0	25	9471	504
75	4	FG	100	6	0.7	9161	0	300	9161	539
76	6	FG	100	6	0.7	9471	0	23	9471	506
77	4	FG	100	6	0.5	9272	0	246	9272	482
78	6	FG	100	6	0.5	9464	0	10	9464	526
79	4	FG	100	6	0.3	9270	0	198	9270	532
80	6	FG	100	6	0.3	9490	0	7	9490	503
81	4	FG	100	6	0.1	9360	0	146	9360	494
82	6	FG	100	6	0.1	9481	0	8	9481	511
83	4	FG	100	6	0	9329	0	136	9329	535
84	6	FG	100	6	0	9492	0	4	9492	504
85	4	FG	50	1	1	6895	373	2229	6895	497
86	6	FG	50	1	1	7983	132	1382	7983	502
87	4	FG	50	1	0.9	6952	342	2215	6952	487
88	6	FG	50	1	0.9	8131	101	1268	8131	500
89	4	FG	50	1	0.7	7329	252	1945	7329	471
90	6	FG	50	1	0.7	8347	118	1070	8347	465
91	4	FG	50	1	0.5	7555	253	1707	7555	485
92	6	FG	50	1	0.5	8478	101	929	8478	492
93	4	FG	50	1	0.3	7751	215	1547	7751	487
94	6	FG	50	1	0.3	8625	107	766	8625	502
95	4	FG	50	1	0.1	7940	236	1317	7940	507
96	6	FG	50	1	0.1	8790	114	577	8790	518

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
97	4	FG	50	1	0	7941	294	1288	7941	477
98	6	FG	50	1	0	8882	115	519	8882	484
99	4	FG	100	1	1	9201	33	288	9201	478
100	6	FG	100	1	1	9369	3	81	9369	547
101	4	FG	100	1	0.9	9231	15	270	9231	484
102	6	FG	100	1	0.9	9455	1	62	9455	482
103	4	FG	100	1	0.7	9261	5	192	9261	542
104	6	FG	100	1	0.7	9440	1	30	9440	529
105	4	FG	100	1	0.5	9363	5	121	9363	511
106	6	FG	100	1	0.5	9460	0	13	9460	527
107	4	FG	100	1	0.3	9399	2	88	9399	511
108	6	FG	100	1	0.3	9488	0	2	9488	510
109	4	FG	100	1	0.1	9385	1	70	9385	544
110	6	FG	100	1	0.1	9534	0	4	9534	462
111	4	FG	100	1	0	9448	3	56	9448	493
112	6	FG	100	1	0	9530	0	6	9530	464
113	4	AV	50	2	1	9054	8	9054	438	500
114	4	AV	50	4	1	9078	0	9078	433	489
115	6	AV	50	2	1	9149	0	9149	398	453
116	6	AV	50	4	1	9156	0	9156	360	484
117	4	AV	50	2	0.9	9091	3	9091	428	478
118	4	AV	50	4	0.9	9087	0	9087	462	451
119	6	AV	50	2	0.9	9111	0	9111	429	460
120	6	AV	50	4	0.9	9149	0	9149	385	466
121	4	AV	50	2	0.7	9064	2	9064	441	493
122	4	AV	50	4	0.7	9084	0	9084	445	471
123	6	AV	50	2	0.7	9105	0	9105	381	514
124	6	AV	50	4	0.7	9127	0	9127	405	468
125	4	AV	50	2	0.5	9042	12	9042	467	479
126	4	AV	50	4	0.5	9040	0	9040	459	501
127	6	AV	50	2	0.5	9079	1	9079	433	487
128	6	AV	50	4	0.5	9089	0	9089	437	474
129	4	AV	50	2	0.3	9054	21	9054	465	460
130	4	AV	50	4	0.3	9049	0	9049	437	514
131	6	AV	50	2	0.3	9057	1	9057	426	516
132	6	AV	50	4	0.3	9113	0	9113	396	491
133	4	AV	50	2	0.1	9017	39	9017	460	484
134	4	AV	50	4	0.1	9017	4	9017	445	534

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
135	6	AV	50	2	0.1	9106	2	9106	427	465
136	6	AV	50	4	0.1	9074	0	9074	406	520
137	4	AV	50	2	0	8943	76	8943	458	523
138	4	AV	50	4	0	9030	12	9030	469	489
139	6	AV	50	2	0	9122	4	9122	444	430
140	6	AV	50	4	0	9111	0	9111	427	462
141	4	AV	100	2	1	9032	0	9032	449	519
142	4	AV	100	4	1	8993	0	8993	483	524
143	6	AV	100	2	1	9055	0	9055	444	501
144	6	AV	100	4	1	9062	0	9062	446	492
145	4	AV	100	2	0.9	8986	0	8986	472	542
146	4	AV	100	4	0.9	8992	0	8992	486	522
147	6	AV	100	2	0.9	9082	0	9082	457	461
148	6	AV	100	4	0.9	9036	0	9036	473	491
149	4	AV	100	2	0.7	9026	0	9026	479	495
150	4	AV	100	4	0.7	9014	0	9014	471	515
151	6	AV	100	2	0.7	9019	0	9019	450	531
152	6	AV	100	4	0.7	9108	0	9108	418	474
153	4	AV	100	2	0.5	9064	0	9064	447	489
154	4	AV	100	4	0.5	9032	0	9032	473	495
155	6	AV	100	2	0.5	9063	0	9063	453	484
156	6	AV	100	4	0.5	9019	0	9019	476	505
157	4	AV	100	2	0.3	9015	0	9015	480	505
158	4	AV	100	4	0.3	9003	0	9003	482	515
159	6	AV	100	2	0.3	9025	0	9025	464	511
160	6	AV	100	4	0.3	9021	0	9021	470	509
161	4	AV	100	2	0.1	9007	0	9007	492	501
162	4	AV	100	4	0.1	9057	0	9057	460	483
163	6	AV	100	2	0.1	9080	0	9080	442	478
164	6	AV	100	4	0.1	9042	0	9042	434	524
165	4	AV	100	2	0	9029	0	9029	461	510
166	4	AV	100	4	0	9024	0	9024	458	518
167	6	AV	100	2	0	9025	0	9025	469	506
168	6	AV	100	4	0	9061	0	9061	433	506
169	4	AV	50	6	1	9065	0	9065	459	476
170	6	AV	50	6	1	9213	0	9213	356	430
171	4	AV	50	6	0.9	9033	0	9033	473	494
172	6	AV	50	6	0.9	9193	1	9193	355	451

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
173	4	AV	50	6	0.7	9109	0	9109	461	430
174	6	AV	50	6	0.7	9203	0	9203	381	416
175	4	AV	50	6	0.5	9079	0	9079	444	477
176	6	AV	50	6	0.5	9135	1	9135	412	452
177	4	AV	50	6	0.3	9068	0	9068	443	489
178	6	AV	50	6	0.3	9136	5	9136	388	471
179	4	AV	50	6	0.1	9076	3	9076	495	426
180	6	AV	50	6	0.1	9187	13	9187	384	411
181	4	AV	50	6	0	9013	3	9013	488	496
182	6	AV	50	6	0	9131	28	9131	386	454
183	4	AV	100	6	1	9054	0	9054	459	487
184	6	AV	100	6	1	9065	0	9065	458	477
185	4	AV	100	6	0.9	8948	0	8948	506	546
186	6	AV	100	6	0.9	9057	0	9057	449	494
187	4	AV	100	6	0.7	9026	0	9026	431	543
188	6	AV	100	6	0.7	9042	0	9042	459	499
189	4	AV	100	6	0.5	9082	0	9082	443	475
190	6	AV	100	6	0.5	8983	0	8983	484	533
191	4	AV	100	6	0.3	8993	0	8993	483	524
192	6	AV	100	6	0.3	9027	0	9027	477	496
193	4	AV	100	6	0.1	9081	0	9081	428	491
194	6	AV	100	6	0.1	9025	0	9025	470	505
195	4	AV	100	6	0	9061	0	9061	404	535
196	6	AV	100	6	0	9063	0	9063	439	498
197	4	AV	50	1	1	8889	125	8889	486	497
198	6	AV	50	1	1	9102	7	9102	395	496
199	4	AV	50	1	0.9	8964	103	8964	443	484
200	6	AV	50	1	0.9	9080	3	9080	427	490
201	4	AV	50	1	0.7	8980	98	8980	446	476
202	6	AV	50	1	0.7	9155	9	9155	375	461
203	4	AV	50	1	0.5	8945	109	8945	458	487
204	6	AV	50	1	0.5	9073	9	9073	434	484
205	4	AV	50	1	0.3	8886	135	8886	484	495
206	6	AV	50	1	0.3	9058	21	9058	422	499
207	4	AV	50	1	0.1	8833	198	8833	471	498
208	6	AV	50	1	0.1	9015	34	9015	437	514
209	4	AV	50	1	0	8746	280	8746	494	479
210	6	AV	50	1	0	9005	73	9005	448	474

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
211	4	AV	100	1	1	8979	20	8979	519	482
212	6	AV	100	1	1	8959	1	8959	495	545
213	4	AV	100	1	0.9	9039	8	9039	468	485
214	6	AV	100	1	0.9	9028	0	9028	492	480
215	4	AV	100	1	0.7	8953	7	8953	501	539
216	6	AV	100	1	0.7	8993	0	8993	474	533
217	4	AV	100	1	0.5	9016	6	9016	465	513
218	6	AV	100	1	0.5	8964	0	8964	505	531
219	4	AV	100	1	0.3	9012	4	9012	461	523
220	6	AV	100	1	0.3	9059	0	9059	431	510
221	4	AV	100	1	0.1	9001	6	9001	454	539
222	6	AV	100	1	0.1	9059	1	9059	479	461
223	4	AV	100	1	0	9028	12	9028	468	492
224	6	AV	100	1	0	9046	2	9046	493	459
225	4	ID	50	2	1	8500	8500	499	454	505
226	4	ID	50	4	1	8586	8586	458	450	505
227	6	ID	50	2	1	8647	8647	476	405	472
228	6	ID	50	4	1	8641	8641	456	385	518
229	4	ID	50	2	0.9	8568	8568	482	432	487
230	4	ID	50	4	0.9	8596	8596	468	467	469
231	6	ID	50	2	0.9	8601	8601	488	434	477
232	6	ID	50	4	0.9	8645	8645	451	415	489
233	4	ID	50	2	0.7	8571	8571	475	443	500
234	4	ID	50	4	0.7	8592	8592	461	456	491
235	6	ID	50	2	0.7	8630	8630	451	395	524
236	6	ID	50	4	0.7	8630	8630	453	431	486
237	4	ID	50	2	0.5	8525	8525	515	469	486
238	4	ID	50	4	0.5	8555	8555	457	473	515
239	6	ID	50	2	0.5	8580	8580	489	437	494
240	6	ID	50	4	0.5	8591	8591	469	447	493
241	4	ID	50	2	0.3	8614	8614	455	471	458
242	4	ID	50	4	0.3	8560	8560	471	446	522
243	6	ID	50	2	0.3	8546	8546	496	427	531
244	6	ID	50	4	0.3	8607	8607	465	415	513
245	4	ID	50	2	0.1	8589	8589	453	471	485
246	4	ID	50	4	0.1	8578	8578	429	450	543
247	6	ID	50	2	0.1	8612	8612	487	431	470
248	6	ID	50	4	0.1	8601	8601	435	425	539
249	4	ID	50	2	0	8520	8520	490	461	528

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
250	4	ID	50	4	0	8569	8569	464	471	495
251	6	ID	50	2	0	8640	8640	459	446	455
252	6	ID	50	4	0	8645	8645	430	445	480
253	4	ID	100	2	1	8501	8501	530	450	519
254	4	ID	100	4	1	8469	8469	527	485	519
255	6	ID	100	2	1	8496	8496	557	446	501
256	6	ID	100	4	1	8551	8551	502	453	494
257	4	ID	100	2	0.9	8452	8452	533	475	540
258	4	ID	100	4	0.9	8492	8492	499	485	524
259	6	ID	100	2	0.9	8506	8506	582	456	456
260	6	ID	100	4	0.9	8496	8496	531	474	499
261	4	ID	100	2	0.7	8483	8483	552	475	490
262	4	ID	100	4	0.7	8528	8528	488	473	511
263	6	ID	100	2	0.7	8487	8487	528	447	538
264	6	ID	100	4	0.7	8540	8540	553	423	484
265	4	ID	100	2	0.5	8564	8564	509	441	486
266	4	ID	100	4	0.5	8528	8528	507	470	495
267	6	ID	100	2	0.5	8536	8536	529	453	482
268	6	ID	100	4	0.5	8467	8467	543	480	510
269	4	ID	100	2	0.3	8482	8482	527	484	507
270	4	ID	100	4	0.3	8538	8538	466	483	513
271	6	ID	100	2	0.3	8451	8451	560	470	519
272	6	ID	100	4	0.3	8505	8505	518	466	511
273	4	ID	100	2	0.1	8530	8530	473	497	500
274	4	ID	100	4	0.1	8601	8601	452	456	491
275	6	ID	100	2	0.1	8529	8529	556	438	477
276	6	ID	100	4	0.1	8570	8570	460	440	530
277	4	ID	100	2	0	8548	8548	491	453	508
278	4	ID	100	4	0	8535	8535	493	457	515
279	6	ID	100	2	0	8552	8552	461	473	514
280	6	ID	100	4	0	8556	8556	498	441	505
281	4	ID	50	6	1	8551	8551	463	482	504
282	6	ID	50	6	1	8690	8690	428	414	468
283	4	ID	50	6	0.9	8540	8540	451	487	522
284	6	ID	50	6	0.9	8643	8643	466	410	481
285	4	ID	50	6	0.7	8608	8608	461	477	454
286	6	ID	50	6	0.7	8668	8668	443	440	449
287	4	ID	50	6	0.5	8562	8562	476	463	499
288	6	ID	50	6	0.5	8601	8601	466	442	491
289	4	ID	50	6	0.3	8591	8591	446	460	503

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
290	6	ID	50	6	0.3	8624	8624	426	442	508
291	4	ID	50	6	0.1	8630	8630	421	501	448
292	6	ID	50	6	0.1	8655	8655	471	427	447
293	4	ID	50	6	0	8554	8554	449	496	501
294	6	ID	50	6	0	8656	8656	434	426	484
295	4	ID	100	6	1	8551	8551	494	468	487
296	6	ID	100	6	1	8547	8547	496	468	489
297	4	ID	100	6	0.9	8454	8454	501	505	540
298	6	ID	100	6	0.9	8519	8519	517	457	507
299	4	ID	100	6	0.7	8550	8550	476	430	544
300	6	ID	100	6	0.7	8520	8520	495	472	513
301	4	ID	100	6	0.5	8581	8581	491	448	480
302	6	ID	100	6	0.5	8455	8455	519	494	532
303	4	ID	100	6	0.3	8544	8544	449	481	526
304	6	ID	100	6	0.3	8528	8528	487	480	505
305	4	ID	100	6	0.1	8621	8621	456	427	496
306	6	ID	100	6	0.1	8549	8549	475	470	506
307	4	ID	100	6	0	8597	8597	458	407	538
308	6	ID	100	6	0	8583	8583	458	447	512
309	4	ID	50	1	1	8260	8260	508	488	495
310	6	ID	50	1	1	8524	8524	562	398	506
311	4	ID	50	1	0.9	8342	8342	544	446	495
312	6	ID	50	1	0.9	8564	8564	504	434	494
313	4	ID	50	1	0.7	8520	8520	484	449	479
314	6	ID	50	1	0.7	8641	8641	512	381	462
315	4	ID	50	1	0.5	8532	8532	483	462	491
316	6	ID	50	1	0.5	8587	8587	479	434	500
317	4	ID	50	1	0.3	8507	8507	500	486	489
318	6	ID	50	1	0.3	8585	8585	475	430	510
319	4	ID	50	1	0.1	8514	8514	504	472	506
320	6	ID	50	1	0.1	8532	8532	510	441	516
321	4	ID	50	1	0	8544	8544	477	499	477
322	6	ID	50	1	0	8576	8576	480	461	483
323	4	ID	100	1	1	8454	8454	546	524	472
324	6	ID	100	1	1	8417	8417	543	492	548
325	4	ID	100	1	0.9	8487	8487	547	472	493
326	6	ID	100	1	0.9	8488	8488	539	490	483
327	4	ID	100	1	0.7	8447	8447	513	497	543
328	6	ID	100	1	0.7	8436	8436	559	478	527
329	4	ID	100	1	0.5	8481	8481	542	462	515

	Species	True Model	Units	Monos	Even Mixes	Correct (/10k)	ID	AV	FG	FULL
330	6	ID	100	1	0.5	8400	8400	564	503	533
331	4	ID	100	1	0.3	8495	8495	529	456	520
332	6	ID	100	1	0.3	8507	8507	550	434	509
333	4	ID	100	1	0.1	8507	8507	494	453	546
334	6	ID	100	1	0.1	8528	8528	535	478	459
335	4	ID	100	1	0	8515	8515	523	468	494
336	6	ID	100	1	0	8531	8531	521	488	460

B The R code used in the simulation study for testing experimental design principles

This R code is designed for use with a particular excel workbook formatting (column names) and would require adaption for use with Excel workbooks of any other formatting. The columns required are, in order: “vNb”, “spNb”, “FGs”, “trueModel”, “units”, “monoReps”, “centroidRatio”, “”, “runs”, “correct”, “”, “ID”, “AV”, “E”, “ADD”, “FG”, “FULL”, “errors”.

Blank columns were added for the sake of clarity visually.

```
#Loop through each version
for(vNb in 1:336){
wb <- openxlsx::loadWorkbook("SimulationStudy_notes.xlsx")
options <- openxlsx::read.xlsx("SimulationStudy_notes.xlsx", "Summary")[vNb, ]

#Extract version specifics from xlsx file
spNb <- options$spNb #Number of species
units <- options$units #Number of experimental units
FGs <- strsplit(options$FGs, ", ")[[1]] #Functional groups of species
runs <- options$runs #Number of times to run this design
trueModel <- options$trueModel #True underlying interaction structure
monoReps <- options$monoReps #Number of monoculture replicates (+1 for looping)
centroidRatio <- options$centroidRatio #Proportion of equi-proportional mixtures
#####
##Set True Underlying Parameters
#####

set.seed(111)

correct <- errors <- ID <- AV <- E <- ADD <- FG <- FULL <- 0

#ID effects
beta1 <- 1034.91
beta2 <- 554.74
beta3 <- 654.36
beta4 <- 713.83
```

```

beta5 <- 251.54

beta6 <- 421.12

#Between functional group effects

bfg1 <- 827.14

bfg2 <- 1280.20

bfg3 <- -292.24

#Within functional group effects

wfg1 <- 443.23

wfg2 <- -359.797

wfg3 <- 105.37

#Average pairwise interaction effect

av <- 750.18


#Set empty matrix, which will be filled with the experiment design

expDes <- matrix(ncol=spNb, nrow=0, dimnames=list(NULL, c(paste0("p", 1:spNb))))


#Randomly assign species proportions according to design principles of this version

for(r in 1:runs){

  #Make Dataset Layout

  props <- data.frame(matrix(ncol=spNb, nrow=0, dimnames=list(NULL, c(paste0("p", 1:spNb)))))

  index <- 1 #row number

  #Monocultures

  for(i in 1:spNb){

    for(j in 1:monoReps){

      props[index, i] <- 1

      index <- index + 1

    }

  }

  unitsLeft <- units - index + 1 #Number of mixture plots

  indexEqual <- 1

  #Even Mixtures

  while(indexEqual <= round(unitsLeft*centroidRatio)){

```

```

rich <- sample(2:spNb, 1) #select richness level for this unit

sp <- sample(1:spNb, rich) #randomly pick species from pool

for(i in sp){

  props[index, i] <- 1/rich #equal proportions

}

index <- index + 1

indexEqual <- indexEqual + 1

}

#Uneven Mixtures

while(index <= units){

  rich <- sample(2:spNb, 1) #select richness level for this unit

  sp <- sample(1:spNb, rich) #randomly pick species from pool

  #Generate list of random numbers that sum to 10, divide by 10 to get proportions

  spProps <- rnorm(rich, 10/rich, 5)

  if(abs(sum(spProps)) < 0.01) {

    spProps <- spProps + 1

  }

  spProps <- round(spProps / sum(spProps) * 10)

  deviation <- 10 - sum(spProps)

  for(. in seq_len(abs(deviation))) {

    spProps[i] <- spProps[i] + sample(rich, 1) + sign(deviation)

  }

  #Ensure all values are positive

  while (any(spProps < 0)) {

    negs <- spProps < 0

    pos <- spProps > 0

    spProps[negs][i] <- spProps[negs][i] + sample(sum(negs), 1) + 1

    spProps[pos][i] <- spProps[pos][i] + sample(sum(pos), 1) - 1

  }

  spProps <- spProps / 10

  #Check that it produced an uneven mixture, if not, redo this index of loop

  if(length(which(spProps != 0)) > 1 & var(spProps[which(spProps != 0)]) != 0){

```

```

    props[index, sp] <- spProps

    index <- index + 1

  }

}

#Fill blanks with 0

props[is.na(props)] <- 0

expDes <- rbind(expDes, props)

##Generate Y values

#####

if(trueModel == "FG"){

  switch(as.character(spNb),

    "4" = props$y <- beta1*props$p1 + beta2*props$p2 + beta3*props$p3 + beta4*props$p4
+ bfg1*(props$p1*props$p3 + props$p1*props$p4 + props$p2*props$p3 + props$p2*props$p4)
+ wfg1*(props$p1*props$p2) + wfg2*(props$p3*props$p4)
+ rnorm(units, 0, 124),

    "6" = props$y <- beta1*props$p1 + beta2*props$p2 + beta3*props$p3 + beta4*props$p4 +
beta5*props$p5 + beta6*props$p6
+ bfg1*(props$p1*props$p3 + props$p1*props$p4 + props$p2*props$p3 + props$p2*props$p4) +
bfg2*(props$p1*props$p5 + props$p1*props$p6 + props$p2*props$p5 + props$p2*props$p6) +
bfg3*(props$p3*props$p5 + props$p3*props$p6 + props$p4*props$p5 + props$p4*props$p6)
+ wfg1*(props$p1*props$p2) + wfg2*(props$p3*props$p4) + wfg3*(props$p5*props$p6)
+ rnorm(units, 0, 124)

  )}

else if(trueModel == "AV"){

  switch(as.character(spNb),

    "4" = props$y <- beta1*props$p1 + beta2*props$p2 + beta3*props$p3 + beta4*props$p4
+ av*(props$p1*props$p2 + props$p1*props$p3 + props$p1*props$p4 +
props$p2*props$p3 + props$p2*props$p4 +
props$p3*props$p4)
+ rnorm(units, 0, 124),

```

```

"6" = props$y <- beta1*props$p1 + beta2*props$p2 + beta3*props$p3 + beta4*props$p4 +
beta5*props$p5 + beta6*props$p6
+ av*(props$p1*props$p2 + props$p1*props$p3 + props$p1*props$p4 + props$p1*props$p5 +
props$p1*props$p6
+ props$p2*props$p3 + props$p2*props$p4 + props$p2*props$p5 + props$p2*props$p6
+ props$p3*props$p4 + props$p3*props$p5 + props$p3*props$p6
+ props$p4*props$p5 + props$p4*props$p6
+ props$p5*props$p6)
+ rnorm(units, 0, 124)
)}

else if(trueModel == "ID"){
  switch(as.character(spNb),
    "4" = props$y <- beta1*props$p1 + beta2*props$p2 + beta3*props$p3 + beta4*props$p4
+ rnorm(units, 0, 124),
    "6" = props$y <- beta1*props$p1 + beta2*props$p2 + beta3*props$p3 + beta4*props$p4 +
beta5*props$p5 + beta6*props$p6
+ rnorm(units, 0, 124)
  )}

##DI Model Selection
#####

tryCatch({
  selected <- DImodels::autoDI(y = "y", prop = paste0("p", 1:spNb), FG = FGs, data = props)

  ##Compare to True Model
  #####

  if(selected$DIcall$DImodel == trueModel){
    correct <- correct + 1
  }

  switch(selected$DIcall$DImodel,
    "ID" = ID <- ID + 1,
    "AV" = AV <- AV + 1,
    "E" = E <- E + 1,
    "ADD" = ADD <- ADD + 1,

```

```

    "FG" = FG <- FG + 1,

    "FULL" = FULL <- FULL + 1

  )},

error = function(e) {

  errors <- errors + 1 # Increment number of singularities incurred

  })} #loop ends here

##Coverage #####

print(correct/runs)

gc()

## Save Info #####

FGs <- paste(FGs, collapse = ", ")

versionInfo <- data.frame(vNb, spNb, FGs, trueModel, units, monoReps, centroidRatio, "", runs, correct, "",
ID, AV, E, ADD, FG, FULL, errors)

openxlsx::writeData(x = versionInfo, wb = wb, sheet = 1, colNames = FALSE, rowNames = FALSE, startRow
= vNb+1, startCol = 1)


expDes <- dplyr::tally(dplyr::group_by_all(expDes))
expDes <- as.data.frame(expDes)
expDes <- dplyr::arrange(expDes, desc(n))
openxlsx::addWorksheet(wb, sheetName = paste0("v", vNb, "_Design"))
openxlsx::writeData(x = expDes, wb = wb, sheet = paste0("v", vNb, "_Design"))
openxlsx::saveWorkbook(wb, "SimulationStudy_notes.xlsx", overwrite = T)

#####

gc()}

```

C The linear transformation applied to the case study ecosystem functions

I follow the linear transformation recommendations laid out in Dooley *et al.*, 2015.

In agronomic settings, nitrogen content and sown biomass are ecosystem functions for which large positive values are desirable while near-zero values are desirable for weed biomass. In order for the desirability of all functions to match in direction, I transform weed biomass into a weed suppression index, for which a large positive value is desirable. To do this, I multiply the value of the weed biomass by -1 and then add the maximum weed biomass value (from the original scale).

To linearly transform all ecosystem functions to lie on the same scale, I convert them to a percentage of the top 10% recorded values for each function, i.e. multiply each value by 100 then divide by the mean of the top 12 values (30 plots over 4 harvests) for that ecosystem function.

D The results of each step in the multivariate repeated measures case study model selection process

The results of the model selection process for the Agrodiversity Study analysis.

When comparing models with varying fixed effects, the maximum likelihood (“ML”) estimation method was used, while if fixed effects did not vary, the restricted maximum likelihood (“REML”) method was used instead, with the final selected model being refit using “REML” in order to produce unbiased estimates.

When models were nested, likelihood ratio tests were used, with an α value of 0.05, otherwise, corrected AIC values were used, where the lower valued model is more desirable. In the case of testing between the θ estimation models, if AICc values were close and θ was estimated close to its boundaries, which leads to fitting issues, the $\theta = 1$ model was chosen instead. The FULL interaction structure was used to estimate and test values of θ .

Structure being tested	Estimation method used	Test/Information criteria	Result
(Function@Time) UN@CS vs. UN@AR(1) vs. UN@UN	REML	AICc: 2454.435 vs. 2476.236 vs. 2133.367	(Function@Time) UN@UN
Interactions: ID vs. AV	ML	Likelihood Ratio: 187.3788 p-value: <.0001	AV
Estimate θ: Sown - 0.01 N - 0.01 Weed - 0.433	ML	AICc: ($\theta = 1$ vs. $\theta \neq 1$): 2272.065 vs. 2274.398 2272.065 vs. 2231.447 2272.065 vs. 2260.965	Sown - 1 N - 1 Weed - 0.433
Interactions: AV vs. FG	ML	Likelihood Ratio: 95.80238 p-value: <.0001	FG
Interactions: FG vs. FULL	ML	Likelihood Ratio: 36.3683 p-value: 0.4515	FG
Treatment: Density incl. vs. Density excl.	ML	Likelihood Ratio: 23.48212 p-value: 0.0239	Density incl.

E Comparison between final multivariate repeated measures and ‘equivalent’ univariate repeated measures models for chapter 3 case study

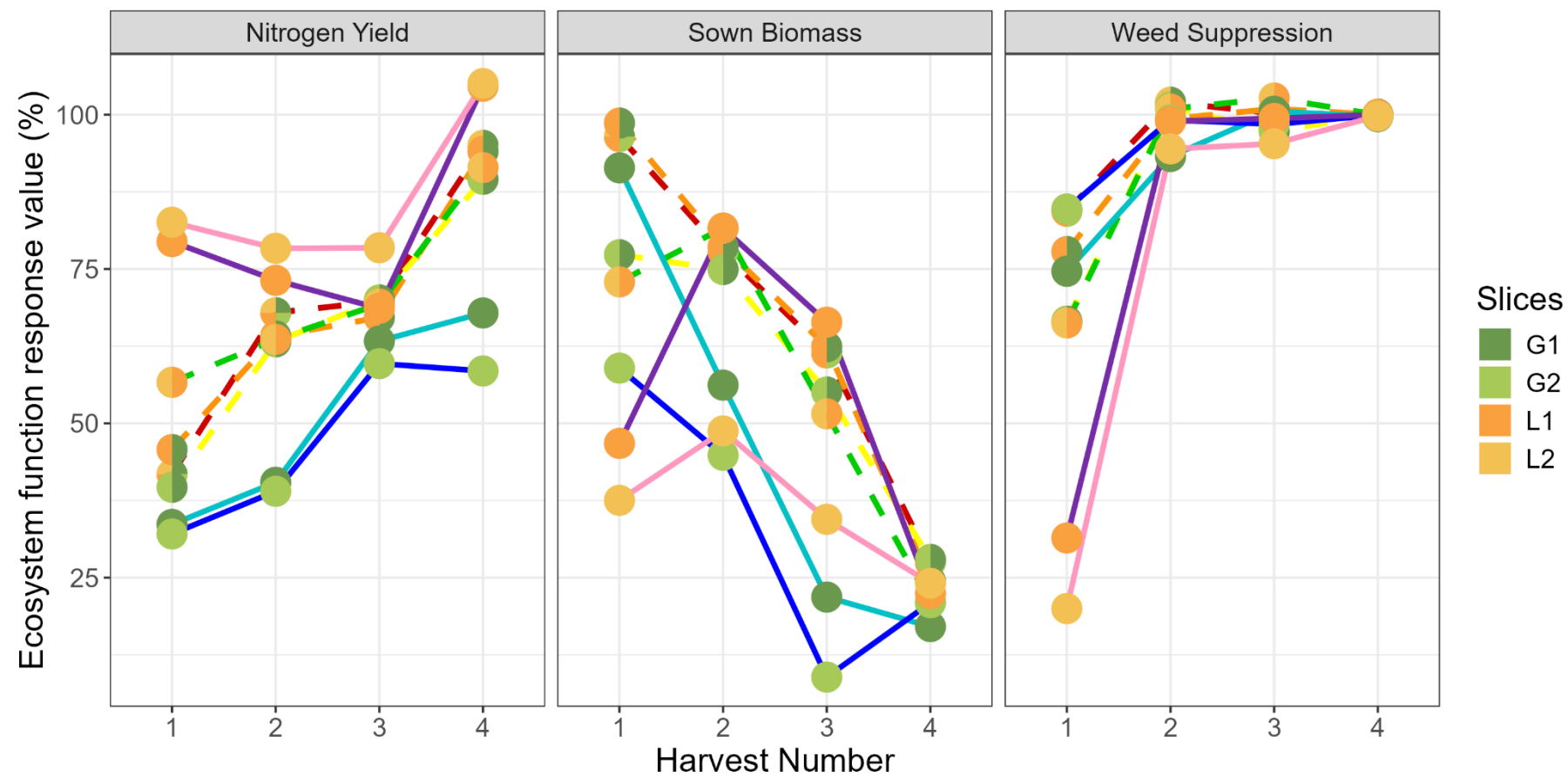
The final selected multivariate repeated measures model from the case study in chapter 3 was adapted to instead be three separate univariate repeated measures models, one for each ecosystem function. The non-linear parameter θ was re-estimated for each model, using the final selected FG DI model with density included as a structural term. The univariately estimated values of θ were all close to zero and therefore caused fitting issues, where values outside the possible range were commonly predicted, therefore all values were set to one. The same issues occurred when the FULL and AV models were instead used for estimation.

I acknowledge that the values selected for the non-linear parameter in the case study that this appendix references may not reflect the true nature of the relationship as the estimation and selection methods used have not been rigorously tested. The estimation issues found with θ throughout this thesis point to a need for the estimation, selection, and interpretation of θ for multivariate and/or bound response variables to be explored further in future work.

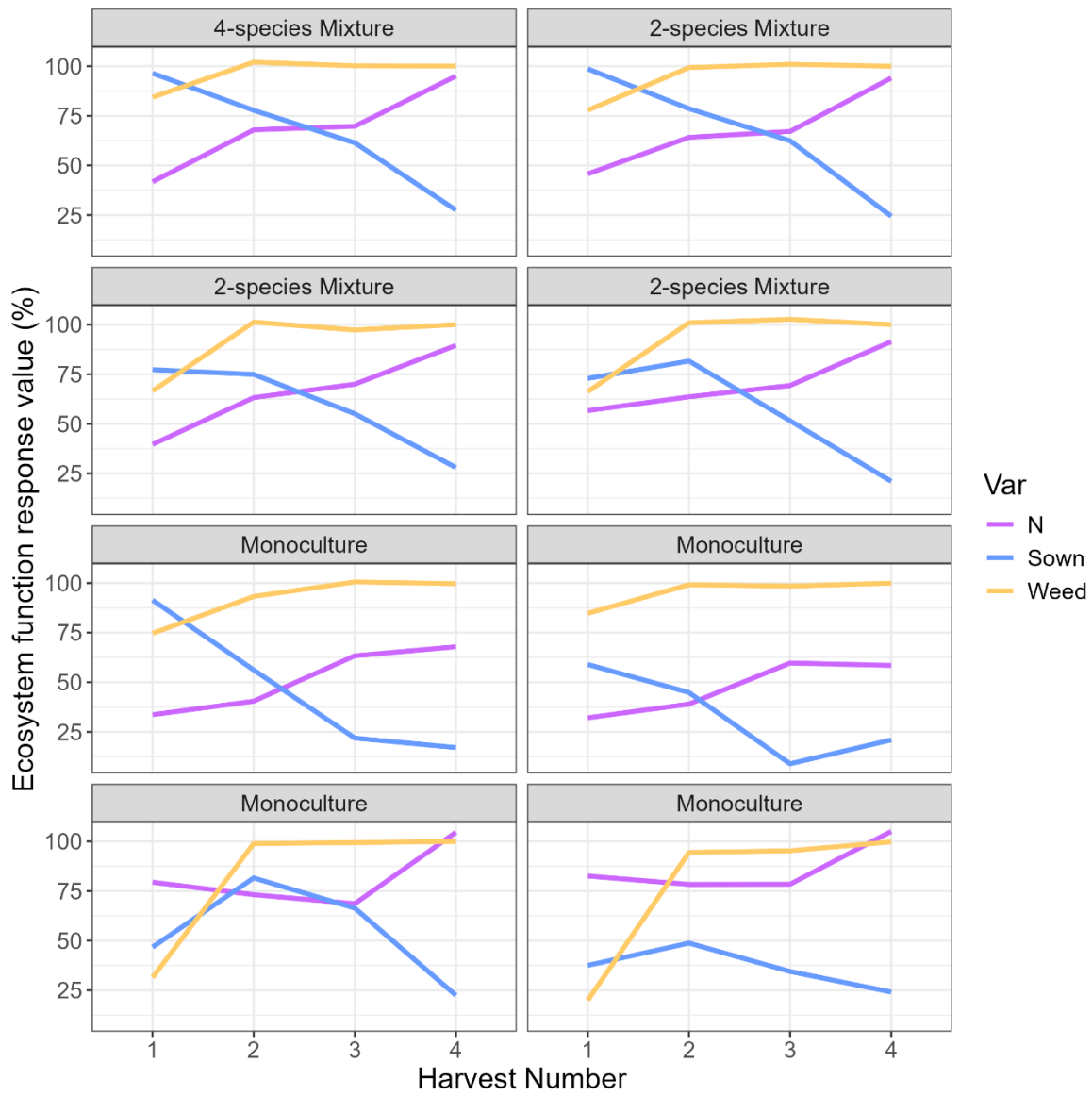
Figure 3.2 is replicated below, instead using the predictions from the univariate models. We can see that the models predict a very similar pattern for each ecosystem function when compared to the multivariate model from chapter 3. The main differences of note occur in the Weed response for mixture communities, noticeably for the 2-species mixture of the legumes, as the θ value for this multivariate response was estimated at a value of roughly 0.43 compared to its univariate value of one. This comparison illustrates that while the same trends may be noted regardless of whether the univariate or multivariate approach is used, the multivariate model maintains an advantage as it can be used to formally test coefficients estimated across functions.

If the estimation of the non-linear parameter θ was not included in either the multivariate or univariate analyses, then it is expected that the models would estimate the same corresponding coefficients .

(a)



(b)



F The code used for the multivariate repeated measures case study

```
## R set up #####

library(DImodelsMulti)

library(reshape2)

library(dplyr)

library(ggplot2)

library(PieGlyph)

## Read & format data

#####

biomass <- read.csv("biomass.csv")

biomass <- biomass[which(biomass$SITE == 1), c(4, 6, 8:9, 11:14, 16, 21:22)]

forage <- read.csv("forage_quality.csv")

forage <- forage[which(forage$SITE == 1), c(4, 6, 8:9, 11:14, 16, 21:22)]

##Use N_L (local) for nitrogen readings

forage$N <- forage$N_L

forage <- forage[, c(1:9, 12)]

Belgium <- merge(biomass, forage,

  by = c("YEARN", "HARVEST", "PLOT", "DENS", "TREAT", "G1", "G2", "L1", "L2"))

#Exclude treated plots

Belgium <- Belgium[which(Belgium$TREAT == 1), -5]

#Exclude Years 2 & 3

Belgium <- Belgium[which(Belgium$YEARN == 1), -1]

#Remove weed biomass from total to form sown biomass

Belgium$HARV_YIELD <- Belgium$HARV_YIELD - Belgium$WEED_Y

#Change response names & order

Belgium <- cbind(Belgium[, c(1:7, 9:10)], Belgium[, 8, drop = FALSE])

names(Belgium)[8:10] <- c("Sown", "N", "Weed")

#Transform factors

Belgium$DENS <- relevel(as.factor(Belgium$DENS), ref = "Low")

Belgium$HARVEST <- as.factor(Belgium$HARVEST)

#Transform weed biomass to weed suppression
```

```

Belgium$Weed <- -1*Belgium$Weed + max(Belgium$Weed)

#Transform to stacked data format

Belgium <- reshape2::melt(Belgium, id = colnames(Belgium)[1:7],
                          variable.name = "Var", value.name = "Y")

#Standardise responses

top <- Belgium %>%      #Top 10% (12) values for each EF at each time
  arrange(desc(Y)) %>%
  group_by(Var) %>%
  slice(1:12)

top <- aggregate(Y ~ Var, data = top, FUN = "mean") #Average of top values

Belgium$Y <- 100*Belgium$Y

#Sown

condition <- which(Belgium$Var == "Sown")

Belgium$Y[condition] <- Belgium$Y[condition] / top[1, "Y"]

#N

condition <- which(Belgium$Var == "N")

Belgium$Y[condition] <- Belgium$Y[condition] / top[2, "Y"]

#Weed

condition <- which(Belgium$Var == "Weed")

Belgium$Y[condition] <- Belgium$Y[condition] / top[3, "Y"]

##Fit simple model

#####

IDmodel_CS <- DImulti(y = "Y", eco_func = c("Var", "UN"), time = c("HARVEST", "CS"),
                     unit_IDs = "PLOT", prop = 4:7, DImodel = "ID",
                     method = "REML", data = Belgium)

## Compare autocorrelation structures

#####

IDmodel_AR1 <- DImulti(y = "Y", eco_func = c("Var", "UN"),
                      time = c("HARVEST", "AR1"), unit_IDs = "PLOT", prop = 4:7,
                      DImodel = "ID", method = "REML", data = Belgium)

IDmodel_UN <- DImulti(y = "Y", eco_func = c("Var", "UN"), time = c("HARVEST", "UN"),
                     unit_IDs = "PLOT", prop = 4:7, DImodel = "ID",

```

```

        method = "REML", data = Belgium)

AICc(IDmodel_CS); AICc(IDmodel_AR1); AICc(IDmodel_UN)

##Compare interaction structures

#####

IDmodel <- DImulti(y = "Y", eco_func = c("Var", "UN"), time = c("HARVEST", "UN"),
                  unit_IDs = "PLOT", prop = 4:7, DImodel = "ID",
                  method = "ML", data = Belgium)

AVmodel <- DImulti(y = "Y", eco_func = c("Var", "UN"), time = c("HARVEST", "UN"),
                  unit_IDs = "PLOT", prop = 4:7, DImodel = "AV",
                  method = "ML", data = Belgium)

anova(IDmodel, AVmodel)

## Estimate theta values

FULLmodel_theta <- DImulti(y = "Y", eco_func = c("Var", "UN"),
                           time = c("HARVEST", "UN"), unit_IDs = "PLOT",
                           prop = 4:7, DImodel = "FULL", estimate_theta = TRUE,
                           method = "ML", data = Belgium)

thetaEst <- FULLmodel_theta$theta

thetaSown<- c(thetaEst[1], 1, 1)

thetaN <- c(1, thetaEst[2], 1)

thetaWeed <- c(1, 1, thetaEst[3])

FULLmodel <- DImulti(y = "Y", eco_func = c("Var", "UN"), time = c("HARVEST", "UN"),
                    unit_IDs = "PLOT", prop = 4:7, DImodel = "FULL",
                    theta = 1, method = "ML", data = Belgium)

FULLmodel_thetaSown <- DImulti(y = "Y", eco_func = c("Var", "UN"),
                              time = c("HARVEST", "UN"), unit_IDs = "PLOT",
                              prop = 4:7, DImodel = "FULL", theta = thetaSown,
                              method = "ML", data = Belgium)

AICc(FULLmodel); AICc(FULLmodel_thetaSown)

##Estimated at bounds, set to 1

thetaEst[1] <- 1

FULLmodel_thetaN <- DImulti(y = "Y", eco_func = c("Var", "UN"),
                            time = c("HARVEST", "UN"), unit_IDs = "PLOT",

```

```

prop = 4:7, DImodel = "FULL", theta = thetaN,
method = "ML", data = Belgium)
AICc(FULLmodel); AICc(FULLmodel_thetaN)
thetaEst[2] <- 1
FULLmodel_thetaWeed <- DImulti(y = "Y", eco_func = c("Var", "UN"),
time = c("HARVEST", "UN"), unit_IDs = "PLOT",
prop = 4:7, DImodel = "FULL", theta = thetaWeed,
method = "ML", data = Belgium)
AICc(FULLmodel); AICc(FULLmodel_thetaWeed)
thetaEst[3] <- thetaWeed[3]
AVmodel_theta <- DImulti(y = "Y", eco_func = c("Var", "UN"),
time = c("HARVEST", "UN"), unit_IDs = "PLOT",
prop = 4:7, DImodel = "AV", theta = thetaEst,
method = "ML", data = Belgium)
AICc(AVmodel); AICc(AVmodel_theta)
#####
FGmodel <- DImulti(y = "Y", eco_func = c("Var", "UN"), time = c("HARVEST", "UN"),
unit_IDs = "PLOT", prop = 4:7, DImodel = "FG",
FG = c("Gr", "Gr", "Le", "Le"), theta = thetaEst,
method = "ML", data = Belgium)
anova(AVmodel_theta, FGmodel)
FULLmodel <- DImulti(y = "Y", eco_func = c("Var", "UN"), time = c("HARVEST", "UN"),
unit_IDs = "PLOT", prop = 4:7, DImodel = "FULL",
theta = thetaEst, method = "ML", data = Belgium)
anova(FGmodel, FULLmodel)
## Test DENS #####
FGmodel_DENS <- DImulti(y = "Y", eco_func = c("Var", "UN"),
time = c("HARVEST", "UN"), unit_IDs = "PLOT",
prop = 4:7, DImodel = "FG", extra_fixed = ~ DENS,
FG = c("Gr", "Gr", "Le", "Le"), theta = thetaEst,
method = "ML", data = Belgium)
anova(FGmodel, FGmodel_DENS)

```

```

## Refit Model #####
FinalModel <- DImulti(y = "Y", eco_func = c("Var", "UN"), time = c("HARVEST", "UN"),
  unit_IDs = "PLOT", prop = 4:7, DImodel = "FG", extra_fixed = ~ DENS,
  FG = c("Gr", "Gr", "Le", "Le"), theta = thetaEst,
  method = "REML", data = Belgium)

qqnorm(FinalModel)

## View Model and Predict
#####

print(FinalModel)

comms <- data.frame(PLOT = c(1, 2, 3, 4, 5, 6, 7, 8),
  G1 = c(1/4, 1/2, 1/2, 0, 1, 0, 0, 0),
  G2 = c(1/4, 0, 1/2, 0, 0, 1, 0, 0),
  L1 = c(1/4, 1/2, 0, 1/2, 0, 0, 1, 0),
  L2 = c(1/4, 0, 0, 1/2, 0, 0, 0, 1),
  MonoMix = c("Mix", "Mix", "Mix", "Mix", "Mono", "Mono", "Mono", "Mono")
)

comms$DENS <- "Low"

comms$PLOT <- as.factor(comms$PLOT)

preds_comms <- predict(FinalModel, newdata = comms)

preds_comms <- merge(comms, preds_comms)

preds_comms <- preds_comms %>% tidyr::separate("Ytype", sep = ":", into = c("Var", "HARVEST"))

piePlot_Low <- ggplot(data = preds_comms[which(preds_comms$PLOT %in% 1:8), ],
  aes(x = HARVEST, y = Yvalue, group = PLOT)) +
  geom_line(aes(linetype = MonoMix, linewidth = MonoMix, colour = PLOT), show.legend = FALSE) +
  scale_color_manual(values = c("#cc0001", "#fb940b", "#ffff01", "#01cc00", "#03c0c6", "#0000fe",
    "#762ca7", "#fe98bf")) +
  scale_linetype_manual(values = c("dashed", "solid")) +
  scale_linewidth_manual(values = c(1.25, 1.2)) +
  geom_pie_glyph(slices = c("G1", "G2", "L1", "L2"), radius = 0.25) +
  theme_bw() +

```

```

scale_fill_manual(values = c("#6a994e", "#a7c957", "#f9a03f", "#f3c053")) +
ylab("Ecosystem function response value (%)") +
xlab("Harvest Number") +
theme(text = element_text(size = 15)) +
facet_wrap(~Var, labeller = labeller(Var =
      c("N" = "Nitrogen Yield",
        "Sown" = "Sown Biomass",
        "Weed" = "Weed Suppression"))))

piePlot_Low

ggsave("Predicted Communities - LOW.png", plot = piePlot_Low,
      width = 10, height = 5, units = "in")

piePlot_Low_Wide <- ggplot(data = preds_comms[which(preds_comms$PLOT %in% 1:8), ],
      aes(x = HARVEST, y = Yvalue, group = Var), show.legend = FALSE) +
geom_line(linewidth = 1.2, aes(colour = Var), show.legend = TRUE) +
scale_color_manual(values = c("#cb66ff", "#669aff", "#ffcb66")) +
theme_bw() +
ylab("Ecosystem function response value (%)") +
xlab("Harvest Number") +
theme(text = element_text(size = 15)) +
facet_wrap(~PLOT, ncol = 2, labeller = labeller(PLOT =
      c("1" = "4-species Mixture",
        "2" = "2-species Mixture",
        "3" = "2-species Mixture",
        "4" = "2-species Mixture",
        "5" = "Monoculture",
        "6" = "Monoculture",
        "7" = "Monoculture",
        "8" = "Monoculture"))))

piePlot_Low_Wide

ggsave("Predicted Communities - LOW - Wide.png", plot = piePlot_Low_Wide,
      width = 8, height = 8, units = "in")

#####

```

G The results of each step in the realised proportions case study model selection process

The results of the model selection process for the LegacyNet analysis.

Models compared using likelihood ratio tests (LRT) for comparisons between models using different DI interaction structures and the corrected Akaike Information Criteria (AICc), for which a lower value is desirable, for comparisons between models using different values of the non-linear parameter θ (Kirwan *et al.*, 2009; Moral *et al.*, 2023). The FG interaction structure was used to estimate and test values of θ , as the FULL interaction structure could not be fit due to lack of data (Vishwakarma *et al.*, 2023).

Structure being tested	Test/Information criteria	Result
Interactions: ID vs. AV	Likelihood ratio: 178.02 p-value: < 0.001	AV
Estimate θ: DMyield – 0.01 VAWeed – 1.50	AICc: ($\theta = 1$ vs. $\theta \neq 1$): 12009.23 vs. 12640.66 12009.23 vs. 11970.01	DMyield – 1.00 VAWeed – 1.50
Interactions: AV vs. FG	Likelihood ratio: 26.44 p-value: 0.003	FG

H The estimated optimal communities for two metrics in the realised proportions case study

For the LegacyNet analysis, two measures for defining optimal functioning of an ecosystem were presented: ‘Opt (1)’ the sown community which resulted in the least change in species composition and ‘Opt (2)’ the sown community which produced the highest-level yield or biomass for sown species, i.e. excluding weeds. The species included in this case study were:

- ‘Grasses’: *Lolium perenne* (Lp) and *Phleum pratense* (Pp),
- ‘Legumes’: *Trifolium pratense* (Tp) and *Trifolium repens* (Tr), and
- ‘Herbs’: *Cichorium intybus* (Ci) and *Plantago lanceolata* (Pl)

The model given in Eq. 4.7 was used to produce estimates.

For ‘Opt (1)’, the estimated optimal sown community was a herb-dominated five-species mixture, comprised of 50.9% *Lp*, 0.0% *Pp*, 1.1% *Tp*, 0.8% *Tr*, 0.0% *Ci*, and 47.2% *Pl*.

For ‘Opt (2)’, the estimated optimal sown community was a six-species mixture, with four dominating species, at 0.0% *Lp*, 27.9% *Pp*, 40.6% *Tp*, 0.0% *Tr*, 0.0% *Ci*, and 31.5% *Pl*.

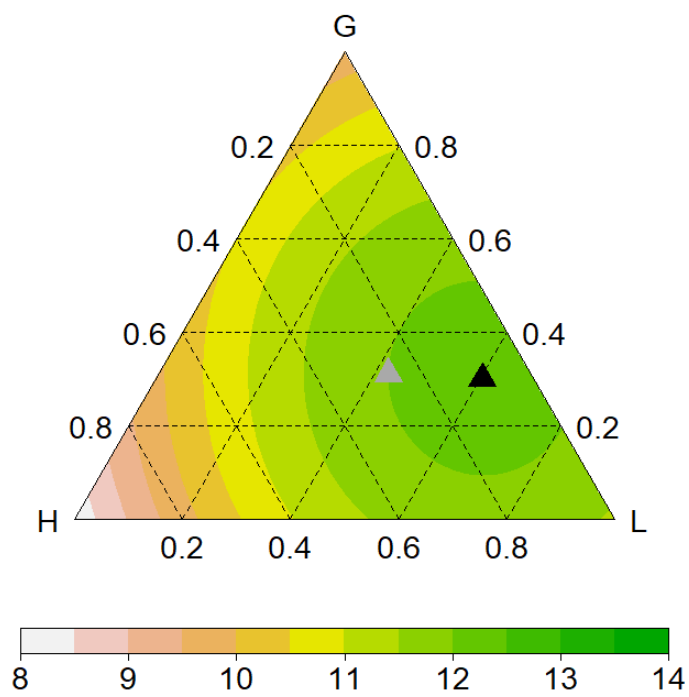
As these communities were selected using a linear regression model, they are quite robust, i.e. small changes to the species composition will likely not have a significant effect on the responses, therefore rare species may be excluded or included at a higher percentage with little trade-off. From these predictions, we can see that mixtures outperform monocultures in both metrics, suggesting that species interactions, particularly between functional groups, and perhaps niche partitioning, are important for species composition and proportional stability and biomass production.

I Ternary diagrams showcasing biomass predictions from the realised proportions case study

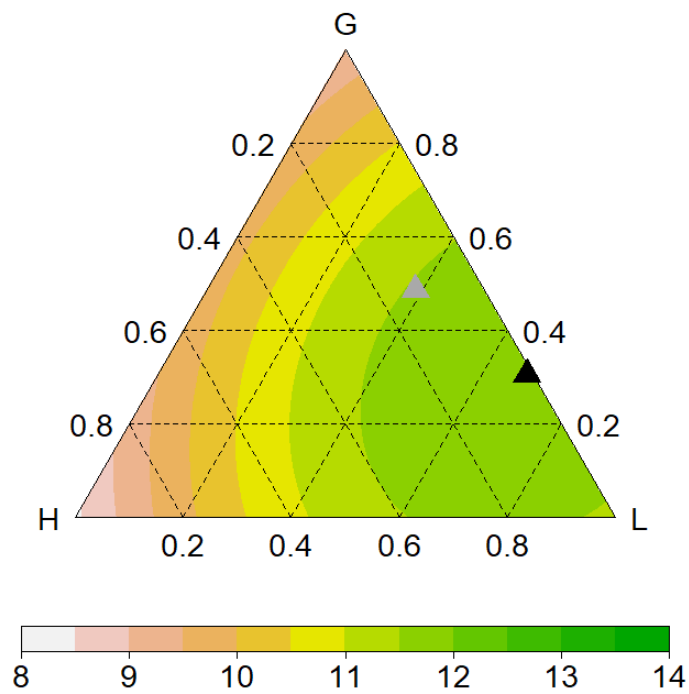
For the LegacyNet analysis, I use the model given in Eq. 4.7 to produce predictions across the simplex space. Here, I present two ternary diagrams, representing the predicted dry matter yield (biomass) produced for different compositions of the three sown functional groups (six species). For figure (a), the species within each functional group are sown at an even rate, e.g. if ‘Grass’ is sown at 50%, species *Lp* and *Pp* are sown at 25% each. For figure (b), the first presented species is used to represent each functional group, i.e. species *Lp* represents the grasses, *Tp* the legumes, and *Ci* the herbs. The optimal community for biomass production can be seen in each figure, represented by a black triangle, while the community’s realised proportions (ignoring unsown species) is represented as a grey triangle.

From these figures, we can see that not only does biomass benefit from multiple functional groups being present, with a high presence of legumes, but also from the presence of multiple species within each functional group. It also appears that the optimal communities in both figure (a) (which had an unshown weed proportion of 0.09) figure (b) (with a weed proportion of 0.02) are robust as the estimates surrounding the points do not differ greatly in the immediate surrounding area, although the optimal community found in figure (b) does appear to be less stable across time, i.e. the shift between the sown and realised communities is larger, than figure (a), indicating that there is either a difference in the ID effects of the two species within each functional group or that there are significant within-FG interactions at play.

(a) – Six-species:



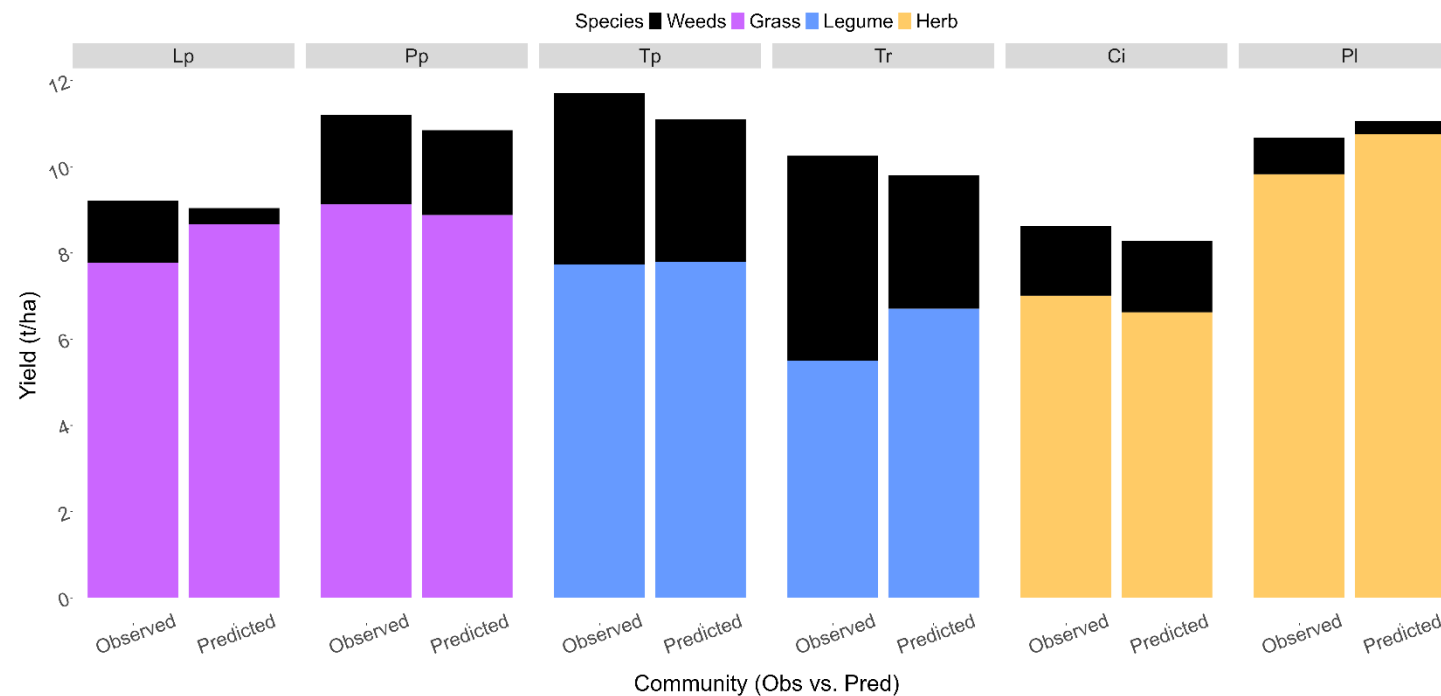
(b) – Three-species:



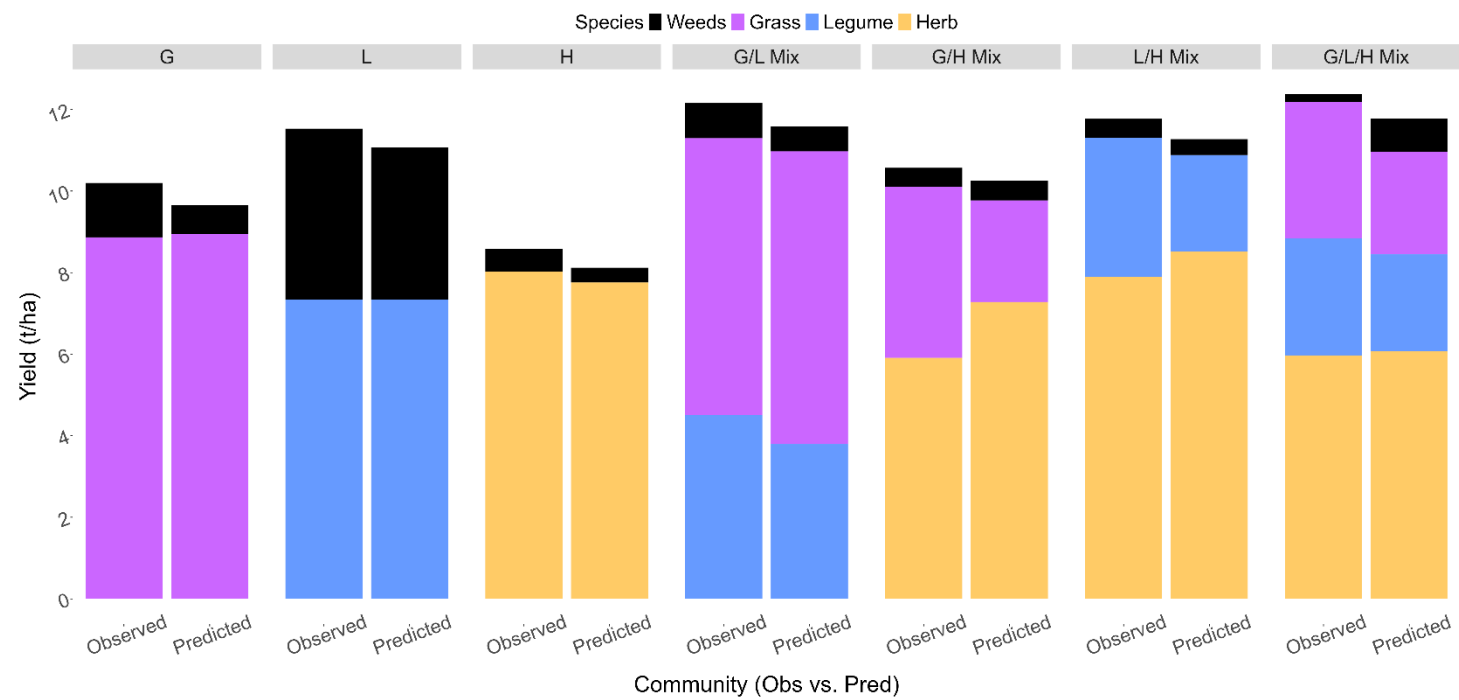
J Stacked barchart showcasing observed vs. predicted biomass from the realised proportions case study

For the LegacyNet analysis, I use the model given in Eq. 4.7 to produce predictions across the simplex space. Here, I present a grouped stacked barchart for comparison between the mean observed biomass from a selection of sown (a) monoculture and (b) mixture communities and the biomass predictions made from my model. As we can see, the fitting is imperfect, however, the story of the data is captured, with legume-only communities seeing larger levels of invasion by weeds while between functional group interactions in mixtures resulting in lower invasion levels.

(a) - Monocultures:



(b) - Mixtures:



K The code used for the realised proportions case study

The code included here is specific to the case study presented in section 4 of this thesis and may require adaptation for use with other datasets. Specifically, the realised proportion of unsown species must currently be grouped into a single proportion if species interactions are desired and the set of predictor variables must currently be in the form of a DI model. This code is intended to be generalised and included in the `DImodelsMulti` R package, presented in section 3 of this thesis, at a later time.

```
#DIach is a function which can be used to fit DI models to realised proportions

#yp is the column names or indices of the set of observed realised proportions

#yt is the column name or index of the total ecosystem function response

#thetaT and thetaW are the theta values assigned to the total and weed responses, respectively

#hessianEst and iterNum are options for the optim() function used to estimate the model coefficients using
MLE

#prop, DImodel, FG, ID, estimate_theta, and data follow the use of the parameters of the same name in the
DImodels (Moral et al., 2023) and DImodelsMulti (section 3) R packages

#extra_fixed is intended to be used as in the DImodelsMulti (section 3) R package, but is currently serves no
function

DIach <- function(yp, yt, prop, DImodel, FG = NULL, ID = NULL, thetaT = 1, thetaW = 1, estimate_theta =
FALSE, extra_fixed = NULL, data, hessianEst = FALSE, iterNum = 5000) {

  #if/else statements to ensure parameter inputs are correct

  if(DImodel == "STR"){

    stop("STR currently not implemented, please change the DImodel tag used")}

  if(!is.null(extra_fixed)){

    warning("extra_fixed has not yet been implemented and this input will be ignored")}

  extraFixed <- extra_fixed

  ##data

  if(missing(data)) {

    stop(crayon::bold(crayon::red("You must supply a dataframe through the argument 'data'.\n"))))}
```

```

if(inherits(data, 'tbl_df')){
  data <- as.data.frame(data)}

if(any(c("AV", "E", "NULL", "NA") %in% colnames(data))){
  stop(crayon::bold(crayon::red("You must not have any columns named \"AV\" or \"E\", \"NA\" or \"NULL\"
in your dataset provided through the parameter 'data'"))))}

if(any(startsWith(colnames(data), c("FULL.", "FG."))){
  stop(crayon::bold(crayon::red("You must not have any column names beginning with \"FULL.\" or \"FG.\"
in your dataset provided through the parameter 'data'"))))}

if(any(endsWith(colnames(data), c("_add"))){
  stop(crayon::bold(crayon::red("You must not have any column names ending with \"_add\" in your dataset
provided through the parameter 'data'"))))}

##y

if(missing(yp)){
  stop(crayon::bold(crayon::red("You must supply proportional response variable names or column indices
through the argument 'yp'.\n")))}

if(missing(yt)){
  stop(crayon::bold(crayon::red("You must supply a total response variable name or column index through the
argument 'yt'.\n")))}

#If positions are passed, change to names

if(is.numeric(yp[1])){
  yp <- colnames(data)[yp]}

if(is.numeric(yt[1])){
  yt <- colnames(data)[yt] }

#Check names are in dataset

if (!all(yp %in% colnames(data))){
  stop(crayon::bold(crayon::red("One or more of the columns referenced in the parameter 'yp' do not exist in
the dataset specified through the 'data' parameter.\n")))}

if (!all(yt %in% colnames(data))){
  stop(crayon::bold(crayon::red("The column referenced in the parameter 'yt' does not exist in the dataset
specified through the 'data' parameter.\n")))}

if (!all(is.numeric(as.matrix(data[, yp])))){

```

```

    stop(crayon::bold(crayon::red("You must supply numerical response variable column names or indices
through the argument 'yp'.\n"))))}

if (!all(is.numeric(as.matrix(data[, yt]))))){

  stop(crayon::bold(crayon::red("You must supply numerical response variable column names or indices
through the argument 'yt'.\n"))))}

##prop
if(missing(prop)){

  stop(crayon::bold(crayon::red("You must supply species proportion variable names or column indices
through the argument 'prop'.\n"))))}

if(length(prop) < 2){

  stop(crayon::bold(crayon::red("You must supply multiple species proportion variable names or column
indices through the argument 'prop'.\n"))))}

#If positions are passed, change to names
if(is.numeric(prop[1])){

  prop <- colnames(data)[prop]}

#Check names are in dataset
if(!all(prop %in% colnames(data))){

  stop(crayon::bold(crayon::red("One or more of the columns referenced in the parameter 'prop' do not exist
in the dataset specified through the 'data' parameter.\n"))))}

#Check that initial proportions sum to one
pi_sums <- apply(data[, prop], 1, sum)

if(any(pi_sums > 1.0001) | any(pi_sums < 0.9999) & !is.null(extraFixed)){

  warning(crayon::bold(crayon::red("One or more rows have species proportions that do not sum to 1. It is
assumed that this is by design with the missing proportions specified through the parameter 'extra_fixed', but
please ensure this.\n"))))}

else if(any(pi_sums > 1.0001) | any(pi_sums < 0.9999) & is.null(extraFixed)){

  stop(crayon::bold(crayon::red("One or more rows have species proportions that do not sum to 1. Please
correct this prior to analysis.\n"))))}

else if(any(pi_sums < 1 & pi_sums > 0.9999) | any(pi_sums > 1 & pi_sums < 1.0001)){

  Pi_sums <- apply(data[, prop], 1, sum)

  Pi_sums <- ifelse(Pi_sums == 0, 1, Pi_sums)

  data[, prop] <- data[, prop] / Pi_sums

```

```

warning(crayon::italic(crayon::red("One or more rows have species proportions that sum to approximately
1, but not exactly 1. This is typically a rounding issue, and has been corrected internally prior to analysis.\n"))))}

##DImodel

if(length(DImodel) != 1){

  stop(crayon::bold(crayon::red("You must enter a single model type through the parameter 'DImodel'")))}

if(!(DImodel %in% c("STR", "ID", "FULL", "E", "AV", "ADD", "FG"))){

  stop(crayon::bold(crayon::red("You have entered an unknown argument through the 'DImodel' parameter.
\nThe options for the 'DImodel' parameter are: \"STR\", \"ID\", \"FULL\", \"E\", \"AV\", \"ADD\", \"FG\"")))}

##FG

if(!is.null(FG)){

  cond0 <- length(FG) != length(prop)

  cond1 <- length(grep(":", FG)) > 0

  cond2 <- any(FG == "_")

  cond3 <- any(FG == "i")

  cond4 <- any(FG == "n")

  cond5 <- any(FG == "f")

  cond6 <- any(FG == "g")

  cond7 <- any(FG == "_i")

  cond8 <- any(FG == "in")

  cond9 <- any(FG == "nf")

  cond10 <- any(FG == "fg")

  cond11 <- any(FG == "g_")

  cond12 <- any(FG == "_in")

  cond13 <- any(FG == "inf")

  cond14 <- any(FG == "nfg")

  cond15 <- any(FG == "fg_")

  cond16 <- any(FG == "_inf")

  cond17 <- any(FG == "infg")

  cond18 <- any(FG == "nfg_")

  cond19 <- any(FG == "_infg")

  cond20 <- any(FG == "infg_")

  cond21 <- any(FG == "_infg_")

```

```

if(cond1 | cond2 | cond3 | cond4 | cond5 | cond6 | cond7 |
  cond8 | cond9 | cond10 | cond11 | cond12 | cond13 | cond14 |
  cond15 | cond16 | cond17 | cond18 | cond19 | cond20 |
  cond21){
  stop(crayon::bold(crayon::red("Please give your functional groups a different name.",
    " Names should not include colons (':'), or any single or multiple",
    " character combination of the expression '_infg_'.",
    " This expression is reserved for computing functional groups internally.")))}

if(cond0){
  stop(crayon::bold(crayon::red("The number of Functional Groups supplied does not match the number of
species supplied through 'prop'")))}

##extra_fixed

if(!is.null(extraFixed) & !plyr::is.formula(extraFixed)){
  stop(crayon::bold(crayon::red("You must supply a formula through the argument 'extra_fixed'\n")))}

#####

## Interaction Term checks

#DImodel parameter checks

STRflag <- if("STR" %in% DImodel) TRUE else FALSE
IDflag <- if("ID" %in% DImodel) TRUE else FALSE
FULLflag <- if("FULL" %in% DImodel) TRUE else FALSE
Eflag <- if("E" %in% DImodel) TRUE else FALSE
AVflag <- if("AV" %in% DImodel) TRUE else FALSE
ADDflag <- if("ADD" %in% DImodel) TRUE else FALSE
FGflag <- if("FG" %in% DImodel) TRUE else FALSE

#Conditions checks

Econdflag <- if(length(prop) > 2) TRUE else FALSE
AVcondflag <- if(length(prop) > 2) TRUE else FALSE
ADDcondflag <- if(length(prop) > 3) TRUE else FALSE
FGcondflag <- if(!is.null(FG)) TRUE else FALSE

#Match conditions with DImodel

if(Eflag && !Econdflag) {

```

```

    stop(crayon::bold(crayon::red("Less than 3 species present, E model will not be produced as it is
uninformative\n"))))}

else if(AVflag && !AVcondflag) {

  stop(crayon::bold(crayon::red("Less than 3 species present, AV model will not be produced as it is
uninformative\n"))))}

else if(ADDflag && !ADDcondflag){

  stop(crayon::bold(crayon::red("Less than 4 species present, ADD model will not be produced as it is
uninformative\n"))))}

else if(FGflag && !FGcondflag){

  stop(crayon::bold(crayon::red("No functional groups given, FG model will not be produced\n"))))}

else if(FGflag & any(!is.character(FG))){

  stop(crayon::bold(crayon::red("FG argument takes character strings with functional",

                                " group names referring to each species, in order"))))}

#####

## build matrices for compositional dependents (y) and independents/predictors (x)

ytName <- yt

yt <- data[, yt] # total response

ypNames <- yp

yp <- data[, yp] # proportional response

#Add _ID to each prop colname

propCols <- data.frame(data[, prop])

colnames(propCols) <- paste0(prop, "_ID")

data <- cbind(data[, which(!(colnames(data) %in% prop))], propCols)

if(!(DImodel %in% c("STR", "ID"))){ # Interactions?

  if(!estimate_theta){ #Create interaction columns with supplied theta values

    intColsT <- DImodels::DI_data(prop = colnames(propCols), FG = FG, data = data, theta = thetaT, what =
DImodel)

    intColsW <- DImodels::DI_data(prop = colnames(propCols), FG = FG, data = data, theta = thetaW, what
= DImodel)}

  else {#estimate theta

    #Find good starting values

```

```

thetaT <- suppressMessages(
  DImodels::DI(y = ytName, prop = colnames(propCols), FG = FG, data = data,
    estimate_theta = TRUE, DImodel = DImodel))$coefficients[["theta"]]
thetaW <- suppressMessages(
  DImodels::DI(y = ypNames[length(ypNames)], prop = colnames(propCols), FG = FG, data = data,
    estimate_theta = TRUE, DImodel = DImodel))$coefficients[["theta"]]
theta_Est <- function(thetaVals)
{
  thetaValT <- thetaVals[1]
  thetaValW <- thetaVals[2]

  fit <- DIach(yp = ypNames, yt = ytName, prop = colnames(propCols), DImodel = DImodel, FG = FG, ID
= ID, thetaT = thetaValT, thetaW = thetaValW, extra_fixed = extra_fixed, data = data)

  return(-as.numeric(logLik(fit))) } #minimise negative logLik (maximise logLik)
thetaVals <- stats::optim(c(thetaT, thetaW), theta_Est, hessian = FALSE,
  lower = rep(0.01, times = 2), upper = rep(1.5, times = 2),
  method = "L-BFGS-B", control = list(maxit = iterNum))$par
thetaT <- thetaVals[1]
intColsT <- DImodels::DI_data(prop = colnames(propCols), FG = FG, data = data, theta = thetaT, what =
DImodel)
thetaW <- thetaVals[2]
intColsW <- DImodels::DI_data(prop = colnames(propCols), FG = FG, data = data, theta = thetaW, what
= DImodel)}

#Change new column names
if(FULLflag){
  intColsT <- data.frame(intColsT)
  for(i in 1:ncol(intColsT)){
    colnames(intColsT)[i] <- paste0("FULL.", colnames(intColsT)[i]) }
  dataT <- cbind(data, intColsT)
  intColsW <- data.frame(intColsW)
  for(i in 1:ncol(intColsW)){
    colnames(intColsW)[i] <- paste0("FULL.", colnames(intColsW)[i]) }
  dataW <- cbind(data, intColsW) }

```

```

else if(FGflag){
  intColsT <- data.frame(intColsT)
  for(i in 1:ncol(intColsT)) {
    colnames(intColsT)[i] <- paste0("FG.", colnames(intColsT)[i]) }
  dataT <- cbind(data, intColsT)
  intColsW <- data.frame(intColsW)
  for(i in 1:ncol(intColsW)) {
    colnames(intColsW)[i] <- paste0("FG.", colnames(intColsW)[i]) }
  dataW <- cbind(data, intColsW)}
else if(AVflag | Eflag) {
  intColsT <- data.frame(intColsT)
  colnames(intColsT)[1] <- DImodel
  dataT <- cbind(data, intColsT)
  intColsW <- data.frame(intColsW)
  colnames(intColsW)[1] <- DImodel
  dataW <- cbind(data, intColsW) }
else if(ADDflag) {
  intColsT <- data.frame(intColsT)
  dataT <- cbind(data, intColsT)
  intColsW <- data.frame(intColsW)
  dataW <- cbind(data, intColsW)}

preds <- c(colnames(propCols), names(intColsT))

xT <- dataT[, preds] #Total yield predictors
xP <- dataW[, preds]} #Proportional response predictors
else if(DImodel == "ID"){
  preds <- colnames(propCols)
  xT <- xP <- data[, preds] }

#####

dm <- dim(yp)
D <- dm[2] ## Numbers of prop responses (C from algebraic model)
n <- dm[1] ## sample size

```

```

xiniT <- xT

xiniP <- xP

xT <- model.matrix(~., as.data.frame(xT) )

xP <- model.matrix(~., as.data.frame(xP) )

## Separate yp and x (0s and non-0s)

beta.iniP <- NULL

beta.iniT <- NULL

xT <- xT[, -1, drop = FALSE] #remove intercept column

xP <- xP[, -1, drop = FALSE]

##YT INITIAL VALUES #####

beta.iniTModel <- suppressMessages(DImodels::DI(y = ytName, prop = colnames(propCols), DImodel =
DImodel, FG = FG, theta = thetaT, #ID = ID,
                                data = data))

if(thetaT != 1) #remove theta coefficient{

  beta.iniT <- beta.iniTModel$coefficients[-length(beta.iniTModel$coefficients)] }

else{

  beta.iniT <- beta.iniTModel$coefficients}

#####

##YP INITIAL VALUES #####

for (i in 1:D){ # for each prop response

  tryCatch({

    #remove interactions

    xP_temp <- xP[, 1:ncol(propCols)]

    if(i == D){ #Weeds have interactions

      invisible(capture.output(suppressMessages({suppressWarnings({

        beta.iniP <- c(beta.iniP, Rfast::prop.reg(yp[, i], xP)$info[-1, 1]) } })))}

    else{ #Sown have no interactions

      invisible(capture.output(suppressMessages({suppressWarnings({

        beta.iniP <- c(beta.iniP, Rfast::prop.reg(yp[, i], xP_temp)$info[-1, 1]) } })))

      beta.iniP <- c(beta.iniP, rep(0, (ncol(xP) - ncol(xP_temp))))}, #pad with 0s for interactions

    error = function(e){

```

```
print("An error has occurred when estimating starting parameters for the proportional responses, setting
values equal to 1.")
```

```
beta.iniP <- c(beta.iniP, rep(1, (ncol(xP)))) } }
```

```
#Format into matrix
```

```
beta.iniP <- matrix(beta.iniP, ncol = D)
```

```
beta.iniP.mat <- beta.iniP #copy matrix version
```

```
#Format into vector for optim
```

```
beta.iniP <- c(beta.iniP) #by column
```

```
beta.iniP <- beta.iniP[beta.iniP != 0] #Remove 0s as they will not be estimated
```

```
#####
```

```
a1 <- which(Rfast::rowsums(yp > 0) == D)
```

```
a2 <- which(Rfast::rowsums(yp > 0) != D)
```

```
n1 <- length(a1) # sample size of the compositional vectors with no zeros
```

```
n2 <- n - n1 # sample size of the compositional vectors with zeros
```

```
za <- yp[a2, , drop = FALSE]
```

```
za[za == 0] <- 1
```

```
za[za < 1] <- 0
```

```
#Marginal prob constant
```

```
marg <- table(apply(za, 1, paste, collapse = ","))
```

```
marg <- as.vector(marg)
```

```
const <- n1 * log(n1/n) + sum(marg * log(marg/n))
```

```
yp1 <- yp[a1, , drop = FALSE]
```

```
lyp1 <- log(yp1)
```

```
x1 <- xP[a1, , drop = FALSE]
```

```
lyp2 <- log(yp[a2, , drop = FALSE])
```

```
x2 <- xP[a2, , drop = FALSE]
```

```
n1 <- nrow(yp1) ; n2 <- n - n1
```

```
##Sigma Initial Values#####
```

```
sigma.ini <- stats::cov(cbind(yt, yp))
```

```
sigma.ini <- as.matrix(Matrix::nearPD(sigma.ini)$mat)
```

```
sigma.ini <- sigma.ini[upper.tri(sigma.ini, diag = TRUE)]
```

```
#####
```

```

#optim parameters

beta.ini <- c(beta.iniT, beta.iniP, sigma.ini) # yt, yp, and sigma .ini's

z <- list(yp = yp,

  lyp1 = lyp1, #log of non-zero responses

  lyp2 = lyp2, #log of zero responses

  yt = yt,    #YTotal column

  x1 = x1,    #X columns/rows for yp1

  x2 = x2,    #X columns/rows for yp2

  xT = xT,    #X matrix with total biomass theta

  xP = xP,    #X matrix with unsown species theta

  a1 = a1,    #row numbers for x1 from X

  a2 = a2,    #row numbers for x2 from X

  beta.iniP.mat = beta.iniP.mat) #matrix version of beta.iniP

optimError <- FALSE

tryCatch({suppressWarnings({

  qa <- optim(beta.ini, .DImixreg, z = z, propCol = colnames(propCols), method = "Nelder-Mead",

    control = list(maxit = iterNum),

    hessian = hessianEst)})),

error = function(e){

  print("An error has occurred with optim, likely optimisation cannot begin with the initial parameters, here's

the message:")

  print(e)

  optimError <-<- TRUE })

if(optimError){ #return empty model object with large negative logLik

  model <- list(

    loglik = -10^10000,

    be = NA, thetaT = thetaT, thetaW = thetaW,

    hessian = NA, seb = NA,

    sigma = NA)

  attr(model, "ytName") <- ytName

  attr(model, "ypNames") <- ypNames

  attr(model, "beT") <- NA

```

```

attr(model, "beP")    <- NA

attr(model, "DImodel") <- DImodel

attr(model, "prop")    <- prop

attr(model, "FG")      <- FG

attr(model, "coefNum") <- length(beta.ini)

attr(model, "sampNum") <- nrow(data) * (D+1)

class(model) <- c("DIach", "list")

return(model)}

par <- qa$par #parameters of optim at MLE point

hessian <- qa$hessian #hessian if hessianEst = TRUE (from optim)

hessianNumDeriv <- numDeriv::hessian(f = .DImixreg, x = par, z = z, propCol = colnames(propCols))

#hessian using numDeriv

#####

#Separate par into components#####

beT <- par[1:ncol(xT)] #beta Total

names(beT) <- colnames(xT)

par <- par[-1:ncol(xT)]

beP <- par[1:length(beta.iniP)] #beta Proportions

beta.iniP.mat[which(beta.iniP.mat != 0)] <- beP #replace non-zero coefficients in matrix

beP <- beta.iniP.mat

rownames(beP) <- colnames(xP); colnames(beP) <- colnames(yp)

beP[-1:-length(propCols), -D] <- 0 #interactions are 0

par <- par[-1:-(length(beta.iniP))]

be <- cbind(beT, beP)

colnames(be) <- c(ytName, colnames(yp))

rownames(be) <- colnames(xT)

sigma <- matrix(0, nrow = (D + 1), ncol = (D + 1)) #Square matrix of 0s

sigma[upper.tri(sigma, diag = TRUE)] <- par

t.sigma <- t(sigma)

```

```

sigma[lower.tri(sigma)] <- t.sigma[lower.tri(t.sigma)]

colnames(sigma) <- c(ytName, colnames(yp))

rownames(sigma) <- c(ytName, colnames(yp))

diag(sigma) <- abs(diag(sigma)) #symmetrical vcov matrix for responses

#Format model to be returned

model <- list(

  loglik = -qa$value + const,

  be = be, thetaT = thetaT, thetaW = thetaW,

  hessian = hessian, hessianNumDeriv = hessianNumDeriv,

  sigma = sigma)

attr(model, "ytName") <- ytName

attr(model, "ypNames") <- ypNames

attr(model, "beT") <- be[, 1, drop = FALSE]

attr(model, "beP") <- be[, -1, drop = FALSE]

attr(model, "DImodel") <- DImodel

attr(model, "prop") <- prop

attr(model, "FG") <- FG

if(thetaT != 1 & thetaW != 1){

  attr(model, "coefNum") <- length(beta.ini) + 2}

else if(thetaT != 1 | thetaW != 1) {

  attr(model, "coefNum") <- length(beta.ini) + 1}

else{

  attr(model, "coefNum") <- length(beta.ini) }

attr(model, "sampNum") <- nrow(xT) * ncol(be)

class(model) <- c("DIach", "list")

return(model)}

#####

##The function optim() uses to maximise likelihood, adapted from zadr() function of Compositional R
package#####

.DImixreg <- function(param, z, propCol) {

  par <- param #save a copy of param input

  yt <- z$yt

```

```

yp <- z$yp
xT <- z$xT
xP <- z$xP
beta.iniP.mat <- z$beta.iniP.mat
bet <- matrix(par[1:(dim(xT)[2])], ncol = 1) # beta values for total response
par <- par[-1:-(dim(xT)[2])] #remove bet from par
mut <- xT %*% bet #estimates for total response
lyp1 <- z$lyp1
lyp2 <- z$lyp2
x1 <- z$x1
x2 <- z$x2
a1 <- z$a1
a2 <- z$a2
dm <- dim(lyp1)
d <- dm[2] # Number of prop responses (C algebraically)

## bep is the matrix of the zadr betas
bep <- par[1:(length(which(beta.iniP.mat != 0)))]
beta.iniP.mat[beta.iniP.mat != 0] <- bep
bep <- beta.iniP.mat
par <- par[-1:-(length(which(beta.iniP.mat != 0)))] #remove bep from par
## next we separate the compositional vectors, those which contain
## zeros and those without. The same separation is performed for the
## independent variable(s)
mu1 <- exp(x1 %*% bep) # w1 (with exp transformation, also avoids negs)
phi1 <- Rfast::rowsums(mu1) ## W for each row of fitted values
## next we find the fitted values for the compositional vectors with zeros
lyp3 <- lyp2
lyp3[sapply(lyp2, is.infinite)] <- 0
mu2 <- exp(x2 %*% bep) # w2
mu2[sapply(lyp2, is.infinite)] <- 0

```

```

phi2 <- Rfast::rowsums(mu2) # W (total shape param)

sigma <- matrix(0, nrow = (d+1), ncol = (d+1)) #Square matrix of 0s

sigma[upper.tri(sigma, diag = TRUE)] <- par

t.sigma <- t(sigma)

sigma[lower.tri(sigma)] <- t.sigma[lower.tri(t.sigma)]

sigma <- as.matrix(Matrix::nearPD(sigma)$mat) #vcov for responses

if(!matrixcalc::is.positive.definite(sigma)){ #if not PD matrix, return huge LL
  return(10^10000)}

mu12 <- rbind(mu1, mu2)

mu12 <- mu12[order(as.numeric(row.names(mu12))), ] #reorder matrix


mu12 <- t(apply(mu12, 1, function(i){i/sum(i)})) #change to proportions (yc not wc)

#standardise yp - mu12 residuals here if required

yAll <- cbind(yt, yp)

muAll <- cbind(mut, mu12)

{#Adapted code from dmnorm

log <- TRUE

tol <- 1e-06

if(is.vector(yAll)){
  yAll <- t(as.matrix(yAll)) }

n <- length(muAll)

if(is.vector(muAll)) {
  p <- length(muAll)

  if(is.matrix(yAll)) {
    muAll <- matrix(rep(muAll, nrow(yAll)), ncol = p, byrow = TRUE) }}

else{
  p <- ncol(muAll) }

if(!all(dim(sigma) == c(p, p)) || nrow(yAll) != nrow(muAll)) {
  stop("incompatible arguments")}

z <- t(yAll - muAll)

logdetS <- try(determinant(sigma, logarithm = TRUE)$modulus, silent = TRUE)

attributes(logdetS) <- NULL

```

```

iS <- MASS::ginv(sigma)

ssq <- diag(t(z) %*% iS %*% z)

loglik <- -(n * (log(2 * pi)) + logdetS + ssq)/2}

llt <- - sum(loglik) #LogLik line 1


llp_zero <- - sum( lgamma(phi2) ) + sum( lgamma(mu2[mu2>0]), na.rm = TRUE ) - sum( (mu2 - 1) * lyp3,
na.rm = TRUE ) #LogLik line 3

llp_nonzero <- - sum( lgamma(phi1) ) + sum( lgamma( mu1[mu1>0] ), na.rm = TRUE ) - sum( (mu1 - 1) *
lyp1, na.rm = TRUE ) #LogLik line 2

f <- llt + (llp_zero + llp_nonzero)

return(f)}

#####

##Useful functions#####

predict.DIach <- function(model, xnew, structZero = FALSE){

  if(structZero) { #structZero not yet implemented

    warning("Structural zero redistribution has not yet been implemented")

    structZero <- FALSE}

  singleRow <- FALSE

  if(nrow(xnew) == 1) { #If one row, double to avoid error

    singleRow <- TRUE

    xnew <- rbind(xnew, xnew)}

  beT <- attr(model, "beT")

  beP <- attr(model, "beP")

  ytName <- attr(model, "ytName")

  ypNames <- attr(model, "ypNames")

  ##Rebuild predictor columns

  #Add _ID to each prop colname

  propCols <- data.frame(xnew[, attr(model, "prop")])

  colnames(propCols) <- paste0(attr(model, "prop"), "_ID")

  xnewTemp <- cbind(xnew[, which(!(colnames(xnew) %in% attr(model, "prop")))], propCols)


  if(!(attr(model, "DImodel") %in% c("STR", "ID"))){ # Interactions?

```

```

intColsT <- DImodels::DI_data(prop = colnames(propCols), FG = attr(model, "FG"), data = xnewTemp,
theta = model$thetaT, what = attr(model, "DImodel"))

intColsW <- DImodels::DI_data(prop = colnames(propCols), FG = attr(model, "FG"), data = xnewTemp,
theta = model$thetaW, what = attr(model, "DImodel"))

#Change new column names
if(attr(model, "DImodel") == "FULL"){
  intColsT <- data.frame(intColsT)
  intColsW <- data.frame(intColsW)

  for(i in 1:ncol(intColsT)) {
    colnames(intColsT)[i] <- paste0("FULL.", colnames(intColsT)[i]) }

  xnewTempT <- cbind(xnewTemp, intColsT)

  for(i in 1:ncol(intColsW)) {
    colnames(intColsW)[i] <- paste0("FULL.", colnames(intColsW)[i]) }

  xnewTempW <- cbind(xnewTemp, intColsW) }
else if(attr(model, "DImodel") == "FG"){
  intColsT <- data.frame(intColsT)
  intColsW <- data.frame(intColsW)

  for(i in 1:ncol(intColsT)) {
    colnames(intColsT)[i] <- paste0("FG.", colnames(intColsT)[i])}

  xnewTempT <- cbind(xnewTemp, intColsT)

  for(i in 1:ncol(intColsW)) {
    colnames(intColsW)[i] <- paste0("FG.", colnames(intColsW)[i]) }

  xnewTempW <- cbind(xnewTemp, intColsW)}
else if(attr(model, "DImodel") == "AV" | attr(model, "DImodel") == "E"){
  intColsT <- data.frame(intColsT)
  intColsW <- data.frame(intColsW)

  colnames(intColsT)[1] <- attr(model, "DImodel")
  colnames(intColsW)[1] <- attr(model, "DImodel")

  xnewTempT <- cbind(xnewTemp, intColsT)

  xnewTempW <- cbind(xnewTemp, intColsW) }
else if(attr(model, "DImodel") == "ADD"){
  intColsT <- data.frame(intColsT)

```

```

intColsW <- data.frame(intColsW)

xnewTempT <- cbind(xnewTemp, intColsT)

xnewTempW <- cbind(xnewTemp, intColsW) }

preds <- c(colnames(propCols), names(intColsT))

xnewTempT <- xnewTempT[, preds]

xnewTempW <- xnewTempW[, preds] }

xnewTempT <- model.matrix(~., as.data.frame(xnewTempT))

xnewTempT <- xnewTempT[, -1, drop = FALSE]

xnewTempW <- model.matrix(~., as.data.frame(xnewTempW))

xnewTempW <- xnewTempW[, -1, drop = FALSE]

maP <- exp(xnewTempW %*% beP)

est <- maP / Rfast::rowsums(maP) ## fitted values

maT <- xnewTempT %*% beT

est <- cbind(maT, est)

colnames(est) <- c(ytName, ypNames)

est <- as.data.frame(est)

##if(structZero) #redistribute structural 0 values here


est <- cbind(xnew, est)

if(singleRow) { #Remove doubled row

  est <- est[1, 1:ncol(est), drop = FALSE] }

return(est)}

#####

logLik.DIach <- function(model){

  return(as.numeric(model$loglik))}

#####

LRT.DIach <- function(model1, model2){

  if(attr(model1, "coefNum") > attr(model2, "coefNum")){

    complex <- logLik.DIach(model1)

    complexMod <- model1

    nested <- logLik.DIach(model2)

    nestedMod <- model2}

```

```

else if(attr(model1, "coefNum") < attr(model2, "coefNum")){

  complex <- logLik.DIach(model2)

  complexMod <- model2

  nested <- logLik.DIach(model1)

  nestedMod <- model1}

else{

  stop("The number of parameters in each model is equal, LRT is usable only if models are nested")}

testStat <- -2 * (nested - complex)

dfDif <- attr(complexMod, "coefNum") - attr(nestedMod, "coefNum")

pVal <- stats::pchisq(testStat, df = dfDif, lower.tail = FALSE)

return(pVal)}

#####

AIC.DIach <- function(model){

  K <- attr(model, "coefNum")

  L <- model$loglik

  AIC <- 2*K - 2*L

  return(AIC)}

#####

AICc.DIach <- function(model){

  AIC <- AIC.DIach(model)

  n <- attr(model, "sampNum")

  K <- attr(model, "coefNum")

  AICc <- AIC + (2 * K^2 + 2 * n)/(n - K - 1)

  return(AICc)}

#####

#####

##Code used for the case study analysis#####

library(DImodels)

library(dplyr)

##Set up data#####

```

```

Gdata <- read.csv("Mix_achieved proportions_Harvest scale data.csv")[, c(-2, -7:-10, -17:-23, -25:-26, -34:-
40)]

Gdata <- Gdata[which(Gdata$HarvNb <= 187 & Gdata$Drought == 0 & Gdata$Nfert == 150), c(-2, -4:-5)]

#Keep all year 1 harvests

#Group species by functional groups

Gdata$SownG <- Gdata$SownLp + Gdata$SownPp
Gdata$SownL <- Gdata$SownRc + Gdata$SownWc
Gdata$SownH <- Gdata$SownCi + Gdata$SownPl

for(i in 10:16){
  Gdata[, i] <- as.numeric(Gdata[, i])}

Gdata$VAG <- Gdata$VALp + Gdata$VAPp
Gdata$VAL <- Gdata$VARc + Gdata$VAWc
Gdata$VAH <- Gdata$VACi + Gdata$VAPl
Gdata$DMyield <- as.numeric(Gdata$DMyield)
Gdata[is.na(Gdata)] <- 0

#Change props to continuous, sum harvests, change back to prop#####

Gdata$VAG_cont <- Gdata$VAG * Gdata$DMyield
Gdata$VAL_cont <- Gdata$VAL * Gdata$DMyield
Gdata$VAH_cont <- Gdata$VAH * Gdata$DMyield
Gdata$VAWeed_cont <- Gdata$VAWeed * Gdata$DMyield

Gdata$VAG <- 0
Gdata$VAL <- 0
Gdata$VAH <- 0
Gdata$VAWeed <- 0

Gdata_annual <- aggregate(
  cbind(VAG_cont, VAL_cont, VAH_cont, VAWeed_cont, DMyield) ~
  Plot + SownG + SownL + SownH +
  SownLp + SownPp + SownRc + SownWc + SownCi + SownPl,
  data = Gdata, FUN = "sum")

Gdata_annual$VAG <- Gdata_annual$VAG_cont / Gdata_annual$DMyield
Gdata_annual$VAL <- Gdata_annual$VAL_cont / Gdata_annual$DMyield
Gdata_annual$VAH <- Gdata_annual$VAH_cont / Gdata_annual$DMyield

```

```

Gdata_annual$VAWeed <- Gdata_annual$VAWeed_cont / Gdata_annual$DMyield

#Fix rounding

Gdata_annual[Gdata_annual == 0.167] <- 1/6

Gdata_annual[Gdata_annual == 0.334] <- 1/3

Gdata_annual[c(19, 24), "SownH"] <- 0.2

Gdata_annual[c(19, 24), "SownPl"] <- 0.1

Gdata_annual[c(19, 24), "SownCi"] <- 0.1

#Transform yield from kg to t

Gdata_annual$DMyield <- Gdata_annual$DMyield/1000

#####

##Fit and compare different DI model structures#####

#ID vs AV

GModel_ID <- DIach(yp = c(16:19), yt = 15, prop = 5:10, DImodel = "ID", data = Gdata_annual)

GModel_AV <- DIach(yp = c(16:19), yt = 15, prop = 5:10, DImodel = "AV", data = Gdata_annual)

#LRT.DIach(GModel_ID, GModel_AV) #Alternative

AICc.DIach(GModel_ID); AICc.DIach(GModel_AV)

#Estimate theta

GModel_AV_theta <- DIach(yp = c(16:19), yt = 15, prop = 5:10, DImodel = "AV", estimate_theta = TRUE,
data = Gdata_annual)

thetaEstT <- GModel_AV_theta$thetaT #0.0648120832051962

thetaEstW <- GModel_AV_theta$thetaW #0.517060486213655

GModel_AV_thetaT <- DIach(yp = c(16:19), yt = 15, prop = 5:10, DImodel = "AV", thetaT = thetaEstT, data
= Gdata_annual)

#LRT.DIach(GModel_AV, GModel_AV_thetaT) #Alternative

AICc.DIach(GModel_AV); AICc.DIach(GModel_AV_thetaT)

thetaEstT <- 1 #GModel_AV_theta$thetaT #Retain better value

GModel_AV_thetaW <- DIach(yp = c(16:19), yt = 15, prop = 5:10, DImodel = "AV", thetaW = thetaEstW, data
= Gdata_annual)

#LRT.DIach(GModel_AV, GModel_AV_thetaW) #Alternative

AICc.DIach(GModel_AV); AICc.DIach(GModel_AV_thetaW)

thetaEstW <- 1 #GModel_AV_theta$thetaW #Retain better value

```

```

#AV vs FG

GModel_FG <- DIach(yp = c(16:19), yt = 15, prop = 5:10, DImodel = "FG", thetaT = thetaEstT, thetaW =
thetaEstW,

                FG = c("Gr", "Gr", "Le", "Le", "He", "He"), data = Gdata_annual)

#LRT.DIach(GModel_AV, GModel_FG) #Alternative
AICc.DIach(GModel_AV); AICc.DIach(GModel_FG)

#Refit best model, attempt to estimate hessian using optim (hessianEst)
GModel <- DIach(yp = c(16:19), yt = 15, prop = 5:10, DImodel = "AV",

                thetaT = thetaEstT, thetaW = thetaEstW,

                FG = c("Gr", "Gr", "Le", "Le", "He", "He"), data = Gdata_annual)

                #, hessianEst = TRUE)

#View model

GModel$be
GModel$thetaT
GModel$thetaW
GModel$sigma

#For comparison with total yield univariate model

GModel_UniAuto <- autoDI(y = 15, prop = 5:10, FG = c("Gr", "Gr", "Le", "Le", "He", "He"), data =
Gdata_annual)

#Extract hessian (either optim or numDeriv version), see if standard errors are available
hessian <- as.matrix(GModel$hessianNumDeriv)

hessNames <- c(paste0(rownames(GModel$be), "_DMyield"),

              paste0(rownames(GModel$be)[1:6], "_VAG"),

              paste0(rownames(GModel$be)[1:6], "_VAL"),

              paste0(rownames(GModel$be)[1:6], "_VAH"),

              paste0(rownames(GModel$be), "_VAWeed"))

hessNames <- c(hessNames, rep("", times = (ncol(hessian) - length(hessNames))))

colnames(hessian) <- rownames(hessian) <- hessNames

hessian <- hessian[which(rownames(hessian) != ""), which(colnames(hessian) != "")]

hessian <- as.matrix(Matrix::nearPD(hessian)$mat)

```

```

sigmaFix <- solve(hessian)

sigmaFix <- as.matrix(Matrix::nearPD(sigmaFix)$mat)

seb <- sqrt(diag(sigmaFix))

temp <- GModel$be

temp[temp != 0] <- seb

seb <- temp

seb

#####

##'Best' Index by GA#####

library(GA)

#Max Sown Bio

p <- c(0, 0, 0, 0, 0, 0)

max_MF_GA <- function(xx){

  den <- 1 + sum(exp(xx[1:5]))

  for(i in 1:5) {

    p[i] <- exp(xx[i]) / den}

  p[6] <- 1 - sum(p[1:5])

  comms <- as.data.frame(t(p))

  names(comms) <- c("SownLp", "SownPp", "SownRc", "SownWc", "SownCi", "SownPl")

  preds <- predict.DIach(GModel, xnew = comms)

  preds$Val <- preds$DMyield - (preds$DMyield * preds$VAWeed)

  return(max(preds$Val))}

GA_max <- ga(type = "real-valued", fitness = max_MF_GA,

  lower = c(-20, -20, -20, -20, -20),

  upper = c(20, 20, 20, 20, 20), popSize = 50,

  pmutation = 0.2, pcrossover = 0.8, maxiter = 600)

psol_max <- rep(0, times = 6)

psol_max[1:5] <- exp(GA_max@solution) / (1 + sum(exp(GA_max@solution)))

psol_max[6] <- 1 - sum(psol_max[1:5])

psol_max <- round(psol_max, digits = 3)

best <- as.data.frame(t(psol_max))

```

```

names(best) <- c("SownLp", "SownPp", "SownRc", "SownWc", "SownCi", "SownPl")

best_preds2 <- predict.DIach(GModel, xnew = best)

best_preds2 #Opt(2) community

#Max Comp Stability

p <- c(0, 0, 0, 0, 0, 0)

max_MF_GA <- function(xx){
  den <- 1 + sum(exp(xx[1:5]))

  for(i in 1:5) {
    p[i] <- exp(xx[i]) / den}

  p[6] <- 1 - sum(p[1:5])

  comms <- as.data.frame(t(p))

  names(comms) <- c("SownLp", "SownPp", "SownRc", "SownWc", "SownCi", "SownPl")

  preds <- predict.DIach(GModel, xnew = comms)

  preds$Val <- preds$VAWeed + abs((preds$SownLp + preds$SownPp) - preds$VAG) + abs((preds$SownRc
+ preds$SownWc) - preds$VAL) + abs((preds$SownCi + preds$SownPl) - preds$VAH)

  return(max(-preds$Val))}

GA_max <- ga(type = "real-valued", fitness = max_MF_GA,
  lower = c(-20, -20, -20, -20, -20),
  upper = c(20, 20, 20, 20, 20), popSize = 50,
  pmutation = 0.2, pcrossover = 0.8, maxiter = 600)

psol_max <- rep(0, times = 6)

psol_max[1:5] <- exp(GA_max@solution) / (1 + sum(exp(GA_max@solution)))

psol_max[6] <- 1 - sum(psol_max[1:5])

psol_max <- round(psol_max, digits = 3)

best <- as.data.frame(t(psol_max))

names(best) <- c("SownLp", "SownPp", "SownRc", "SownWc", "SownCi", "SownPl")

best_preds1 <- predict.DIach(GModel, xnew = best)

best_preds1 #Opt(1) community

#####

comms <- data.frame(

```

```

"Community" = c("G", "L", "H", "G/L Mix", "G/H Mix", "L/H Mix", "G/L/H Mix"),
"Richness" = c("2", "2", "2", "4", "4", "4", "6"),
"SownLp" = c(0.5, 0, 0, 0.25, 0.25, 0, 1/6),
"SownPp" = c(0.5, 0, 0, 0.25, 0.25, 0, 1/6),
"SownRc" = c(0, 0.5, 0, 0.25, 0, 0.25, 1/6),
"SownWc" = c(0, 0.5, 0, 0.25, 0, 0.25, 1/6),
"SownCi" = c(0, 0, 0.5, 0, 0.25, 0.25, 1/6),
"SownPl" = c(0, 0, 0.5, 0, 0.25, 0.25, 1/6)
)

comms <- rbind(comms, c("Community" = "Opt (1)", "Richness" = "4", best_preds1[1:6]))
comms <- rbind(comms, c("Community" = "Opt (2)", "Richness" = "3", best_preds2[1:6]))

preds <- predict.DIach(GModel, xnew = comms)

colnames(preds)[10:13] <- c("Grass", "Legume", "Herb", "Weeds")

for(i in 1:nrow(preds)){ #Redistribute structural zeros
  if((preds[i, "SownLp"] == 0 & preds[i, "SownPp"] == 0) & preds[i, "Grass"] != 0) {
    preds[i, "Grass"] <- 0
    preds[i, "Legume"] <- preds[i, "Legume"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])
    preds[i, "Herb"] <- preds[i, "Herb"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])
    preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])}
  if((preds[i, "SownRc"] == 0 & preds[i, "SownWc"] == 0) & preds[i, "Legume"] != 0) {
    preds[i, "Legume"] <- 0
    preds[i, "Grass"] <- preds[i, "Grass"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])
    preds[i, "Herb"] <- preds[i, "Herb"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])
    preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])}
  if((preds[i, "SownCi"] == 0 & preds[i, "SownPl"] == 0) & preds[i, "Herb"] != 0) {
    preds[i, "Herb"] <- 0
    preds[i, "Legume"] <- preds[i, "Legume"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])
    preds[i, "Grass"] <- preds[i, "Grass"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])
    preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])}
}

preds_stack <- reshape2::melt(preds,
  id.vars = c("Community", "Richness", "SownLp", "SownPp", "SownRc", "SownWc",
    "SownCi", "SownPl", "DMyield"),

```

```

        variable.name = "Species",

        value.name = "Percent")

preds_stack$Yield <- preds_stack$DMyield * preds_stack$Percent

#####

library(ggplot2)

preds_stack$Community <- factor(preds_stack$Community,
                                levels = c("G", "L", "H", "G/L Mix", "G/H Mix", "L/H Mix", "G/L/H Mix", "Opt (2)",
                                "Opt (1)"))

preds_stack$Species <- relevel(preds_stack$Species, 'Weeds')

plotBreak <- ggplot(preds_stack, aes(fill = Species, y = Yield, x = Community)) +
  geom_bar(position = "stack", stat = "identity") +
  facet_grid(~Richness, scales = "free_x", space = "free_x") +
  theme(axis.text = element_text(size = 24, angle = 20),
        axis.title = element_text(size = 28),
        legend.position="top",
        legend.text = element_text(size = 24),
        legend.title = element_text(size = 24),
        strip.text.x = element_text(size = 24),
        panel.border = element_blank(),
        panel.grid.major = element_blank(),
        panel.grid.minor = element_blank(),
        panel.background = element_rect(fill = "white"),
        panel.spacing = unit(0.5, "cm")) +
  labs(x = "Community", y = "Yield (t/ha)") +
  scale_y_continuous(breaks = round(seq(min(preds_stack$Yield),
                                       max(preds_stack$Yield*2),
                                       by = 2), 1)) +
  scale_fill_manual(values = c("black", "#cb66ff", "#669aff", "#ffcb66"))
plotBreak #Plot of selected community yield, broken down by species props
ggsave("LegacyNetIE2_Break.png", plot = plotBreak,

```

```
width = 12, height = 8, units = "in")
```

```
#####
```

```
comms <- data.frame(

  "L_pc"      = c("0%", "10%", "20%", "30%", "40%", "50%", "60%", "70%", "80%", "90%", "100%"),
  "SownLp"    = c(0.25, 0.225, 0.20, 0.175, 0.15, 0.125, 0.10, 0.075, 0.05, 0.025, 0.00),
  "SownPp"    = c(0.25, 0.225, 0.20, 0.175, 0.15, 0.125, 0.10, 0.075, 0.05, 0.025, 0.00),
  "SownRc"    = c(0.00, 0.050, 0.10, 0.150, 0.20, 0.250, 0.30, 0.350, 0.40, 0.450, 0.50),
  "SownWc"    = c(0.00, 0.050, 0.10, 0.150, 0.20, 0.250, 0.30, 0.350, 0.40, 0.450, 0.50),
  "SownCi"    = c(0.25, 0.225, 0.20, 0.175, 0.15, 0.125, 0.10, 0.075, 0.05, 0.025, 0.00),
  "SownPl"    = c(0.25, 0.225, 0.20, 0.175, 0.15, 0.125, 0.10, 0.075, 0.05, 0.025, 0.00)
)

preds <- predict.DIach(GModel, xnew = comms)

colnames(preds)[9:12] <- c("Grass", "Legume", "Herb", "Weeds")

for(i in 1:nrow(preds)){ #Redistribute structural zeros

  if((preds[i, "SownLp"] == 0 & preds[i, "SownPp"] == 0) & preds[i, "Grass"] != 0) {

    preds[i, "Grass"] <- 0

    preds[i, "Legume"] <- preds[i, "Legume"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])
    preds[i, "Herb"] <- preds[i, "Herb"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])
    preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])

    if((preds[i, "SownRc"] == 0 & preds[i, "SownWc"] == 0) & preds[i, "Legume"] != 0) {

      preds[i, "Legume"] <- 0

      preds[i, "Grass"] <- preds[i, "Grass"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])
      preds[i, "Herb"] <- preds[i, "Herb"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])
      preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])

      if((preds[i, "SownCi"] == 0 & preds[i, "SownPl"] == 0) & preds[i, "Herb"] != 0) {

        preds[i, "Herb"] <- 0

        preds[i, "Legume"] <- preds[i, "Legume"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])
        preds[i, "Grass"] <- preds[i, "Grass"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])
        preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])
      }
    }
  }
}
```

```

preds_stack <- reshape2::melt(preds,
                              id.vars = c("L_pc", "SownLp", "SownPp", "SownRc", "SownWc", "SownCi", "SownPl",
                              "DMyield"),
                              variable.name = "Species",
                              value.name = "Percent")

preds_stack$Yield <- preds_stack$DMyield * preds_stack$Percent

#####

preds_stack$L_pc <- factor(preds_stack$L_pc,
                           levels = c("0%", "10%", "20%", "30%", "40%", "50%", "60%", "70%", "80%", "90%",
                           "100%"))

preds_stack$Species <- relevel(preds_stack$Species, 'Weeds')


plotLeg <- ggplot(preds_stack, aes(fill = Species, y = Yield, x = L_pc)) +
  geom_bar(position = "stack", stat = "identity") +
  theme(axis.text = element_text(size = 24, angle = 20),
        axis.title = element_text(size = 28),
        legend.position="top",
        legend.text = element_text(size = 24),
        legend.title = element_text(size = 24),
        strip.text.x = element_text(size = 24),
        panel.border = element_blank(),
        panel.grid.major = element_blank(),
        panel.grid.minor = element_blank(),
        panel.background = element_rect(fill = "white"),
        panel.spacing = unit(0.5, "cm")) +
  labs(x = "Legume Presence", y = "Yield (t/ha)") +
  scale_y_continuous(breaks = round(seq(min(preds_stack$Yield),
                                       max(preds_stack$Yield*2),
                                       by = 2),1)) +
  scale_fill_manual(values = c("black", "#cb66ff", "#669aff", "#ffcb66"))

plotLeg #Plot showing change in yield and breakdown as legumes increase

```

```
ggsave("LegacyNetIE2_Leg.png", plot = plotLeg,  
       width = 12, height = 8, units = "in")
```

```
#####
```

L The code used for the ternary diagrams in Appendix I using the realised proportions case study model

```
library(lattice)

lattice.options(default.theme = standard.theme(color = FALSE))

#Generate variables base and high (maintaining correct relative values) for graphing
#Values 1000 and 870 can be adjusted to make code run quicker but with less resolution

trian <- expand.grid(base = seq(0, 1, l = 1000*2), high = seq(0, sin(pi/3),l = 870*2))

trian <- subset(trian, (base * sin(pi/3) * 2) > high)

trian <- subset(trian, ((1 - base) * sin(pi/3) * 2) > high)

#Create the new proportions from base and high in trian

trian$Sp_G <- trian$high * 2/sqrt(3)

trian$Sp_L <- trian$base - trian$high/sqrt(3)

trian$Sp_H <- 1 - trian$Sp_G - trian$Sp_L

#Predict values for y at proportions p1, p2, p3 using the coefficients from GModel

yhats <- predict(GModel, xnew = data.frame("SownLp" = trian$Sp_G*0.5,

      "SownPp" = trian$Sp_G*0.5,

      "SownRc" = trian$Sp_L*0.5,

      "SownWc" = trian$Sp_L*0.5,

      "SownCi" = trian$Sp_H*0.5,

      "SownPI" = trian$Sp_H*0.5))

#OR#####

yhats <- predict(GModel, xnew = data.frame("SownLp" = trian$Sp_G,

      "SownPp" = 0,

      "SownRc" = trian$Sp_L,

      "SownWc" = 0,

      "SownCi" = trian$Sp_H,

      "SownPI" = 0))

#####

trian$yhat <- yhats[, 7]
```

```
#Setup the basic contour diagram
```

```
grade.trellis <- function(from = 0.2, to = 0.8, step = 0.2, col = 1, lty = 2, lwd = 0.5){  
  x1 <- seq(from, to, step)  
  x2 <- x1/2  
  y2 <- x1 * sqrt(3)/2  
  x3 <- (1 - x1) * 0.5 + x1  
  y3 <- sqrt(3)/2 - x1 * sqrt(3)/2  
  panel.segments(x1, 0, x2, y2, col = col, lty = lty, lwd = lwd)  
  panel.text(x1, 0, label = x1, pos = 1, cex = 1.75)  
  panel.segments(x1, 0, x3, y3, col = col, lty = lty, lwd = lwd)  
  panel.text(x2, y2, label = rev(x1), pos = 2, cex = 1.75)  
  panel.segments(x2, y2, 1 - x2, y2, col = col, lty = lty, lwd = lwd)  
  panel.text(x3, y3, label = rev(x1), pos = 4, cex = 1.75)  
}
```

```
#Colour for the ternary diagram
```

```
revterrain.colors <- function (n, alpha = 1) {  
  if ((n <- as.integer(n[1L])) > 0) {  
    k <- n %/% 2  
    h <- rev(c(4/12, 2/12, 0/12))  
    s <- rev(c(1, 1, 0))  
    v <- rev(c(0.65, 0.90, 0.95))  
    c(hsv(h = seq.int(h[1L], h[2L], length.out = k),  
s = seq.int(s[1L], s[2L], length.out = k), v = seq.int(v[1L], v[2L], length.out = k), alpha = alpha), hsv(h =  
seq.int(h[2L], h[3L], length.out = n - k + 1)[-1L], s = seq.int(s[2L], s[3L], length.out = n - k + 1)[-1L], v =  
seq.int(v[2L], v[3L], length.out = n - k + 1)[-1L], alpha = alpha))}  
  else character() }
```

```
levelplot(yhat ~ base * high, trian, aspect = "iso",  
  xlim = c(-0.1, 1.1), ylim = c(-0.1, 0.96),  
  xlab = NULL, ylab = NULL, contour = TRUE,  
  col.regions = revterrain.colors,  
  colorkey = list(space = "bottom", axis.line = list(col = 1),
```

```

axis.text = list(col = 1, cex = 1.75)),

at = seq(8, 14, 0.5), #use different values if you want different number of contours

par.settings = list(axis.line = list(col = NA), axis.text = list(col = NA)),

panel = function(..., at, contour = TRUE, labels = TRUE){

  panel.levelplot(..., at = at, labels = labels,

    lty = 3, lwd = 0.01, col = 1, cex = 1.75)})

#Put labels for each species on the vertices of the basic contour diagram

trellis.focus("panel", 1, 1, highlight = FALSE)

panel.segments(c(0, 0, 0.5), c(0, 0, sqrt(3)/2), c(1, 1/2, 1), c(0, sqrt(3)/2, 0))

grade.trellis()

panel.text(0, 0, label = "H", pos = 2, cex = 1.75)

panel.text(1/2, sqrt(3)/2, label = "G", pos = 3, cex = 1.75)

panel.text(1, 0, label = "L", pos = 4, cex = 1.75)

xx <- c(0,0,0)

p <- c(0,0,0)

# Set a value for the intercept

intval <- 1

fitness1 <- function(xx){

  den = 1 + sum(exp(xx[1:2]))

  for(i in 1:2){ p[i] = exp(xx[i])/den }

  p[3] = 1 - sum(p[1:2])

#Predict from GModel

ys <- predict(GModel, xnew = data.frame("SownLp" = p[1]*0.5,

  "SownPp" = p[1]*0.5,

  "SownRc" = p[2]*0.5,

  "SownWc" = p[2]*0.5,

  "SownCi" = p[3]*0.5,

  "SownPI" = p[3]*0.5))

#OR#####

ys <- predict(GModel, xnew = data.frame("SownLp" = p[1],

  "SownPp" = 0,

```

```

        "SownRc"  = p[2],
        "SownWc"  = 0,
        "SownCi"  = p[3],
        "SownPl"  = 0))

##### value <- ys[, 7]

}

#Use the GA function to find the maximum value of the fitness function

GA <- ga(type = "real-valued", fitness = fitness1,
        min = c(-20,-20), max = c(20,20), popSize = 50,
        pmutation = 0.2, pcrossover = 0.8, maxiter = 600)#, monitor = plot)

#Calculate the solution in terms of proportions

psol = rep(0, times = 3)

psol[1:2] = exp(GA@solution)/(1 + sum(exp(GA@solution)))

psol[3] = 1 - sum(psol[1:2])

psol

#Predict from this optimum point

opt <- predict(GModel, xnew = data.frame("SownLp"  = psol[1]*0.5,
        "SownPp"  = psol[1]*0.5,
        "SownRc"  = psol[2]*0.5,
        "SownWc"  = psol[2]*0.5,
        "SownCi"  = psol[3]*0.5,
        "SownPl"  = psol[3]*0.5))

#OR#####

opt <- predict(GModel, xnew = data.frame("SownLp"  = psol[1],
        "SownPp"  = 0,
        "SownRc"  = psol[2],
        "SownWc"  = 0,
        "SownCi"  = psol[3],
        "SownPl"  = 0))

#####

```

```
opt <- opt[8:10]

pres <- opt / sum(opt)

#Add in the optimum points (sown and realised)

panel.points(x = (0.5 + (psol[2])/2 - (psol[3])/2),
             y = sin(pi/3) * (psol[1]), col = "black", pch = 17, cex = 2.5)

panel.points(x = (0.5 + (pres[2])/2 - (pres[3])/2),
             y = sin(pi/3)*(pres[1]), col = "darkgrey", pch = 17, cex = 2.5)
```

M The code used for the comparative plots in Appendix J using the realised proportions case study model

```
#Mixtures
```

```
comms <- data.frame(

  "Community" = c("G", "L", "H", "G/L Mix", "G/H Mix", "L/H Mix", "G/L/H Mix"),

  "Richness" = c("2", "2", "2", "4", "4", "4", "6"),

  "SownLp" = c(0.5, 0, 0, 0.25, 0.25, 0, 1/6),

  "SownPp" = c(0.5, 0, 0, 0.25, 0.25, 0, 1/6),

  "SownRc" = c(0, 0.5, 0, 0.25, 0, 0.25, 1/6),

  "SownWc" = c(0, 0.5, 0, 0.25, 0, 0.25, 1/6),

  "SownCi" = c(0, 0, 0.5, 0, 0.25, 0.25, 1/6),

  "SownPl" = c(0, 0, 0.5, 0, 0.25, 0.25, 1/6))

preds <- predict.DIach(GModel, xnew = comms)

colnames(preds)[10:13] <- c("Grass", "Legume", "Herb", "Weeds")

for(i in 1:nrow(preds)){ #Redistribute structural 0 proportions

  if((preds[i, "SownLp"] == 0 & preds[i, "SownPp"] == 0) & preds[i, "Grass"] != 0) {

    preds[i, "Grass"] <- 0

    preds[i, "Legume"] <- preds[i, "Legume"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])

    preds[i, "Herb"] <- preds[i, "Herb"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])

    preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])}

  if((preds[i, "SownRc"] == 0 & preds[i, "SownWc"] == 0) & preds[i, "Legume"] != 0) {

    preds[i, "Legume"] <- 0

    preds[i, "Grass"] <- preds[i, "Grass"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])

    preds[i, "Herb"] <- preds[i, "Herb"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])

    preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])}

  if((preds[i, "SownCi"] == 0 & preds[i, "SownPl"] == 0) & preds[i, "Herb"] != 0) {

    preds[i, "Herb"] <- 0

    preds[i, "Legume"] <- preds[i, "Legume"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])

    preds[i, "Grass"] <- preds[i, "Grass"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])}

}
```

```

  preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])} }
preds_stack <- reshape2::melt(preds,
  id.vars = c("Community", "Richness", "SownLp", "SownPp", "SownRc", "SownWc",
    "SownCi", "SownPl", "DMyield"),
  variable.name = "Species",
  value.name = "Percent")
preds_stack$Yield <- preds_stack$DMyield * preds_stack$Percent
preds_stack$Y <- "Predicted"
preds_stack <- preds_stack[, -2]
Gdata_annual_obs <- Gdata_annual[c(4:5, 12:15, 28:36), ]
Gdata_annual_obs$Community <- c("G", "G", "G/L Mix", "G/L Mix", "L", "L", "G/L/H Mix", "G/L/H Mix",
  "G/L/H Mix", "G/H Mix", "G/H Mix", "L/H Mix", "L/H Mix", "H", "H")
Gdata_annual_obs <- Gdata_annual_obs[, c(5:10, 15:20)]
colnames(Gdata_annual_obs)[8:11] <- c("Grass", "Legume", "Herb", "Weeds")
Gobs_stack <- reshape2::melt(Gdata_annual_obs,
  id.vars = c("Community", "SownLp", "SownPp", "SownRc", "SownWc", "SownCi",
    "SownPl", "DMyield"),
  variable.name = "Species",
  value.name = "Percent")
Gobs_stack$Yield <- Gobs_stack$DMyield * Gobs_stack$Percent
Gobs_stack$Y <- "Observed"
obsPreds <- rbind(preds_stack, Gobs_stack)
obsPreds_ag <- aggregate(Yield ~ Community + Y + Species, data = obsPreds, FUN = "mean")
obsPreds_ag$Community <- factor(obsPreds_ag$Community,
  levels = c("G", "L", "H", "G/L Mix", "G/H Mix", "L/H Mix", "G/L/H Mix"))
obsPreds_ag$Species <- relevel(obsPreds_ag$Species, 'Weeds')

plotComp <- ggplot(obsPreds_ag, aes(fill = Species, y = Yield, x = Y)) +
  geom_bar(position = "stack", stat = "identity") +
  facet_grid(~ Community, scales = "free_x", space = "free_x") +
  theme(axis.text = element_text(size = 24, angle = 20),
    axis.title = element_text(size = 28),

```

```

legend.position="top",

legend.text = element_text(size = 24),

legend.title = element_text(size = 24),

strip.text.x = element_text(size = 24),

panel.border = element_blank(),

panel.grid.major = element_blank(),

panel.grid.minor = element_blank(),

panel.background = element_rect(fill = "white"),

panel.spacing = unit(0.5, "cm")) +

labs(x = "Community (Obs vs. Pred)", y = "Yield (t/ha)") +

scale_y_continuous(breaks = round(seq(min(preds_stack$Yield),

                                     max(preds_stack$Yield*2),

                                     by = 2),1)) +

scale_fill_manual(values = c("black", "#cb66ff", "#669aff", "#ffcb66"))

plotComp #Plot comparing the selected mixture plots from the original experiment

ggsave("LegacyNetIE2_Comp.png", plot = plotComp,

       width = 25, height = 12, units = "in")

#Monocultures

comms <- data.frame(

  "Community" = c("Lp", "Pp", "Tp", "Tr", "Ci", "Pl"),

  "SownLp" = c(1, 0, 0, 0, 0, 0),

  "SownPp" = c(0, 1, 0, 0, 0, 0),

  "SownRc" = c(0, 0, 1, 0, 0, 0),

  "SownWc" = c(0, 0, 0, 1, 0, 0),

  "SownCi" = c(0, 0, 0, 0, 1, 0),

  "SownPl" = c(0, 0, 0, 0, 0, 1))

preds <- predict.DIach(GModel, xnew = comms)

colnames(preds)[9:12] <- c("Grass", "Legume", "Herb", "Weeds")

for(i in 1:nrow(preds)){ #Redistribute structural 0 proportions

  if((preds[i, "SownLp"] == 0 & preds[i, "SownPp"] == 0) & preds[i, "Grass"] != 0){

    preds[i, "Grass"] <- 0

    preds[i, "Legume"] <- preds[i, "Legume"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])
  }
}

```

```

preds[i, "Herb"] <- preds[i, "Herb"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])
preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Legume"] + preds[i, "Herb"] + preds[i, "Weeds"])}
if((preds[i, "SownRc"] == 0 & preds[i, "SownWc"] == 0) & preds[i, "Legume"] != 0) {
  preds[i, "Legume"] <- 0
  preds[i, "Grass"] <- preds[i, "Grass"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])
  preds[i, "Herb"] <- preds[i, "Herb"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])
  preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Grass"] + preds[i, "Herb"] + preds[i, "Weeds"])}
if((preds[i, "SownCi"] == 0 & preds[i, "SownPl"] == 0) & preds[i, "Herb"] != 0) {
  preds[i, "Herb"] <- 0
  preds[i, "Legume"] <- preds[i, "Legume"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])
  preds[i, "Grass"] <- preds[i, "Grass"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])
  preds[i, "Weeds"] <- preds[i, "Weeds"]/(preds[i, "Legume"] + preds[i, "Grass"] + preds[i, "Weeds"])}
preds_stack <- reshape2::melt(preds,
  id.vars = c("Community", "SownLp", "SownPp", "SownRc", "SownWc", "SownCi",
    "SownPl", "DMyield"),
  variable.name = "Species",
  value.name = "Percent")
preds_stack$Yield <- preds_stack$DMyield * preds_stack$Percent
preds_stack$Y <- "Predicted"
Gdata_annual_obs <- Gdata_annual[c(1:3, 6:8, 9:11, 16:18, 21:23, 38:40), ]
Gdata_annual_obs$Community <- rep(c("Lp", "Pp", "Tp", "Tr", "Ci", "Pl"), each = 3)
Gdata_annual_obs <- Gdata_annual_obs[, c(5:10, 15:20)]
colnames(Gdata_annual_obs)[8:11] <- c("Grass", "Legume", "Herb", "Weeds")
Gobs_stack <- reshape2::melt(Gdata_annual_obs,
  id.vars = c("Community", "SownLp", "SownPp", "SownRc", "SownWc", "SownCi",
    "SownPl", "DMyield"), variable.name = "Species", value.name = "Percent")
Gobs_stack$Yield <- Gobs_stack$DMyield * Gobs_stack$Percent
Gobs_stack$Y <- "Observed"
obsPreds <- rbind(preds_stack, Gobs_stack)
obsPreds_ag <- aggregate(Yield ~ Community + Y + Species, data = obsPreds, FUN = "mean")

```

```

obsPreds_ag$Community <- factor(obsPreds_ag$Community,
                                levels = c("Lp", "Pp", "Tp", "Tr", "Ci", "Pl"))
obsPreds_ag$Species <- relevel(obsPreds_ag$Species, 'Weeds')

plotCompMono <- ggplot(obsPreds_ag, aes(fill = Species, y = Yield, x = Y)) +
  geom_bar(position = "stack", stat = "identity") +
  facet_grid(~Community, scales = "free_x", space = "free_x") +
  theme(axis.text = element_text(size = 24, angle = 20),
        axis.title = element_text(size = 28),
        legend.position="top",
        legend.text = element_text(size = 24),
        legend.title = element_text(size = 24),
        strip.text.x = element_text(size = 24),
        panel.border = element_blank(),
        panel.grid.major = element_blank(),
        panel.grid.minor = element_blank(),
        panel.background = element_rect(fill = "white"),
        panel.spacing = unit(0.5, "cm")) +
  labs(x = "Community (Obs vs. Pred)", y = "Yield (t/ha)") +
  scale_y_continuous(breaks = round(seq(min(preds_stack$Yield),
                                       max(preds_stack$Yield*2),
                                       by = 2),1)) +
  scale_fill_manual(values = c("black", "#cb66ff", "#669aff", "#ffcb66"))

plotCompMono #Plot comparing the selected mono plots from the original experiment

ggsave("LegacyNetIE2_CompMono.png", plot = plotCompMono,
       width = 25, height = 12, units = "in")

#####

```

N The code used for the additive partitioning method using the realised proportions case study model

```
#remotes::install_github("BenjaminDelory/bef")

library(bef)

#Extract desired mixture communities from the dataset

Gdata_biomass <- Gdata_annual[, c(2:4, 11:13)]

Gmix_sown <- Gdata_biomass[which(
                                rowSums(Gdata_biomass[, 1:3]!=0)
                                != 1), -1:-3]

colnames(Gmix_sown) <- c("Grass", "Legume", "Herb")

#Extract desired monoculture communities from the dataset

Gmono_sown <- Gdata_biomass[which(
                                rowSums(Gdata_biomass[, 1:3]!=0)
                                == 1), -1:-3]

colnames(Gmono_sown) <- c("Grass", "Legume", "Herb")

#Extract monoculture observations

Gplots_sown <- as.matrix(Gdata_biomass[which(rowSums(Gdata_biomass[, 1:3]!=0) != 1), 1:3])

colnames(Gplots_sown) <- c("Grass", "Legume", "Herb")

#Additive partitioning using traditional methodology

addpart_sown <- apm(mix = Gmix_sown, mono = Gmono_sown,
                   ry = Gplots_sown, method = "loreau")

addpart_sown <- cbind(Gplots_sown, addpart_sown)

#Create desired dataset, may include unsown communities

Gdata_new <- data.frame(
  "Community" = c("G", "L", "H", "G/L Mix", "G/H Mix", "L/H Mix", "G/L/H Mix"),
  "Richness" = c("2", "2", "2", "4", "4", "4", "6"),
  "SownLp" = c(0.5, 0, 0, 0.25, 0.25, 0, 1/6),
  "SownPp" = c(0.5, 0, 0, 0.25, 0.25, 0, 1/6),
  "SownRc" = c(0, 0.5, 0, 0.25, 0, 0.25, 1/6),
  "SownWc" = c(0, 0.5, 0, 0.25, 0, 0.25, 1/6),
```

```

"SownCi" = c(0, 0, 0.5, 0, 0.25, 0.25, 1/6),
"SownPl" = c(0, 0, 0.5, 0, 0.25, 0.25, 1/6)
)

Gdata_new <- rbind(Gdata_new, c("Community" = "Opt (1)", "Richness" = "4", best_preds1[1:6]))
Gdata_new <- rbind(Gdata_new, c("Community" = "Opt (2)", "Richness" = "3", best_preds2[1:6]))

#Remove monocultures and unneeded columns

Gplots_new <- Gdata_new[-1:-3, -1:-2]

#Predict for mixtures

Gmix_new <- predict(GModel, xnew = Gplots_new)[, 7:10]

Gmix_new$Grass <- Gmix_new$DMyield * Gmix_new$VAG
Gmix_new$Legume <- Gmix_new$DMyield * Gmix_new$VAL
Gmix_new$Herb <- Gmix_new$DMyield * Gmix_new$VAH
Gmix_new <- Gmix_new[, -1:-4]

#Summarise species into functional groups

Gplots_new$Grass <- Gplots_new$SownLp + Gplots_new$SownPp
Gplots_new$Legume <- Gplots_new$SownRc + Gplots_new$SownWc
Gplots_new$Herb <- Gplots_new$SownCi + Gplots_new$SownPl
Gplots_new <- as.matrix(Gplots_new[, -1:-6])

#adjust rounding issue, see rowSums(Gplots_new)

Gplots_new[5, 2:3] <- Gplots_new[5, 2:3] - 0.0005

#Predict for monocultures

Gmono_new <- predict(GModel, xnew = Gdata_new[1:3, -1:-2])[, 7:10]

Gmono_new$Grass <- Gmono_new$DMyield * Gmono_new$VAG
Gmono_new$Legume <- Gmono_new$DMyield * Gmono_new$VAL
Gmono_new$Herb <- Gmono_new$DMyield * Gmono_new$VAH
Gmono_new <- Gmono_new[, -1:-4]

#Additive partitioning using DI model predictions

addpart_new <- apm(mix = Gmix_new, mono = Gmono_new, ry = Gplots_new, method = "loreau")

addpart_new <- cbind(Gplots_new, addpart_new)

addpart_new <- cbind(Community = Gdata_new$Community[-1:-3], Richness = Gdata_new$Richness[-1:-3],
addpart_new)

#Prepare data for plotting

```

```

addpart_new_stack <- reshape2::melt(as.data.frame(addpart_new),
                                   id.vars = c("Community", "Richness", "Grass", "Legume", "Herb", "NBE"), variable.name
= "Effect", value.name = "Value")

addpart_new_stack$NBE <- as.numeric(addpart_new_stack$NBE)

addpart_new_stack$Value <- as.numeric(addpart_new_stack$Value)

addpart_new_stack$Community <- factor(addpart_new_stack$Community,
                                     levels = c("G/L Mix", "G/H Mix", "L/H Mix", "G/L/H Mix", "Opt (2)", "Opt (1)"))

#Plot results

addpart_plot <- ggplot(addpart_new_stack, aes(fill = Effect, y = Value, x = Community)) +
  geom_bar(position = "stack", stat = "identity") +
  facet_grid(~Richness, scales = "free_x", space = "free_x") +
  theme(axis.text = element_text(size = 24, angle = 20),
        axis.title = element_text(size = 28),
        legend.position="top",
        legend.text = element_text(size = 24),
        legend.title = element_text(size = 24),
        strip.text.x = element_text(size = 24),
        panel.border = element_blank(),
        panel.grid.major = element_blank(),
        panel.grid.minor = element_blank(),
        panel.background = element_rect(fill = "white"),
        panel.spacing = unit(0.5, "cm")) +
  scale_fill_manual(values = c("#cb66ff", "#ffcb66"))

addpart_plot

ggsave("LegacyNetIE2_AddPart.png", plot = addpart_plot,
       width = 12, height = 8, units = "in")

#####

```