

Get Started with PRIMER 8

A quick guide to get you up and running with PRIMER 8 software by Marti J. Anderson (2026)

- [Citation](#)
- [Table of contents](#)
- [Summary](#)
- [1. Install & activate PRIMER 8](#)
 - [Overview](#)
- [2. Getting data in to PRIMER](#)
 - [Getting example data](#)
 - [Opening example data](#)
 - [Importing data from Excel](#)
 - [Post-import data checks](#)
 - [Save your data & workspace](#)
- [3. Example analysis pathway \(CLUSTER, MDS, ANOSIM\)](#)
 - [An analysis of biotic data](#)
 - [Step 1: Transformation](#)
 - [Step 2: Resemblance](#)
 - [Step 3: Cluster](#)
 - [Step 4: Ordination](#)
 - [Step 5: ANOSIM test](#)
 - [Summary of the pathway](#)
- [4. Some wizards](#)
 - [Basic multivariate analysis](#)

- [Matrix display](#)
- [5. Run a PERMANOVA](#)
 - [Overview](#)
 - [A three-factor hierarchical design](#)
 - [Steps in a PERMANOVA analysis](#)
 - [Step 1: Data selection](#)
 - [Step 2: Jaccard resemblance](#)
 - [Step 3: Specify the design](#)
 - [Step 4: Run PERMANOVA](#)
 - [Step 4 \(continued\): Key additional details about PERMANOVA in PRIMER](#)
 - [Step 5: Ordination of centroids](#)
 - [Summary of the PERMANOVA analysis](#)
- [6. Help & additional resources](#)
 - [Help & support](#)
 - [Further reading](#)
- [References](#)

Citation

Anderson, M.J. (2026). "*Get Started with PRIMER 8.*" *PRIMER-e Learning Hub*. PRIMER-e, Auckland, New Zealand. <https://learninghub.primer-e.com/books/get-started-with-primer-8>.

Table of contents

- **Citation**
- **Summary**
- **1. Install & activate PRIMER 8**
 - Overview
- **2. Getting data in to PRIMER**
 - Getting example data
 - Opening example data
 - Importing data from Excel
 - Post-import data checks
 - Save your data & workspace
- **3. Example analysis pathway (CLUSTER, MDS, ANOSIM)**
 - An analysis of biotic data
 - Step 1: Transformation
 - Step 2: Resemblance
 - Step 3: Cluster
 - Step 4: Ordination
 - Step 5: ANOSIM test
 - Summary of the pathway
- **4. Some wizards**
 - Basic multivariate analysis
 - Matrix display
- **5. Run a PERMANOVA**
 - Overview
 - A three-factor hierarchical design
 - Steps in a PERMANOVA analysis
 - Step 1: Data selection
 - Step 2: Jaccard resemblance
 - Step 3: Specify the design
 - Step 4: Run PERMANOVA
 - Step 4 (continued): Key additional details about PERMANOVA in PRIMER
 - Step 5: Ordination of centroids

- Summary of the PERMANOVA analysis
- **6. Help & additional resources**
 - Help & support
 - Further reading
- **References**

Summary

This book provides a guide to get you up and running with PRIMER version 8. For a more comprehensive look at the many new features in PRIMER 8 that distinguish it from PRIMER 7, see the following resource:

- Anderson, M.J. (2026) "**What's New in PRIMER 8**". *PRIMER-e Learning Hub*. PRIMER-e, Auckland, New Zealand. <https://learninghub.primer-e.com/books/whats-new-in-primer-8>.

If you are looking for help on getting started with PRIMER version 7, see the following resource:

- Anderson, M.J. (2024) "**A Quick Guide to PRIMER 7**". *PRIMER-e Learning Hub*. PRIMER-e, Auckland, New Zealand. <https://learninghub.primer-e.com/books/a-quick-guide-to-primer-7>.

1. Install & activate PRIMER 8

1. Install & activate PRIMER 8

Overview

To get started using PRIMER 8, you will need to follow two essential steps:

- first, *install* the software on your machine;
- second, *activate* the software, using your purchased licence key.[¶]

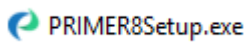
These two steps are outlined in detail below.

Install PRIMER 8

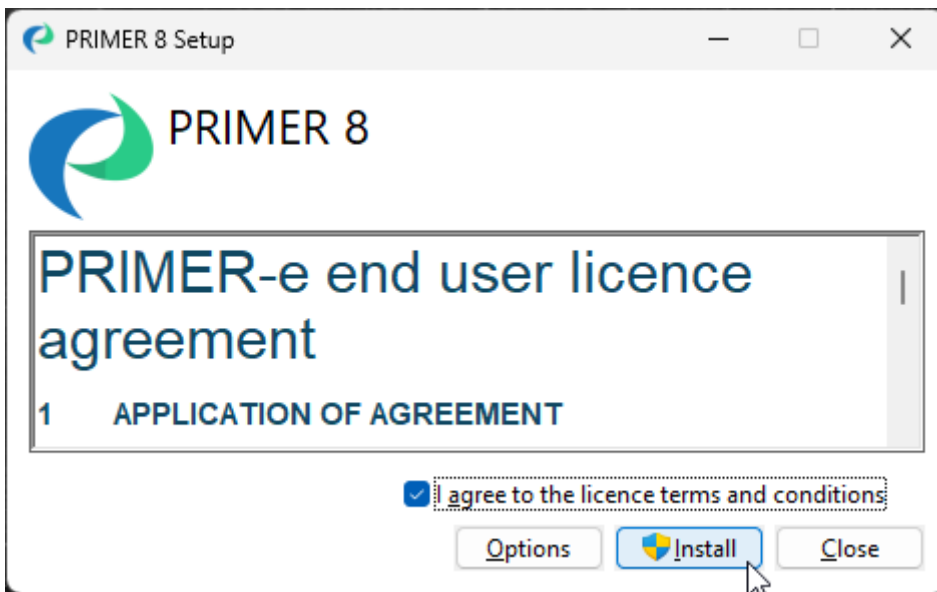
When you purchase a subscription, a direct link to download the PRIMER 8 software installer (an executable file called 'PRIMER8Setup.exe') will be sent to you *via* email, along with your licence key. Alternatively, you can go to our [download page](#) and download it from there. Note that the same installer file is used for all PRIMER 8 subscription plans.

Important note: You will need to have administrator permissions on your user account in Windows to install PRIMER 8 software on your computer. If you do not have these permissions, please contact your IT team who should be able to assist you with the installation.

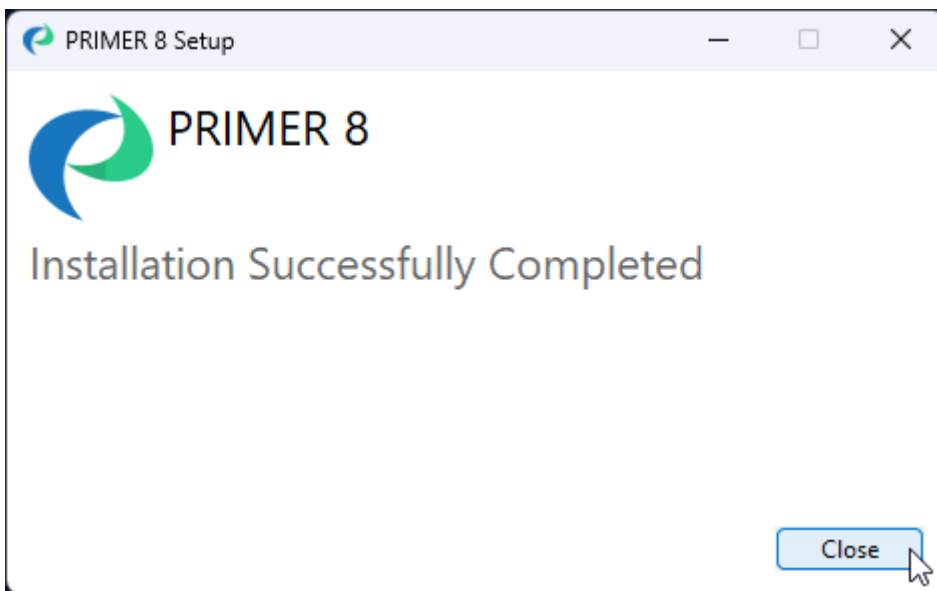
1. Double-click on the executable file to start the installation process.



2. A 'PRIMER 8 Setup' window will appear. Please read the 'PRIMER-e end user licence agreement' ([EULA](#)). If you agree to the terms and conditions of this EULA, click on the tick-box ☒ 'I agree to the licence terms and conditions' and then click '**Install**'.



3. Follow the prompts and once the installation has completed successfully, click '**Close**'.

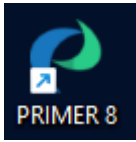


Once the software is installed, you can launch the PRIMER 8 software and it will run in [trial mode](#) for free for 14 days. You will need to activate your licence key in order to enable full functionality of the software in accordance with your purchased subscription plan.

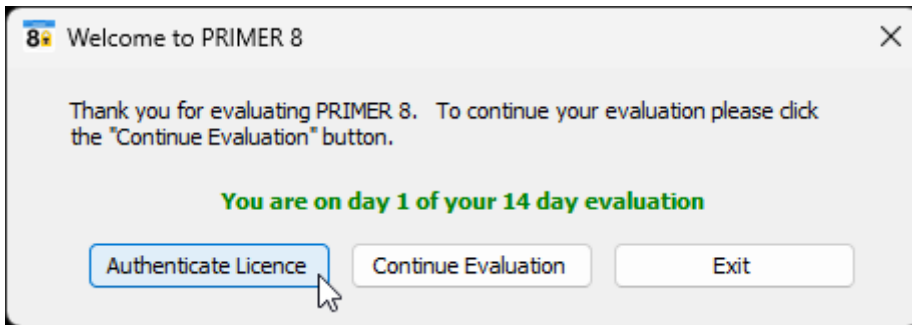
Activate PRIMER 8

Note: The machine on which you are activating PRIMER must be connected to the internet and able to talk to our authentication server.

1. Launch PRIMER 8 (e.g., by double-clicking on the PRIMER 8 icon on your desktop):

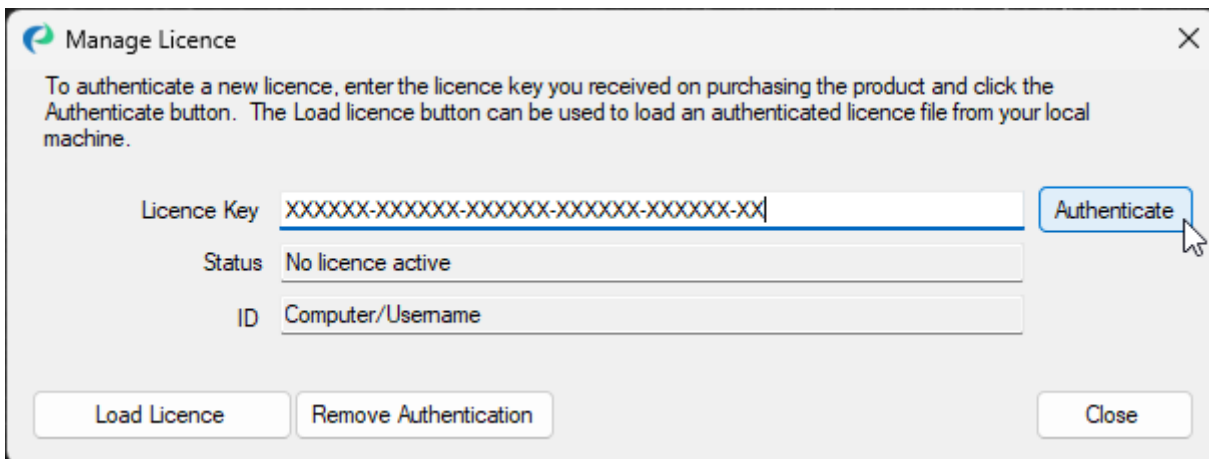


2. Click on the '**Authenticate Licence**' button in the 'Welcome to PRIMER 8' dialog box.



(Note: If you have already exhausted your 14-day trial, a message may appear to advise you that 'Your evaluation period has expired'. If that happens, simply click on 'Authenticate Licence' to carry on with the instructions to activate your licence.)

3. In the 'Manage Licence' dialog window, copy and paste your 37-character licence key code into the 'Licence Key' box and click '**Authenticate**'.



PRIMER 8 will now operate with all of the entitlements and functionality associated with your subscription plan.

For more details regarding installing, activating, and trouble-shooting any issues that could arise here, see [Managing your PRIMER 8 installation](#).

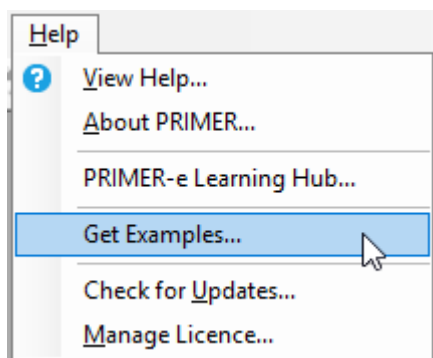
[¶]You may explore the features of PRIMER 8 [in trial mode](#) for free for a period of 14 days without purchasing and activating the software using a licence key, but not all features will be fully functional.

2. Getting data in to PRIMER

2. Getting data in to PRIMER

Getting example data

You can get a full set of **Example datasets** from within the software. To access them click **Help > Get Examples...** and choose the folder to place them in. The default name for this folder is 'Examples_P8'. These datasets are used and referred to in the software documentation (including this book!).



We highly recommend working through the manuals available on our [Learning Hub](#), and replicating the examples given there with the **Example datasets**. This is an excellent way to familiarize yourself with the software and its many routines.

Opening example data

Data from the Fal estuary

Example datasets can be obtained *via* the **Help** menu item. Click **Help > Get Examples...**, and you can download the folder of examples called 'Examples_P8' to a location of your choice. Inside this folder, there are individual folders that hold a plethora of datasets. Here we will focus on data from a study of benthic infaunal communities in soft sediments from 27 sites over five creeks of the Fal estuary, SW England ([Somerfield *et al.* \(1994a\)](#) , [Somerfield *et al.* \(1994b\)](#)). Sediments at these sites were contaminated to varying degrees by heavy metals, from historic mining activities. Both faunal counts and environmental measures were obtained from the same set of sites. The data files from this study are located in a folder called 'Fal_benthic_fauna' inside 'Examples_P8'.

Here, we shall focus on data for the *nematodes*, located in the file named 'Fal_nematode_abundance.pri', which contains abundances of 62 nematode species x 27 sites. Note that the extension *.pri indicates a file containing a data matrix that has been saved in PRIMER's own internal binary format. (These data are also available in Excel format in the file called 'Fal_nematode_abundance.xlsx'.)

Open the Fal data in PRIMER

To open the species-by-samples data matrix, launch PRIMER, then click **File > Open** from the main menu, navigate to the 'Fal_benthic_fauna' folder inside the 'Examples_P8' folder (in the location you had specified for it), and select 'Fal_nematode_abundance.pri'. Click **Open** to display the species matrix.

Properties of the data

Click on **Edit > Properties** and you will see that PRIMER-format *.pri data sheets carry other information about the data matrix as well, including:

- Title,
- Data type,
- Array size (number of columns and rows),
- Orientation (whether samples are found in columns or rows),
- a Description, and
- a History of what pre-treatments (such as a transformation) have been applied to them.

Sample Data Properties

Title:
Fal estuary nematodes

Data type:
☒ Abundance
☐ Biomass
☐ Environmental
☐ Unknown/other

History:

Samples as:
☒ Columns
☐ Rows

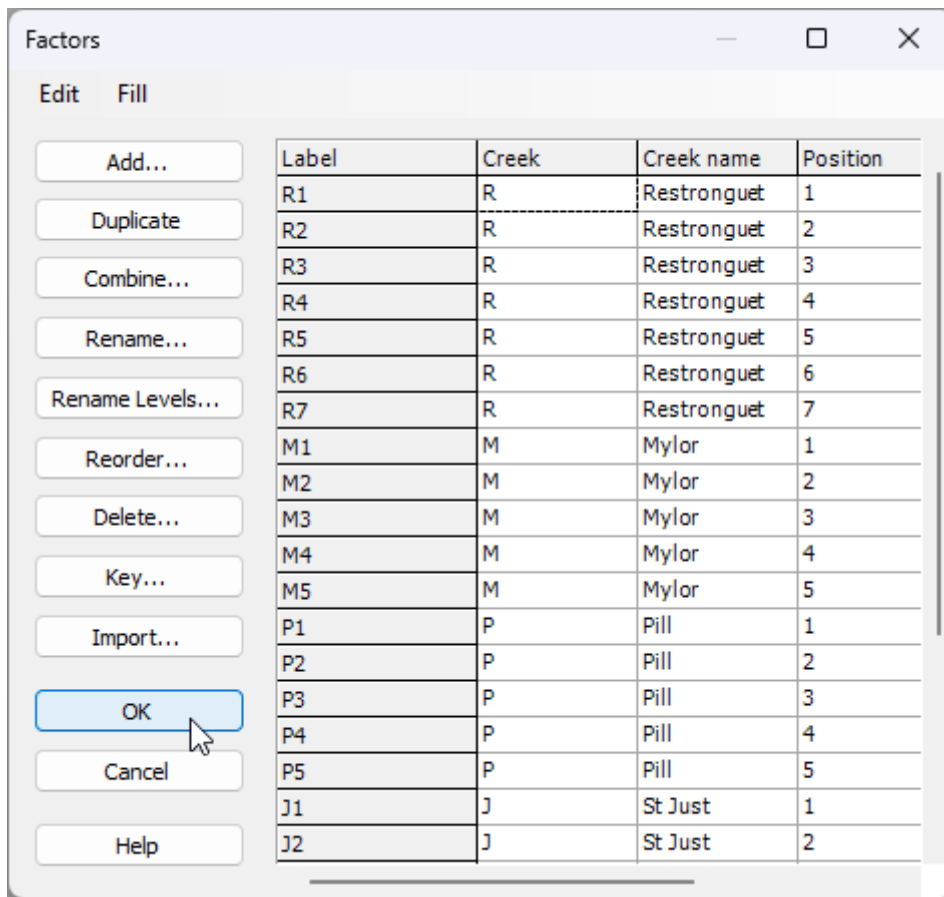
Number of columns: 27
Number of rows: 62

Description:
Description:
Abundance of 64 nematode species from meiofauna cores of the sediments of 27 sites over 5 creeks of the Fal estuary, Cornwall, UK: Restronguet (R), Mylor (M), Pill (P), St Just (J), Percuil (E), collected in November 1991. These sites are subject to differing levels of heavy metal concentrations in sediments from historic tin mining.

OK Cancel Help

Factors associated with the data

Click on **Edit > Factors...**, and you can see that a subsidiary sheet of three factors is also linked to this worksheet: 'Creek', a single-letter abbreviation for the creeks, the full 'Creek name' and a numeric 'Position' factor identifying the location of each sampling site down each creek.



Note that additional factors could be added here by clicking **Add**, and also combinations of levels of existing factors can be created by clicking on **Combine**.

2. Getting data in to PRIMER

Importing data from Excel

Step 1. Ensure your data are in a format suitable for import into PRIMER

Suppose we have a dataset in Excel that is already in a suitable format for import into PRIMER. The environmental data from the Fal estuary provides an example of this. These data are found in the file 'Fal_environment.xls' and consist of values for each of 12 environmental variables measured from sediments collected from 27 sites across 5 tidal creeks in the Fal Estuary (available in the folder called 'Fal_benthic_fauna' inside 'Examples_P8', obtained by clicking **Help > Get Examples...**, as seen in the last section).

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	
1	Fal estuary environmental variables														
2		R1	R2	R3	R4	R5	R6	R7	M1	M2	M3	M4	M5	P1	P2
3	% silt/clay	87	76	71	73	45	70	61	97	92	93	97	97	94	
4	% organic c	6.8	6.5	5.7	5.7	3.9	8.7	6.5	8.5	9	8.5	8.5	8.1	8.8	
5	Ag	4.94	4.91	3.27	3.04	3.09	3.43	2.73	2.34	2.2	2.48	2.23	2.38	1.49	
6	Cd	3.75	4.45	2.49	2.61	1.44	2.59	1.91	1.57	1.48	1.46	1.19	1.63	1.74	
7	Co	27.4	27.2	21.5	18.7	17.3	22.2	18.8	13.3	12.8	12.2	12.8	13.3	11.4	
8	Cr	45.4	43.8	37.9	32.4	34	41.3	36.4	66.3	56.5	63	54.4	58.2	41.4	
9	Cu	3302	3373	2387	2176	2040	2452	1997	1354	1246	1274	1188	1300	690	
10	Fe	67862	68098	53118	49505	47250	55256	49826	42261	41326	41584	38693	39978	32023	
11	Mn	611	586	521	521	507	539	488	405	411	392	373	394	279	
12	Ni	37.3	35.9	28.1	26.5	27.6	29.2	25.3	33.8	31.7	32.6	31.9	33.8	25.9	
13	Pb	271	265	192	180	165	206	185	196	179	186	186	194	134	
14	Zn	5279	5873	3583	3328	2252	3475	2900	1460	1357	1481	1379	1479	1085	
15															
16	Creek	R	R	R	R	R	R	R	M	M	M	M	M	P	P
17	Creek name	Restrongu	Restrongu	Restrongu	Restrongu	Restrongu	Restrongu	Restrongu	Mylor	Mylor	Mylor	Mylor	Mylor	Pill	Pi
18	Position	1	2	3	4	5	6	7	1	2	3	4	5	1	
19															
20															
21															

Important things to note about this file:

- There is a *title* for the dataset ('Fal estuary environmental variables') in the very first (upper left-hand) cell (A1). This title is optional, but handy as a naming convention.
- The cell immediately under the title (cell A2) is *empty*.
- There are *column labels* ('R1', 'R2', ...) in row 2. These are unique labels for the sampling units (Sites in this case).

- There are *row labels* ('%silt/clay', '% organic carbon', ...) in column A. These are unique names associated with each variable.
- The entries for every cell in the matrix of data itself (beginning with cell B3) all contain *numerical values only*. There are no non-numeric characters. This means that you may not use 'NaN' or 'NAN' to denote missing values. If data are missing from a cell, then it should be left **blank**. In addition, symbols such as '<' or '>' (for 'less than' or 'greater than') are similarly not permitted or accepted as valid data values within the data matrix.
- In this example, the variables are rows and the sampling units (sites) are columns. It is perfectly ok to have this formatted the other way around, with variables as columns and sampling units as rows. You will specify the orientation of your data matrix explicitly when you import your Excel file into PRIMER.

This format must be adhered to precisely, with no extra blank rows or columns, or extra headers, otherwise PRIMER will not be able to open it successfully.

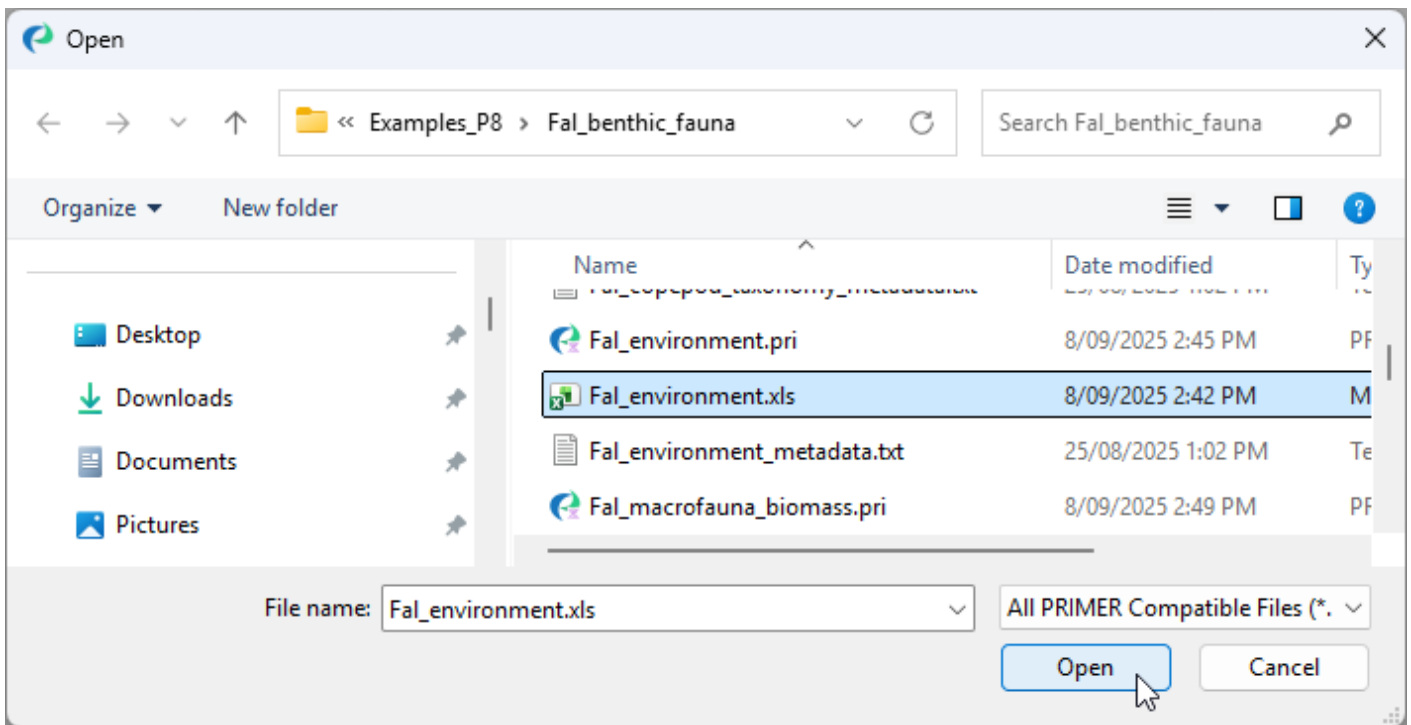
Other things to note:

- **Factors:** You can label the sampling units as belonging to a level of one (or more) factors by **skipping a line** at the bottom of the matrix and placing this 'factor information' there. In the above image you can see that there are three factors: 'Creek' (row 16), 'Creek name' (row 17) and 'Position' (row 18).
- **Indicators:** You can similarly label variables as belonging to particular groups in the same manner; this is done along the *other* margin of the data matrix (e.g., after skipping a column, for this example). This might be useful for doing analyses on subsets of variables belonging to different types, such as physical vs chemical variables. In a case where variables are species, one might want to consider subsets of variables corresponding to families, functional groups, etc.

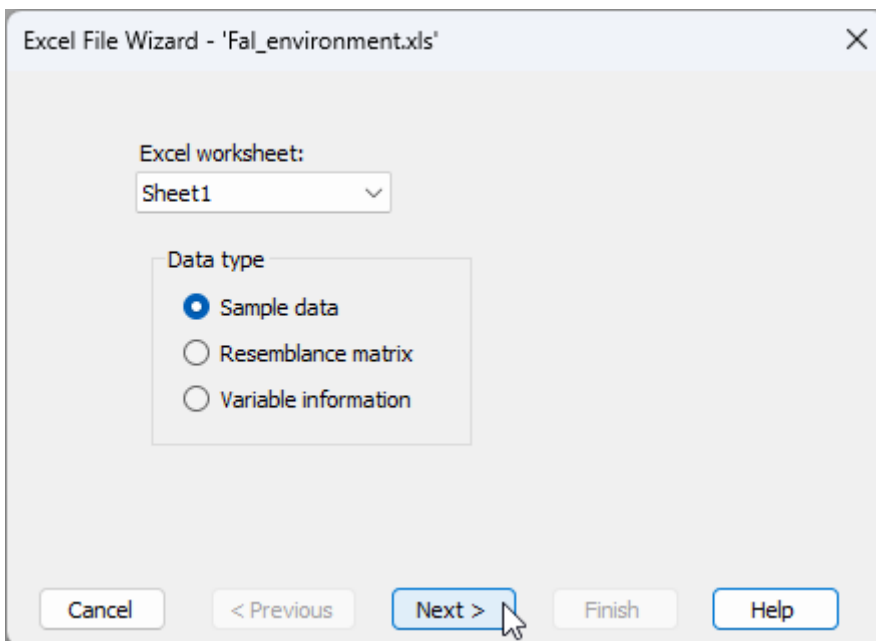
Inclusion of one (or more) *factors* (to specify groups of samples) or *indicators* (to specify groups of variables) is optional. If you have more than one factor, then these are given one after the next (in adjacent rows); do not put blank rows between multiple factors. Similarly, if you have more than one indicator, these would be given one after the next (in adjacent columns, for this orientation). The initial single blank row (or column) is there simply to demarcate the difference between the data matrix itself and additional information about the data matrix upon import.

Step 2. Open PRIMER and import the data from Excel

1. Once your Excel file is ready, open up PRIMER and choose **File > Open**, then click on the name of your Excel file in the browser (here it is 'Fal environment.xls'), then click **Open**.



- This will initiate PRIMER's Excel File data-import Wizard. Choose the name of the specific sheet within your Excel file that contains your data and the type of data you are importing. Here, we have (Excel worksheet: **Sheet1**) & (Data type **Sample data**), then click **Next >**.



- Choose the correct orientation, type of data and the meaning of blank entries (if any). For this example, we have (Orientation **Samples are columns**) & (Data type **Environmental**) & (Blank = **Missing value**), then click **Finish**.

Excel File Wizard - 'Fal_environment.xls'

☒ Title
☒ Row labels

Orientation

☒ Samples as columns
☐ Samples as rows

Data type

☐ Abundance
☐ Biomass
☒ Environmental
☐ Unknown/other

Blank =

☒ Missing value
☐ Zero

Cancel < Previous Next > Finish Help

4. You will now see your data file has been imported and is nicely displayed in the PRIMER workspace. It appears in its own window, and its name also appears in the 'Explorer tree'-type window shown on the left-hand side of the PRIMER desktop.

PRIMER

File Edit Select View Wizards Pre-treatment Analyse Plots PERMANOVA+ Tools Window Help

Workspace

Fal_environment

Fal_environment

Fal estuary environmental variables
Environmental

	Samples							
	R1	R2	R3	R4	R5	R6	R7	M
% silt/clay	87	76	71	73	45	70	61	
% organic carbon	6.8	6.5	5.7	5.7	3.9	8.7	6.5	
Ag	4.94	4.91	3.27	3.04	3.09	3.43	2.73	
Cd	3.75	4.45	2.49	2.61	1.44	2.59	1.91	
Co	27.4	27.2	21.5	18.7	17.3	22.2	18.8	
Cr	45.4	43.8	37.9	32.4	34	41.3	36.4	
Cu	3302	3373	2387	2176	2040	2452	1997	
Fe	67862	68098	53118	49505	47250	55256	49826	
Mn	611	586	521	521	507	539	488	
Ni	37.3	35.9	28.1	26.5	27.6	29.2	25.3	
Pb	271	265	192	180	165	206	185	
Zn	5279	5873	3583	3328	2252	3475	2900	

Now that the data are in PRIMER, it is a good idea to do some post-import data checks.

Post-import data checks

Check the orientation

After import, make sure you have specified the orientation correctly by examining the labels on the columns and rows of the data frame. For example, after importing the Fal environmental data from Excel (see the previous page), you can see that the columns are 'Samples' (a navy blue-coloured strip across the top) and rows are 'Variables' (a teal-coloured strip along the left margin).

	Samples							
	R1	R2	R3	R4	R5	R6	R7	M
% silt/clay	87	76	71	73	45	70	61	
% organic carbon	6.8	6.5	5.7	5.7	3.9	8.7	6.5	
Ag	4.94	4.91	3.27	3.04	3.09	3.43	2.73	
Cd	3.75	4.45	2.49	2.61	1.44	2.59	1.91	
Co	27.4	27.2	21.5	18.7	17.3	22.2	18.8	
Cr	45.4	43.8	37.9	32.4	34	41.3	36.4	
Cu	3302	3373	2387	2176	2040	2452	1997	
Fe	67862	68098	53118	49505	47250	55256	49826	
Mn	611	586	521	521	507	539	488	
Ni	37.3	35.9	28.1	26.5	27.6	29.2	25.3	
Pb	271	265	192	180	165	206	185	
Zn	5279	5873	3583	3328	2252	3475	2900	

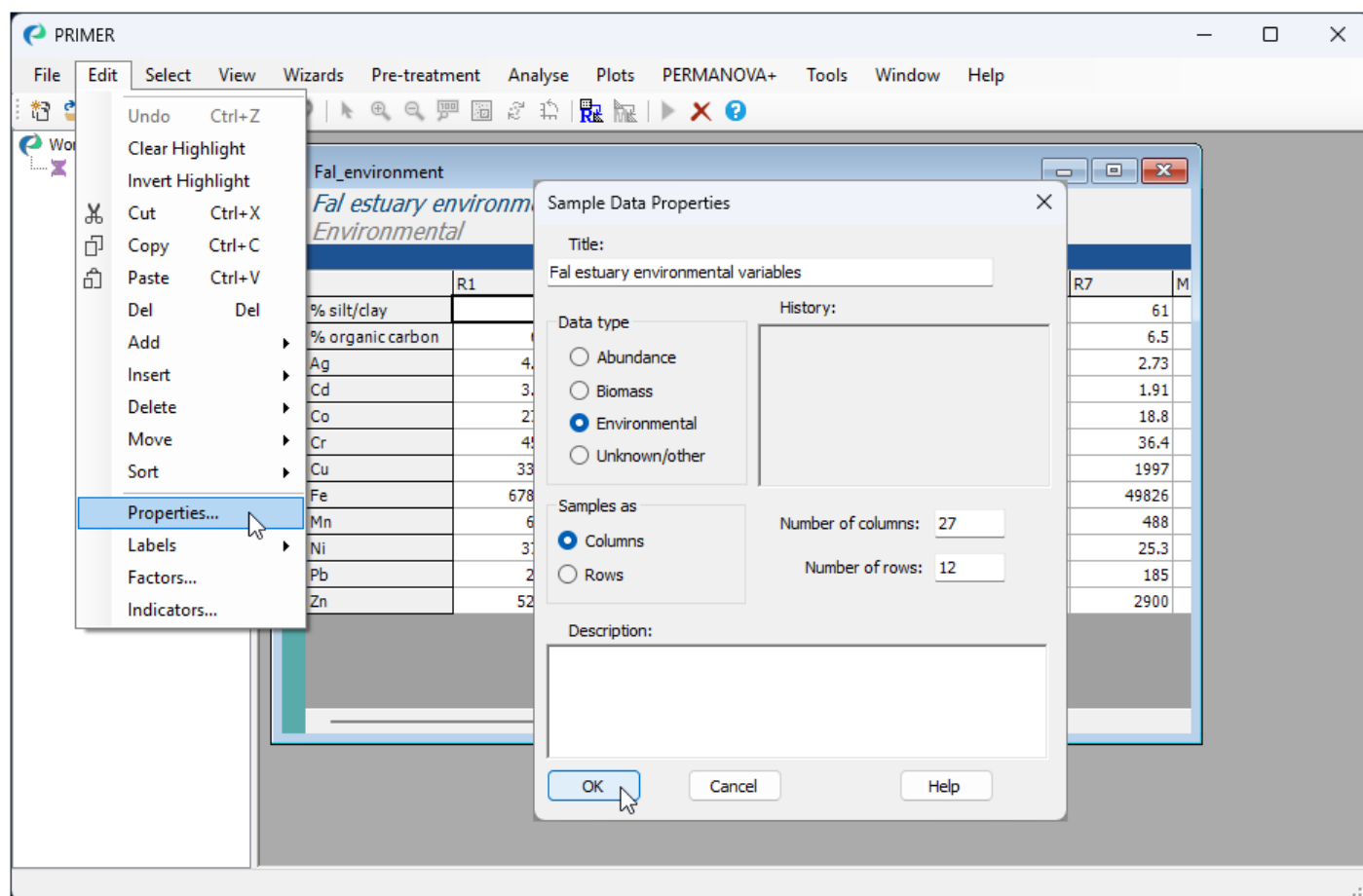
If you happen to get this the wrong way around (e.g., if your variables are actually columns instead of rows), this can easily be changed (swapped around) by choosing **Edit > Properties** and toggling the radio button for 'Samples as' to either '\$\bullet\$Columns' or '\$\bullet\$Rows', whichever is appropriate.

Check the properties, factors and indicators

To be sure that the import has been fully successful, including all data points, factors and indicators that may have been included in your original Excel file, you can see additional

information attached to your data matrix by clicking on your imported dataset in PRIMER, and doing the following:

- Look at the data properties, size of the matrix, etc.: Click **Edit > Properties....**



For the Fal environmental dataset, we can see there are 12 variables and 27 sites, etc. Note that you can add a useful 'Description' of your data into this dialog if you like.

- Look at the Factors (if any): Click **Edit > Factors....** For the Fal environmental dataset, you will see the same three factors of 'Creek', 'Creek name' and 'Position' that we saw in the Excel file.

Factors

Edit Fill

Add... Duplicate Combine... Rename... Rename Levels... Reorder... Delete... Key... Import... OK Cancel Help

Label	Creek	Creek name	Position
R1	R	Restronguet	1
R2	R	Restronguet	2
R3	R	Restronguet	3
R4	R	Restronguet	4
R5	R	Restronguet	5
R6	R	Restronguet	6
R7	R	Restronguet	7
M1	M	Mylor	1
M2	M	Mylor	2
M3	M	Mylor	3
M4	M	Mylor	4
M5	M	Mylor	5
P1	P	Pill	1
P2	P	Pill	2
P3	P	Pill	3
P4	P	Pill	4
P5	P	Pill	5
J1	J	St Just	1
J2	J	St Just	2

- Look at the Indicators (if any): Click **Edit** > **Indicators....** For the Fal environmental dataset, there were no indicators, but you could add some here, if you wish.

Indicators

Edit Fill

Add... Duplicate Combine... Rename... Rename Levels... Reorder... Delete... Key... Import... OK Cancel Help

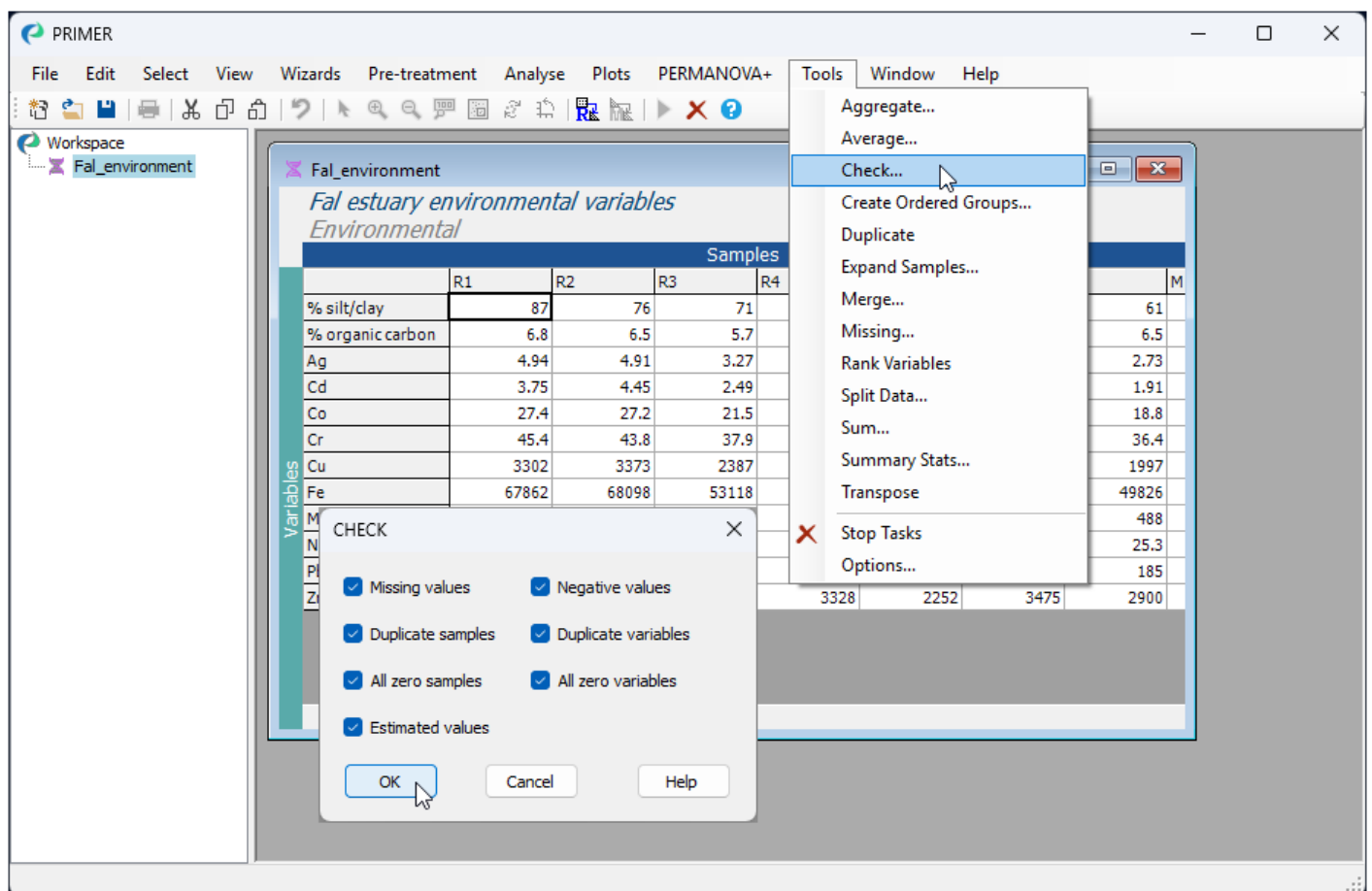
No indicators click 'Add...' button

Run PRIMER's data-checking tool

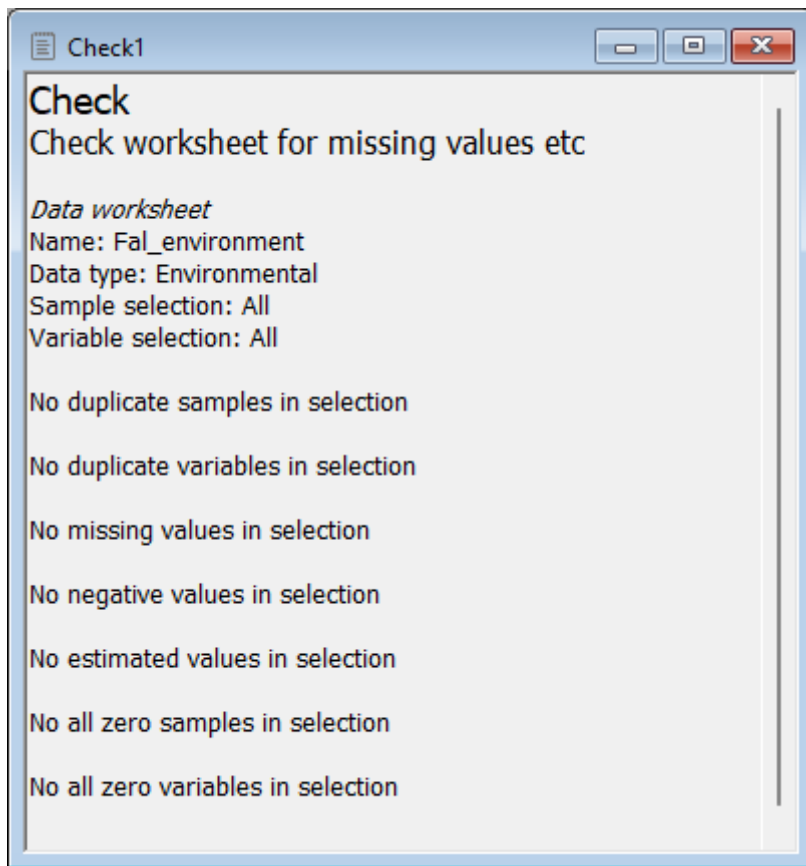
As an additional option, you can run PRIMER's internal data-checking tool to find and identify certain other features that might be present in your data, including:

- Missing values,
- Negative values,
- Duplicate samples,
- Duplicate variables,
- All-zero samples,
- All-zero variables, and/or
- Estimated values

To run this routine, start by clicking on your datasheet, then click on **Tools > Check...**:



For the Fal environmental data, none of these features occurred (see below), and we are ready to proceed with subsequent analyses.

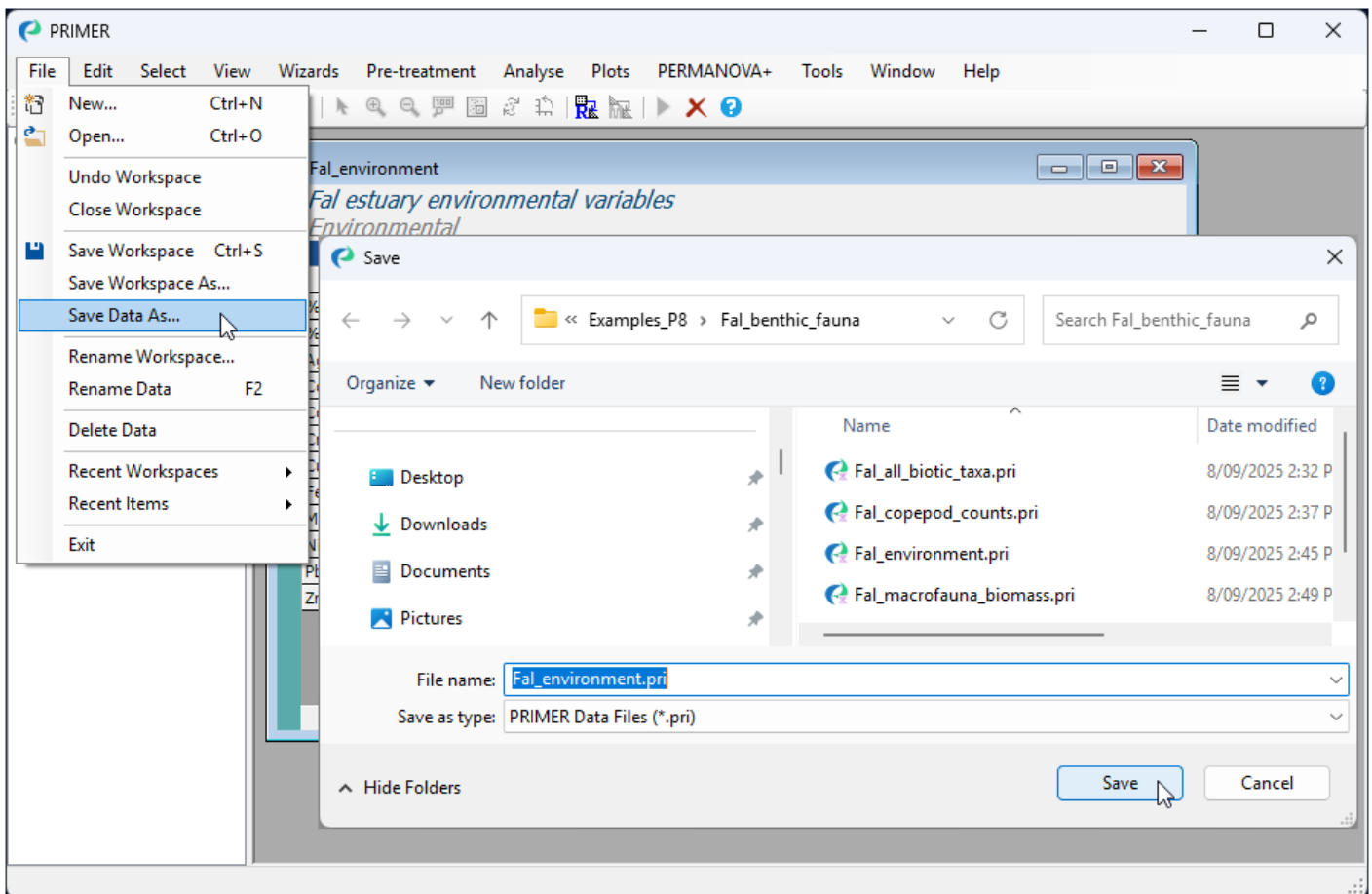


2. Getting data in to PRIMER

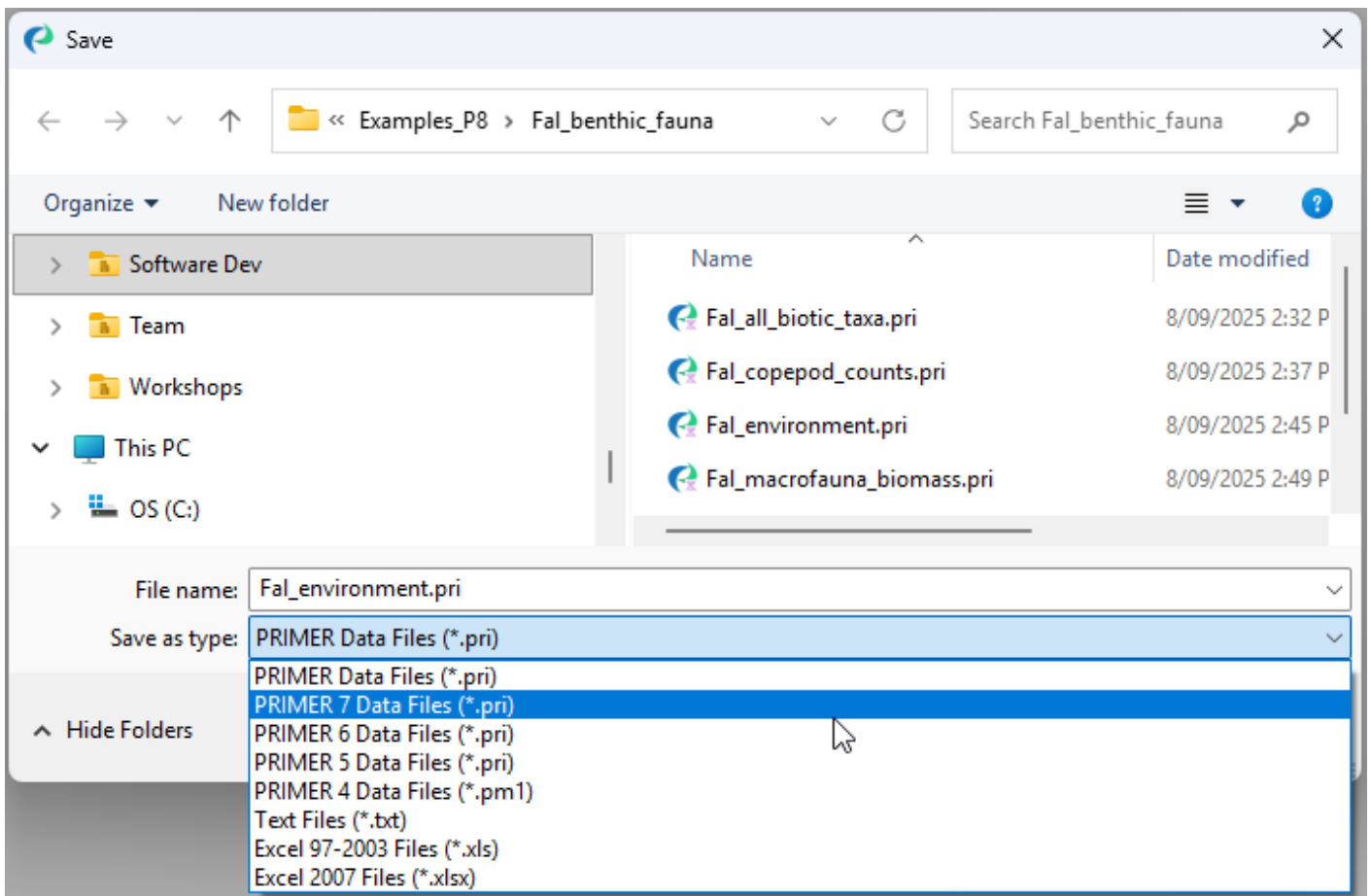
Save your data & workspace

Save your data in PRIMER (*.pri) format

To save a data sheet in PRIMER (*.pri) format, click on the data sheet in the Explorer tree inside the PRIMER workspace you want to save and click **File > Save Data As....**

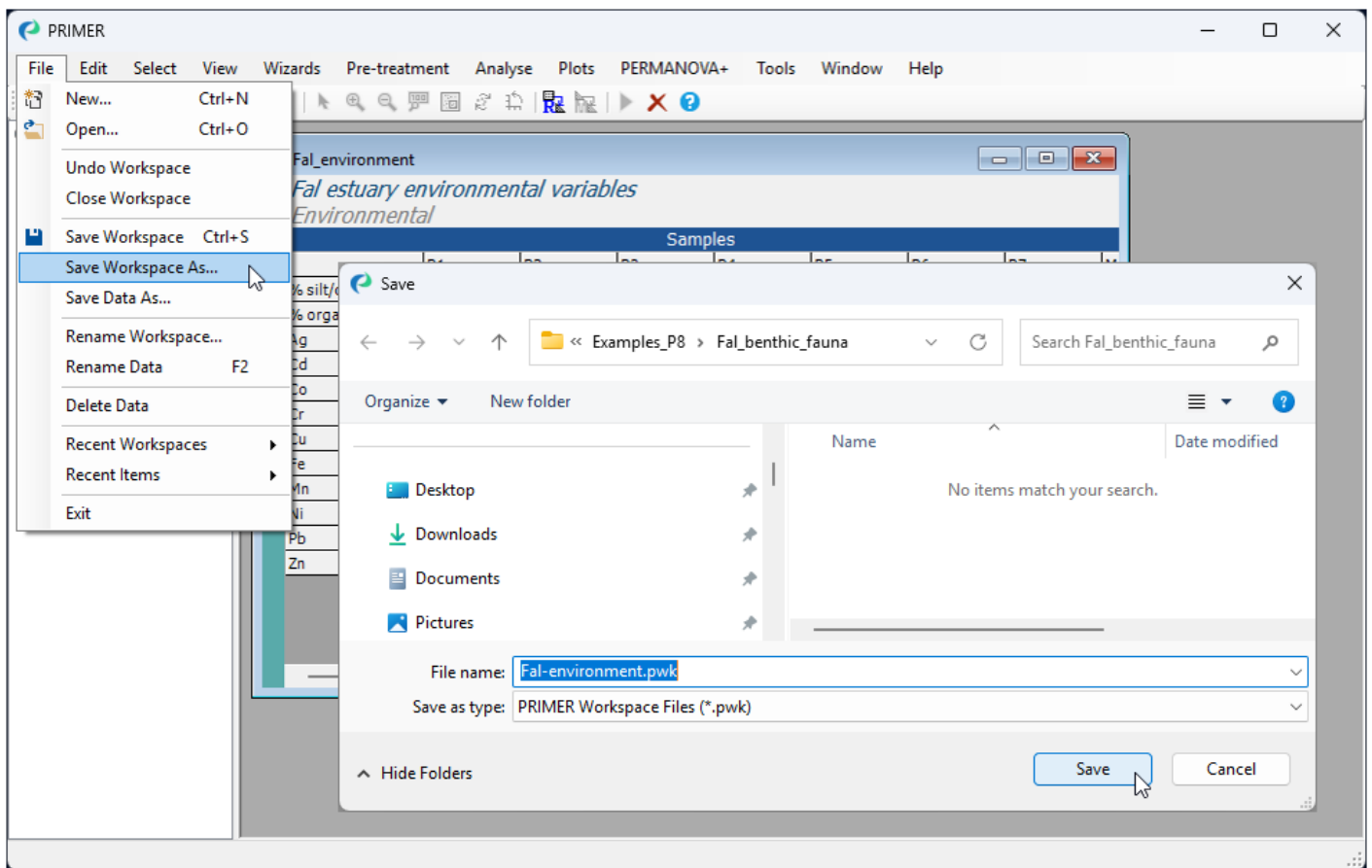


Note the default file-type for saving a data file is 'PRIMER Data Files (*.pri)', as shown in the box labeled 'Save as type:' in the 'Save' dialog. This will save the file in *PRIMER 8 data format*. You can change this (by clicking in the 'Save as type:' box) in order to save the data in some other format of choice, such as *.txt, *.csv, *.xls, *.xlsx, or in a PRIMER data file format compatible with older versions of the PRIMER software (e.g., PRIMER 7 data format):



Save your entire PRIMER workspace (*.pwk)

Generally, you will want to save your data *plus* all of the work you have done in PRIMER 8 to analyse that dataset. To save the entire PRIMER 8 workspace, click **File > Save Workspace As...**



Type the name you wish to give the workspace file (e.g., **Fal_environment.pwk**), and click **Save**. This will save the *entire* PRIMER workspace (not just the data sheet), including all elements you may have created during the PRIMER session (e.g., data, graphics, analyses, etc.), all of which you are able to navigate and see listed (in hierarchical fashion) within the Explorer tree window on the left-hand side of the open PRIMER workspace.

A PRIMER workspace file is identifiable by using the file extension ***.pwk** in the file name. Double-clicking on a file with this extension will open up that entire PRIMER workspace within the PRIMER 8 software.

The default is to save the workspace file in *PRIMER 8 workspace format*. You can change this (by clicking in the 'Save as type:' box) in order to save the workspace in a format compatible with older versions of the PRIMER software, i.e., either PRIMER 7 or PRIMER 6 (versions earlier than PRIMER 6 did not have a workspace). Be aware, however, that some of the graphical elements that are new to PRIMER 8 (e.g., Violin plots, Dot plots, etc.) may not appear when the workspace is saved to older file types.

3. Example analysis pathway (CLUSTER, MDS, ANOSIM)

An analysis of biotic data

Rationale

Multivariate data are very complex and can be difficult to analyse and interpret. Each variable can be considered its own **dimension**, with its own story to tell. Thus, when there are a lot of variables (e.g., if we have sampled occurrence, abundance, or biomass of individual species in a community, and each species is a variable), then we have a very **high-dimensional** system to consider. We can choose to analyse each variable individually, but this would be time-consuming if we have a lot of variables. Also, for biotic data, many species are rare, so our data will be riddled with a large number of zeros - it would be difficult to characterise anything about individual rare species on their own. Furthermore, what we are usually aiming to do is to characterise and analyse the **whole community** *simultaneously* and how it may be changing through time or space, or in response to some experimental treatments (such as human impacts or management scenarios, etc.).

When analysing multivariate biotic data (especially those having a large number of variables), a useful approach is to base the analysis on a resemblance measure (i.e., a dissimilarity or similarity measure), calculated among the sampling units (e.g., [Clarke \(1993\)](#) , [Anderson \(2001a\)](#) , [McArdle & Anderson \(2001\)](#)). A resemblance measure such as Bray-Curtis captures the ecological idea of **turnover** in the identities and relative abundances of species between any pair of sampling units ([Clarke et al. \(2006\)](#) , [Anderson et al. \(2011\)](#)). For many routines in PRIMER, the fundamental starting point is a **resemblance matrix**, which expresses (numerically) the (dis)similarities between every pair of samples.

Before calculating a resemblance matrix, some **pre-treatment** of the data (sometimes in more than one way) is usually desirable. For assemblage data, an **overall transformation** can be used to reduce the dominant contribution of abundant species towards Bray-Curtis similarities. For example, a fourth-root transformation will down-weight the influence of highly-abundant taxa.

From a resemblance matrix, we may readily apply the following methods (Fig. 5.1):

- **Cluster analysis**, to help characterise inter-sample relationships, showing potential groups of similar samples;
- **Ordination**, to help visualise inter-sample relationships in a small number of dimensions (usually 2 or 3); and
- **Hypothesis-testing**, to formally test a particular null hypothesis of interest, such as H_0 : 'There are no differences in community structure among different (*a priori*) groups of samples'.

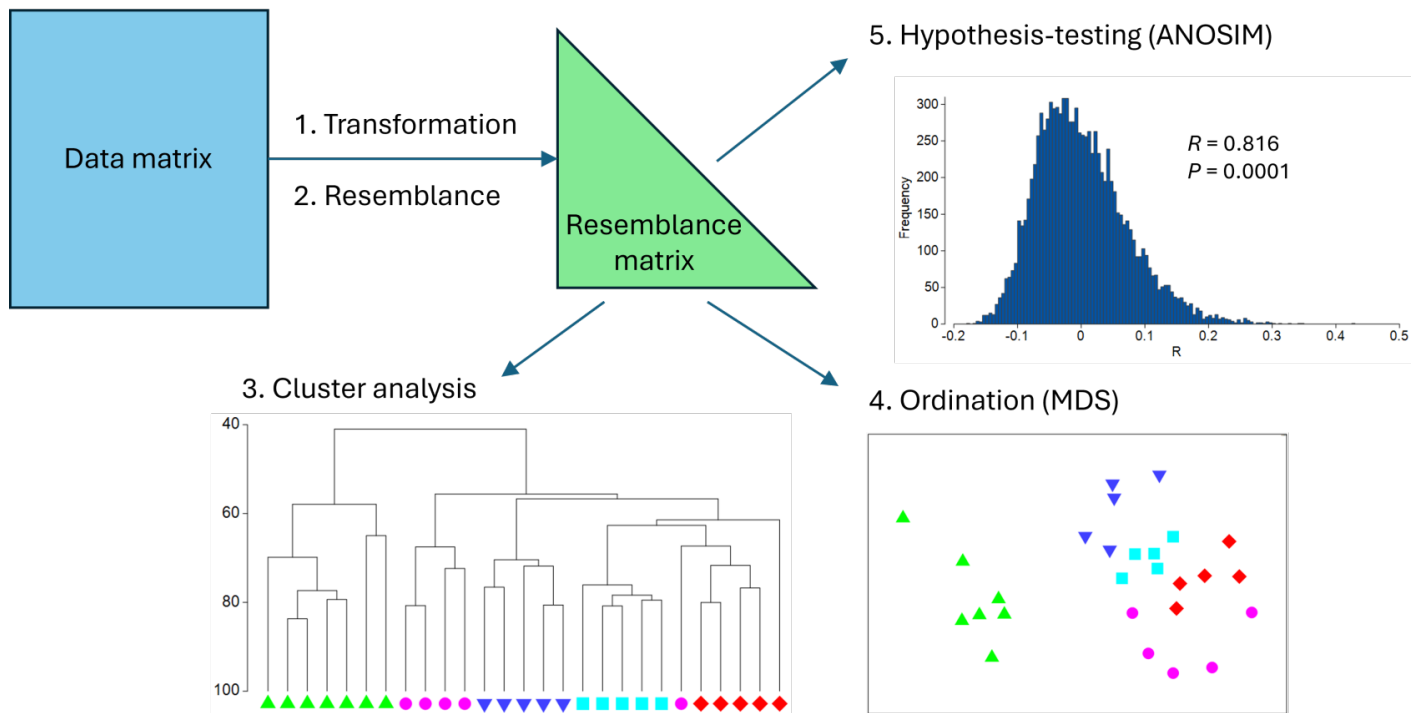


Fig. 5.1. Schematic diagram of a simple multivariate analysis pathway, including transformation, resemblance, clustering, ordination and hypothesis-testing.

Some steps for analysis

In what follows, we shall take the 'Fal nematode.pri' example dataset in PRIMER and proceed along a simple multivariate analysis pathway according to the following steps:

1. Apply a fourth-root **transformation** to the data (all variables).
2. Calculate the Bray-Curtis **resemblance** among all pairs of samples.
3. Calculate a hierarchical group-average **cluster** analysis of the samples to produce a dendrogram.
4. Create a non-metric multi-dimensional scaling (nMDS) **ordination** plot of the samples.
5. Use the analysis of similarities (ANOSIM) to **test the null hypothesis** of no differences in community structure among the 4 creeks of the Fal estuary (with $n = 5$ to 7 replicates per creek).

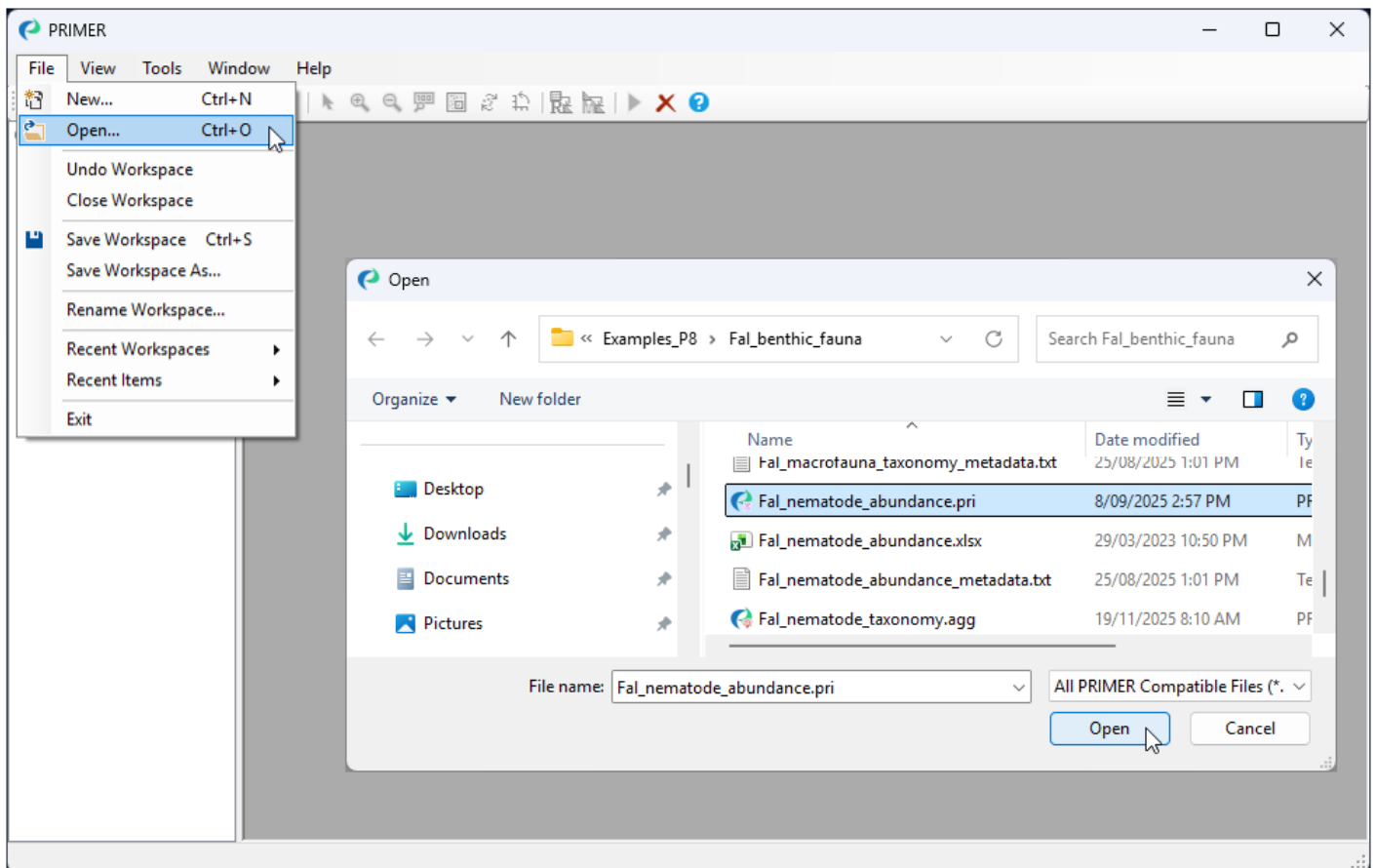
3. Example analysis pathway (CLUSTER, MDS, ANOSIM)

Step 1: Transformation

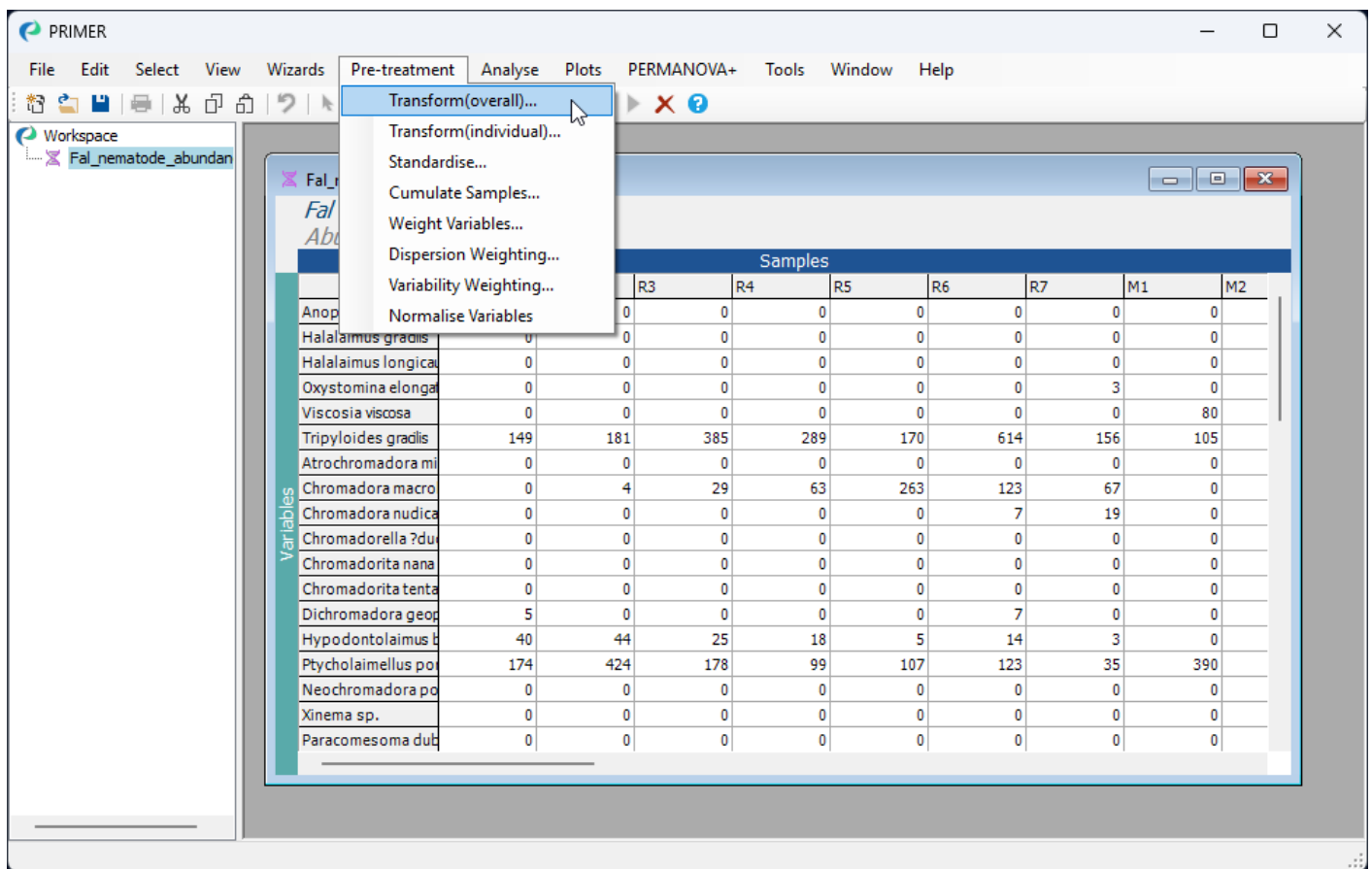
Perform the transformation

There are a range of possible transformations one might use in a pre-treatment of biotic data. For some discussion regarding appropriate pre-treatment transformations, see [Chapter 9](#) in '*Change in Marine Communities*'. In the case of the Fal data set, we may choose to downweight considerably the contribution of highly abundant taxa. Our first (pre-treatment) step will be to transform the data to fourth roots, i.e., $y' = y^{\{0.25\}}$.

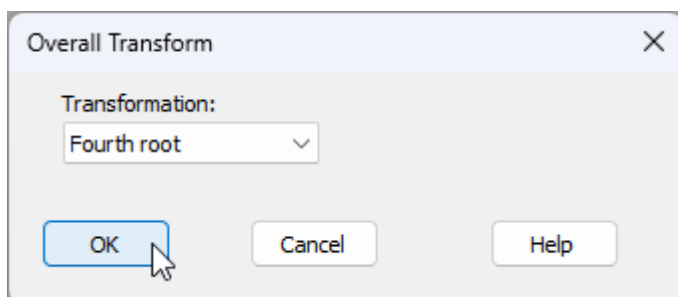
1. Launch PRIMER 8 and click on **File > Open...**, choose the trial data file 'Fal nematode abundance.pri' (found in 'Examples_P8' > 'Fal_benthic_fauna') in the browser window and click **Open**.



2. To perform an overall transformation of the data (i.e., on all variables), click **Pre-treatment > Transform(overall)....**



3. In the 'Overall Transform' dialog box, choose (Transformation: **Fourth root**), then click **OK**.



You should now see the transformed dataset (called '**Data1**') as the 'active window' in the PRIMER workspace. Note that the transformed datasheet has inherited all of the properties, factors and indicators that were associated with the original data as well.

The screenshot shows the PRIMER software interface. The main window is titled 'Overall Transform' and displays a table of data for 'Fal estuary nematodes Abundance'. The table has columns for samples (R1, R2, R3, R4, R5, R6, R7) and rows for various nematode species. The 'Data1' window is also open, showing the same data table. The 'Workspace' panel on the left shows the project structure, including 'Fal_nematode_abundant', 'Overall Transform1', and 'Data1'.

	R1	R2	R3	R4	R5	R6	R7
Anoplostoma vivax	0	0	0	0	0	0	0
Halalaimus gradis	0	0	0	0	0	0	0
Halalaimus longicauda	0	0	0	0	0	0	0
Oxystomina elongata	0	0	0	0	0	0	1.3161
Viscosia viscosa	0	0	0	0	0	0	0
Tripyloides gradis	3.4938	3.6679	4.4296	4.1231	3.6109	4.9779	3.5341
Atrochromadora mi	0	0	0	0	0	0	0
Chromadora macro	0	1.4142	2.3206	2.8173	4.0271	3.3302	2.861
Chromadora nudica	0	0	0	0	0	1.6266	2.0878
Chromadorella ?du	0	0	0	0	0	0	0
Chromadorita nana	0	0	0	0	0	0	0
Chromadorita tenta	0	0	0	0	0	0	0

You will also see a small file (called 'Overall Transform1'), indicated by a little notepad () in the Explorer tree. This shows a little more information regarding the choices you made for this particular run of the '**Transformation(overall)...**' routine.

The screenshot shows the 'Overall Transform1' notepad window. It contains the following information:

Overall Transform
Downweight high values

Data worksheet
Name: Fal_nematode_abundance
Data type: Abundance
Sample selection: All
Variable selection: All

Parameters
Transform: Fourth root

Outputs
Worksheet: Data1

The 'active window'

The active window is identified by slightly darker shading of its border within the PRIMER workspace. Its name is also shaded in the Explorer tree. When you want to run any particular routine in PRIMER, you will need to make sure that the correct dataset (or resemblance matrix) is

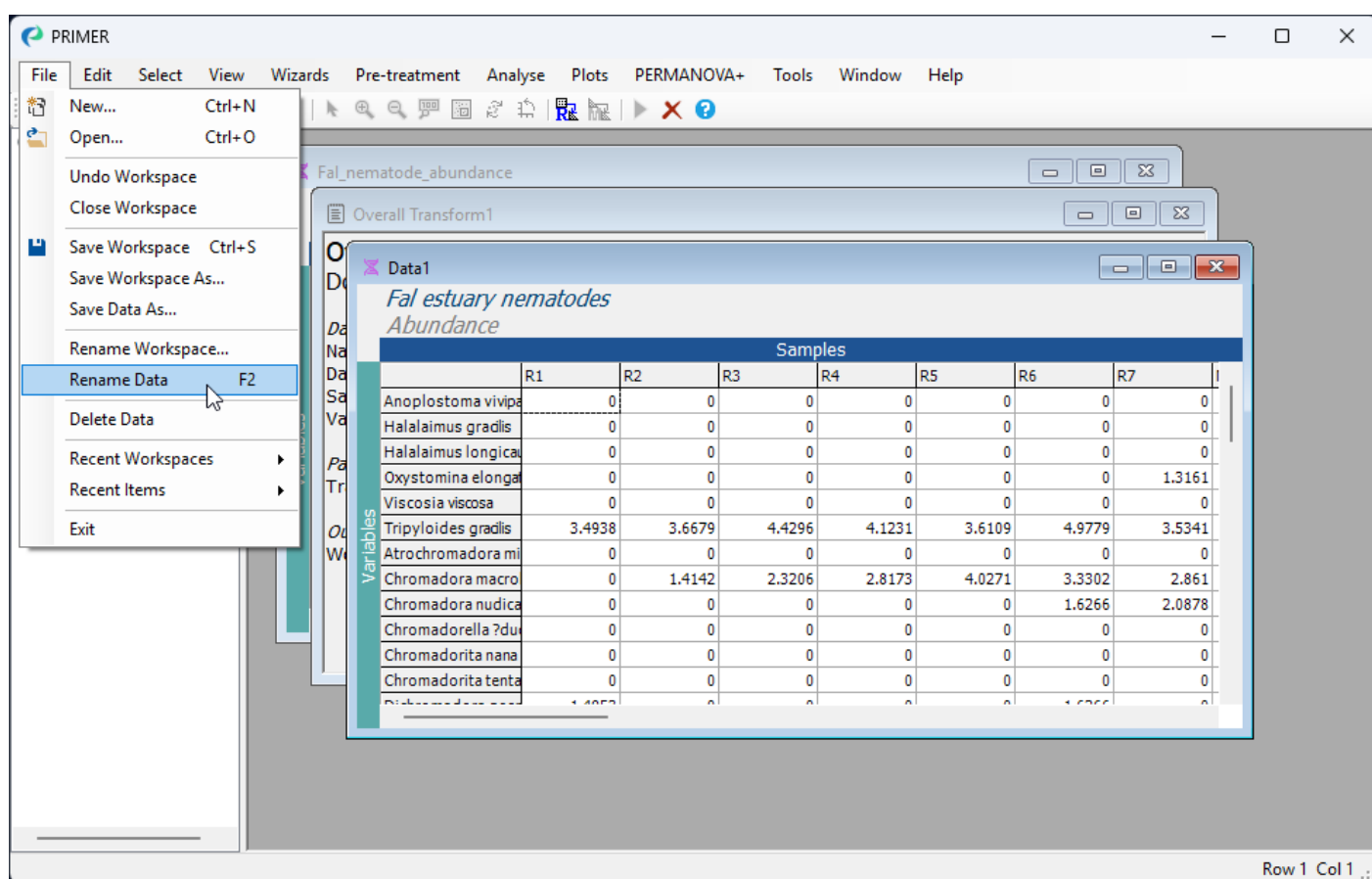
the active window when you launch that routine.

Notice that, although the transformed data ('Data1') is output as the active window after you do a transformation, the original datasheet is still in your workspace, and you can navigate back to it either by clicking on it's name in the Explorer tree (look right under the word 'Workspace' for 'Fal_nematode_abundance'), or by clicking on the window containing the datasheet itself inside the PRIMER workspace.

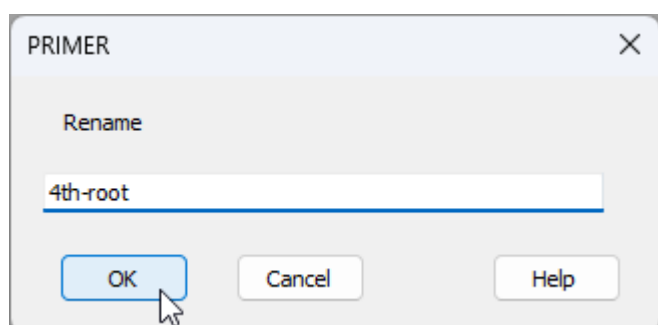
Re-name the transformed data

It is good practice to rename the datasheet corresponding to the transformed data for future reference within this workspace as you proceed with your work.

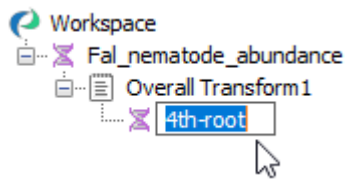
To re-name the datasheet within PRIMER, click on **File > Rename Data**.



In the dialog box, type the new name you want to call it (in this example, we could re-name the sheet containing transformed data as '4th-root', as shown below), then click **OK**.



Alternatively, you can simply click and hover over the word 'Data1' in the Explorer tree window, and you will be able to change the name *in situ* right there, viz:



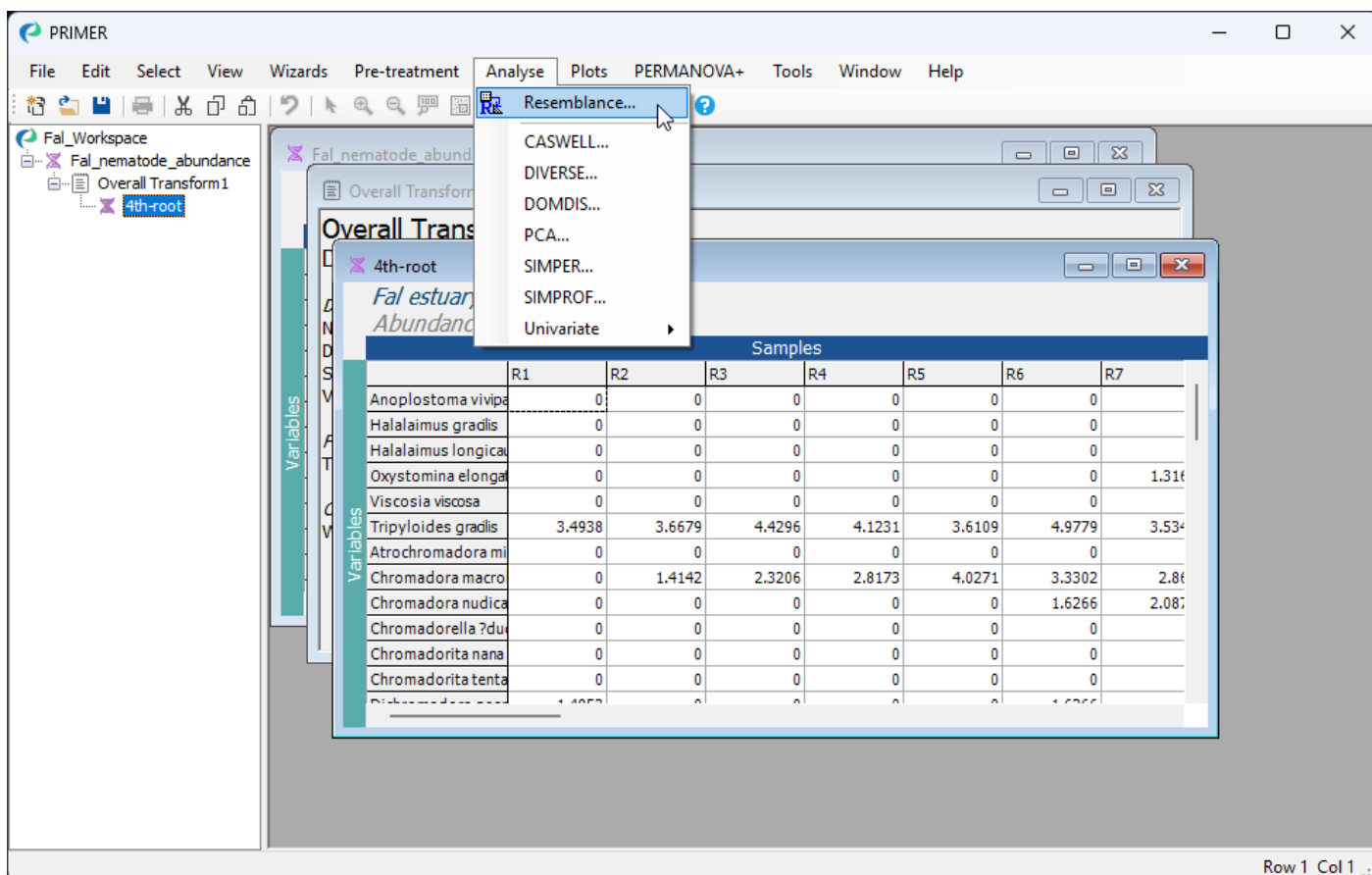
Before heading on to Step 2, you can optionally also now click on **File > Save Workspace** to save your full workspace thus far (you could use the name 'Fal_Workspace.pwk', for example).

3. Example analysis pathway (CLUSTER, MDS, ANOSIM)

Step 2: Resemblance

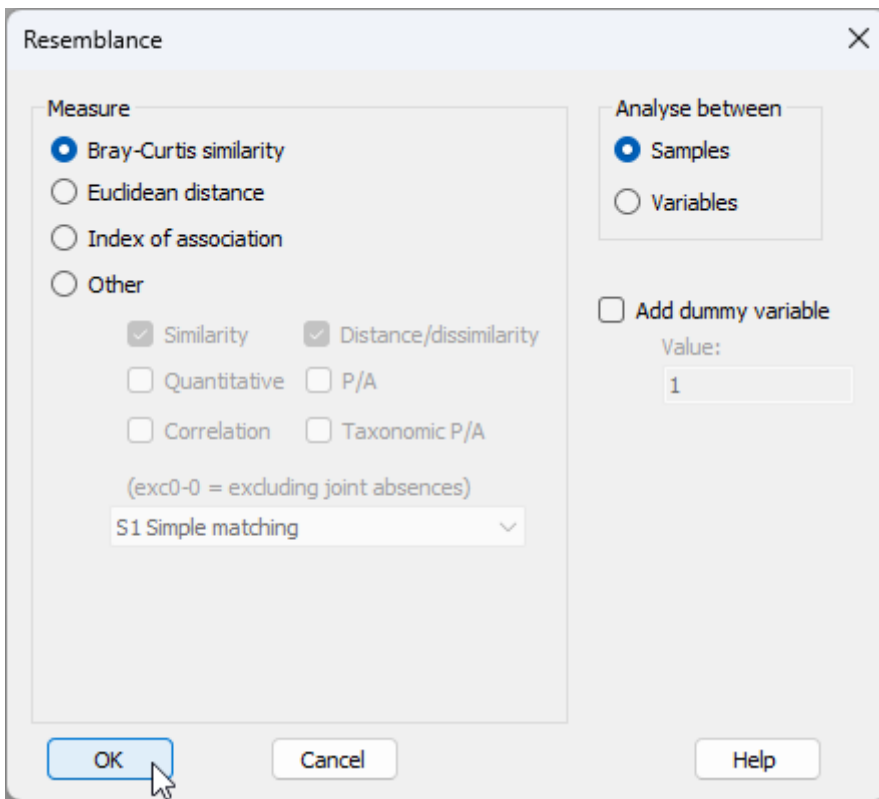
For a description of the Bray-Curtis resemblance measure and the rationale for its use with biotic data, see [Clarke *et al.* \(2006\)](#) and [Chapter 2](#) in '*Change in Marine Communities*'.

1. With the data sheet named '4th-root' from the previous step as the active window, click on **Analyse > Resemblance...**



This dialog has a host of choices for you to consider, which you can uncover by clicking on the '\$\bullet\$Other' radio button. See [Clarke *et al.* \(2006\)](#) and chapter 7 in [Legendre & Legendre \(2012\)](#) for more details regarding the available choices here. For this example, however, we are going to use the Bray-Curtis resemblance measure. (This is the default for data of type 'Abundance').

2. In the 'Resemblance' dialog, under 'Measure', click on \$\bullet\$Bray-Curtis similarity (the default), then click **OK**.



You will see the resulting (lower triangular form of the) matrix of Bray-Curtis similarities between every pair of samples (called 'Resem1').

PRIMER

File Edit Select View Analyse PERMANOVA+ Tools Window Help

Fal_Workspace

- Fal_nematode_abundance
 - Overall Transform1
 - 4th-root
 - Resemblance1
 - Resem1

Fal estuary nematodes
Similarity (0 to 100)

	Samples						
	R1	R2	R3	R4	R5	R6	R7
R1							
R2	64.982						
R3	53.833	67.005					
R4	53.602	78.516	83.701				
R5	54.444	61.138	77.211	78.462			
R6	51.61	60.184	76.57	77.119	79.357		
R7	43.462	55.112	66.678	67.859	70.821	74.032	
M1	40.004	52.708	38.665	46.935	41.623	43.638	39.1
M2	43.811	49.494	52.47	56.677	57.561	58.843	54.0
M3	32.397	37.747	29.797	39.536	32.493	40.112	32.4
M4	36.844	46.684	47.671	56.551	53.946	59.747	53.1
M5	40.491	44.13	38.527	48.595	41.647	44.155	40.7

Row 1 Col 1

Notice that this will inherit relevant properties and also factors that were associated with the dataset from which it was generated. From the 'Resem1' matrix as the active window, click on **Edit** > **Properties** to see the properties associated with this resemblance matrix (shown below); you

can also click on **Edit > Factors** to see the factors (which have not changed).

Resemblance Matrix Properties

Title:
Fal estuary nematodes

Resemblance type

- ☒ Similarity
- ☐ Dissimilarity
- ☐ Distance
- ☐ Distance²
- ☐ Correlation
- ☐ R
- ☐ Rank

Between

- ☒ Samples
- ☐ Variables
- ☐ Other

History:

Transform: Fourth root
Resemblance: S17 Bray-Curtis similarity

Number of row/columns: 27

Description:

File Name(s):
Fal_nematode_abundance.pri (PRIMER 8 format)
Fal_nematode_abundance.xlsx (Excel format)

Dataset Name:

OK Cancel Help

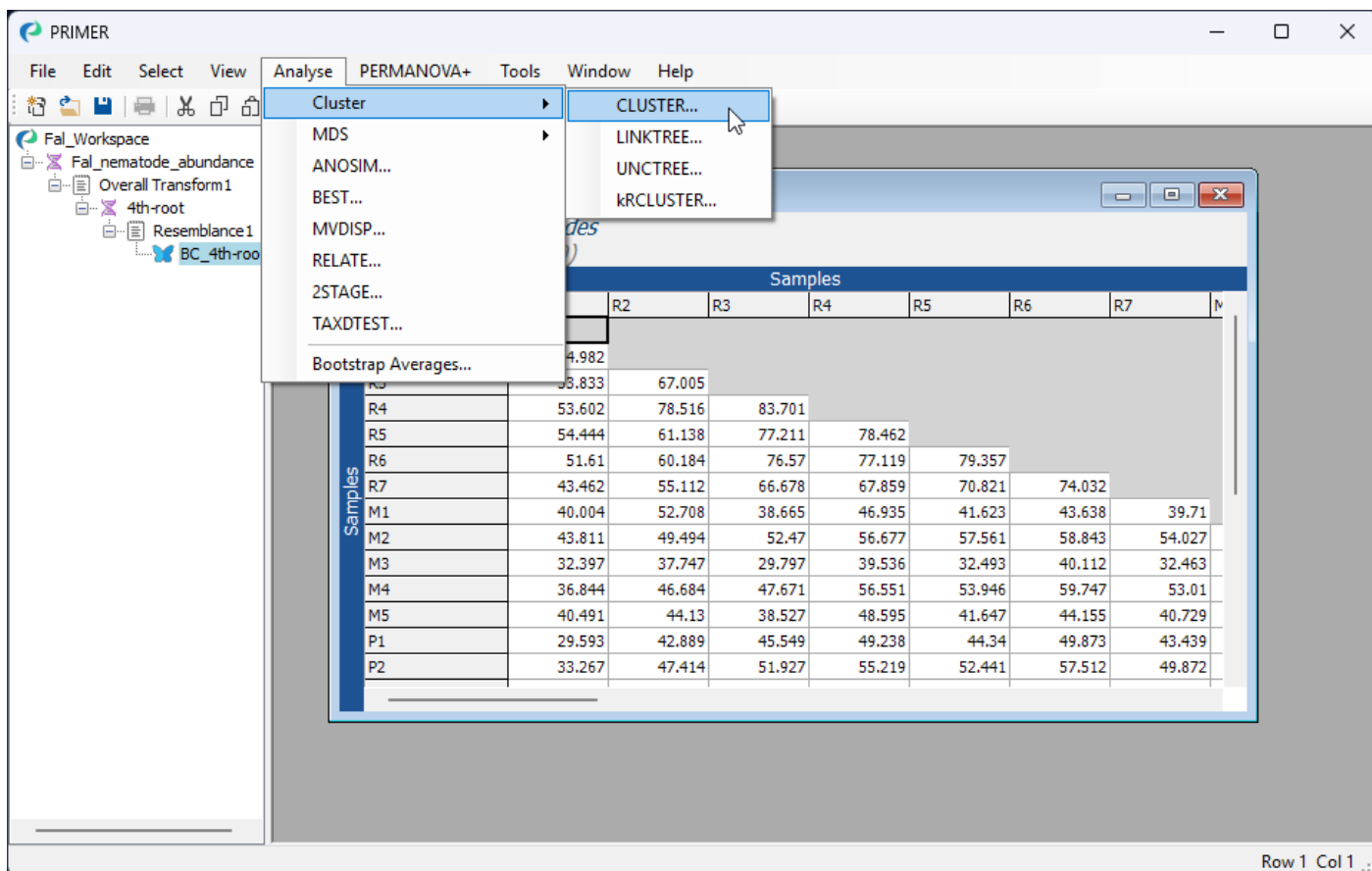
- It is (once again) good practice to re-name the sheet containing the resulting resemblance matrix (currently called 'Resem1') something that would be more useful for future reference. Click on the 'Resem1' sheet, then click on **File > Rename Resem** and call it 'BC_4th-root' for clarity in ensuing analyses.

Step 3: Cluster

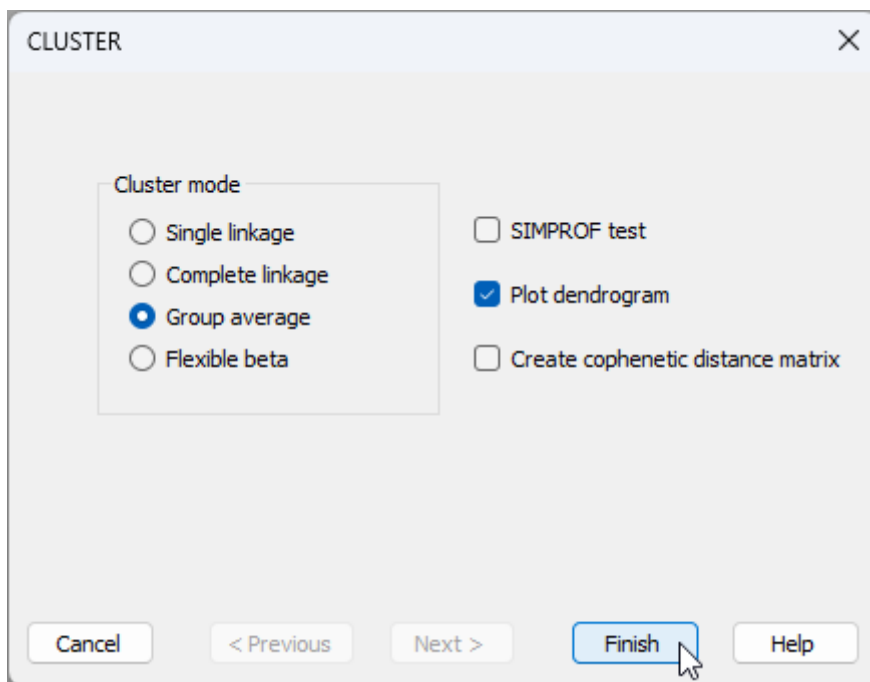
Now that we have a resemblance matrix, we can proceed with a cluster analysis of the data. The purpose of a cluster analysis is to join together samples that are similar to one another, and also to clarify separations of groups (or clusters) of samples that may be dissimilar ([Clarke \(1993\)](#) , [Legendre & Legendre \(2012\)](#)). For descriptions of the clustering methods in PRIMER, see [Chapter 3](#) in '*Change in Marine Communities*'. The starting point for cluster analyses in PRIMER is always the resemblance matrix that quantifies inter-sample (or inter-variable) (dis)similarities.

Perform the cluster analysis

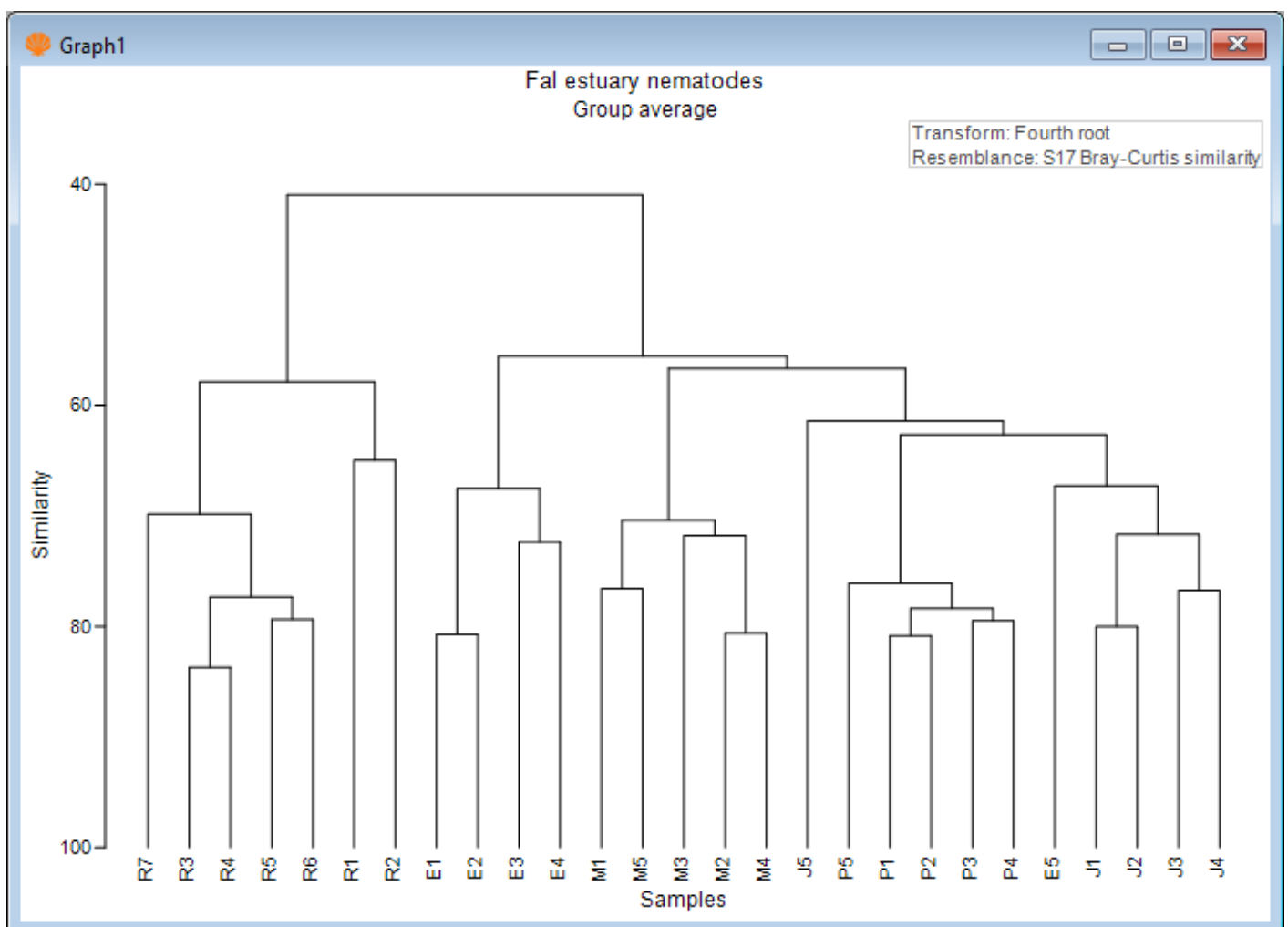
1. In the PRIMER workspace file containing the Fal nematode data, click on the resemblance matrix of Bray-Curtis similarities calculated from fourth-root transformed data (called 'Resem1', or re-named to 'BC_4th-root' in the Explorer tree) to make it the active window, then click on **Analyse > Cluster > CLUSTER...**.



2. To perform hierarchical agglomerative clustering with group-average linkage (also sometimes referred to as 'unweighted pair group method with arithmetic mean' or UPGMA), retain all of the defaults in the 'CLUSTER' dialog and click **Finish**.



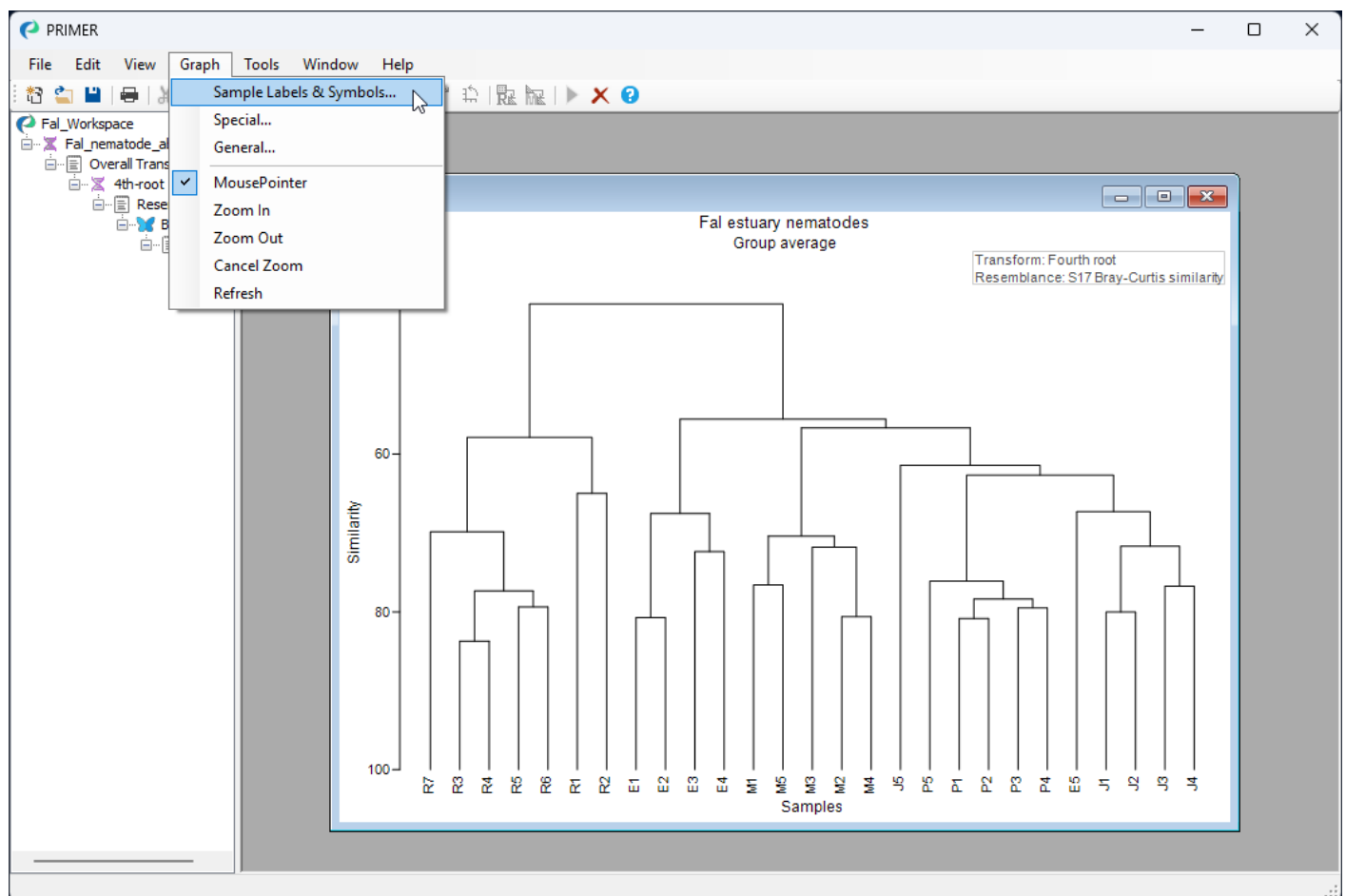
You will see an output file that provides all of the details of the clustering algorithm's pathway (shown as 'CLUSTER1' in the Explorer tree). A dendrogram will also appear as output in a new graphics window (shown as 'Graph1' in the Explorer tree), viz:



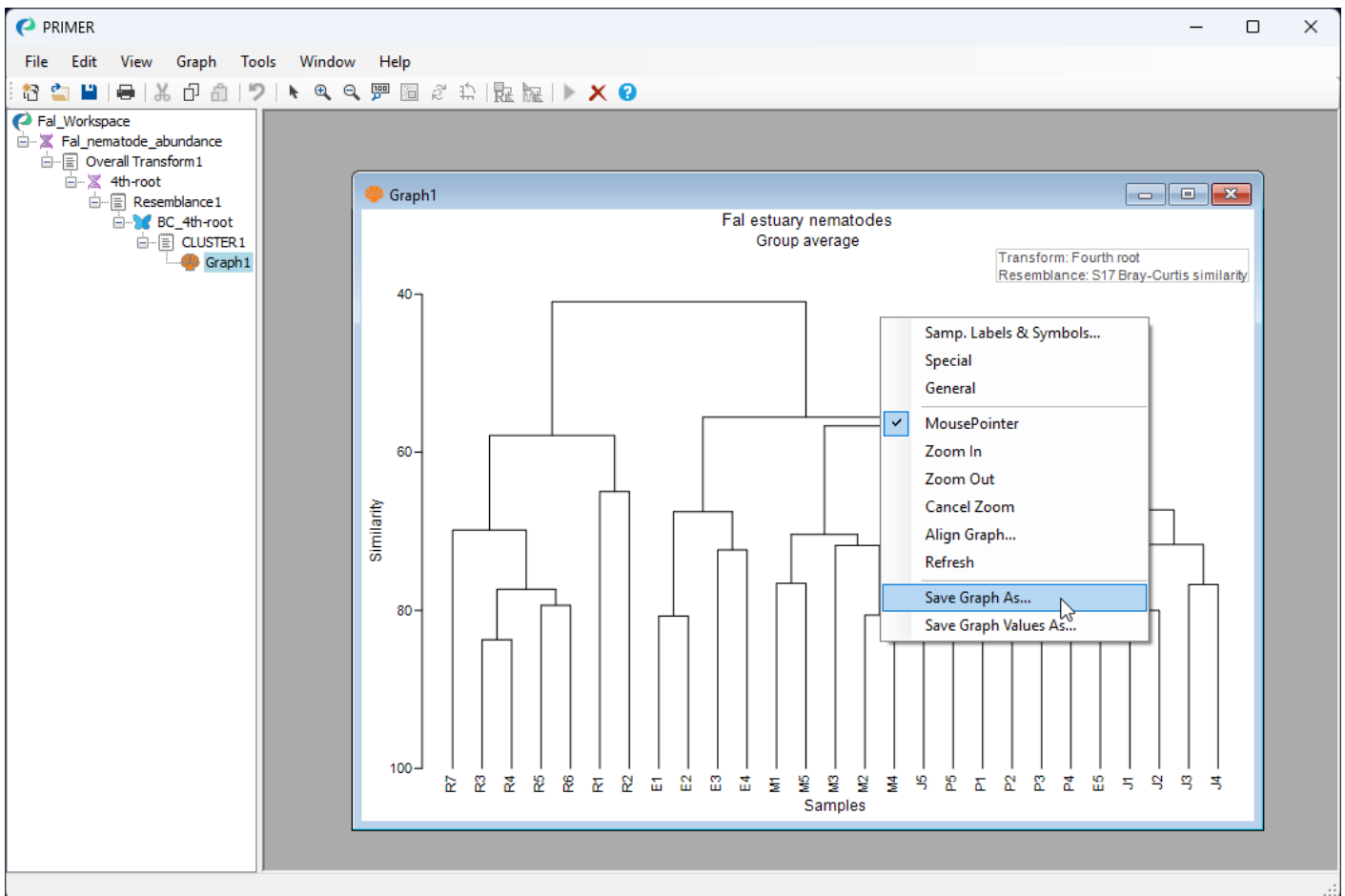
Customise the dendrogram graphic

You can do a number of customisations to the dendrogram graphic.

3. With the dendrogram as the active window, click on **Graph** and you will see a range of options.



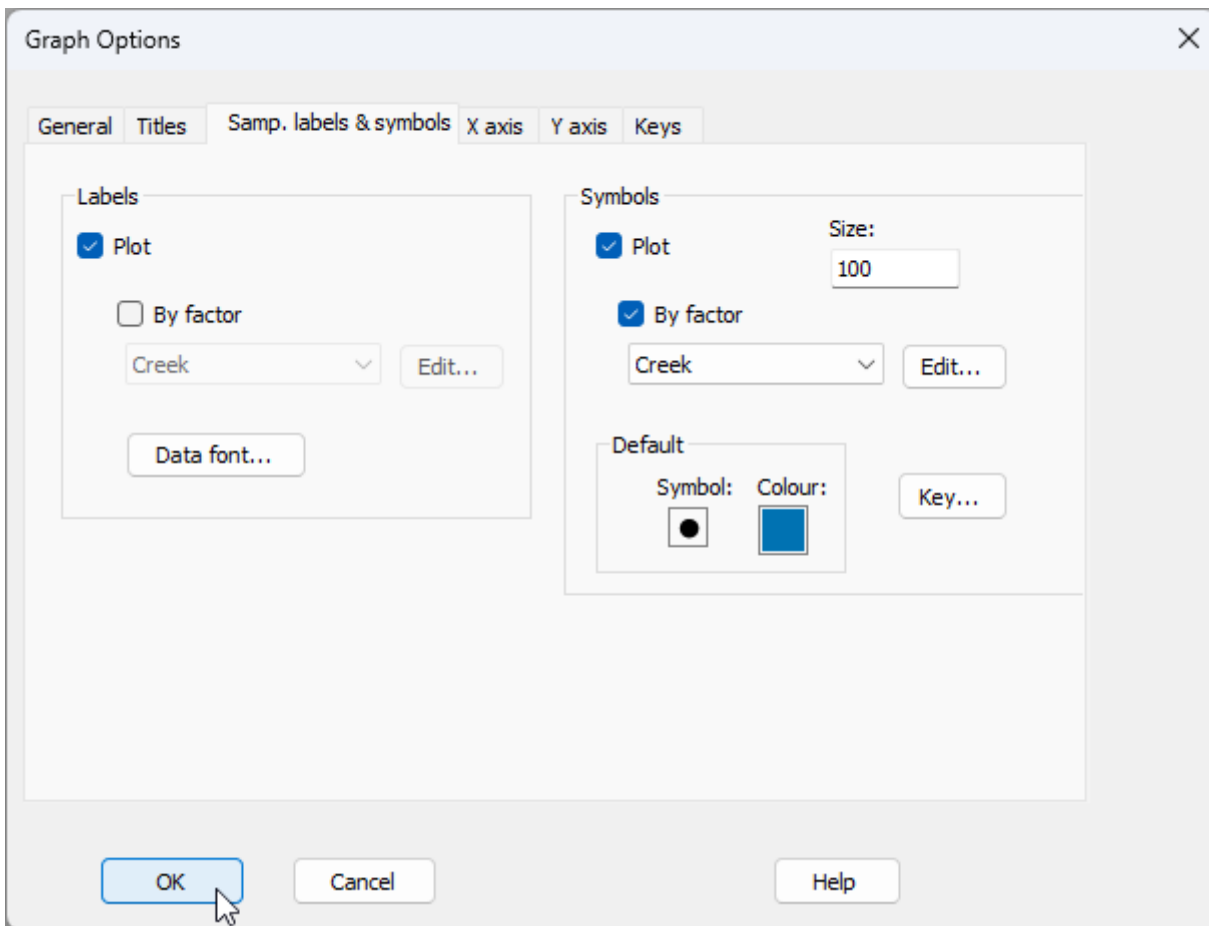
Alternatively, you can right-click anywhere inside the dendrogram's graphic window itself and a similar list pops up, along with additional options for saving the graphic (or graph values) as a separate file, etc. Note that you can also save any graph by clicking on **File > Save Graph As...** with the graphic you want as the active window.



The choices of **Graph > Sample Labels & Symbols...** and also **Graph > General...** will show you a set of options for changing elements that are common to all (or most) graph types in PRIMER (i.e., things like scales and labels to use for the x and y axes, titles and sub-titles, fonts, symbols, colours, etc.). In contrast, the choice of **Graph > Special..** will provide a special menu of graphical options *specific to that particular type of graph* - in this case, a dendrogram.

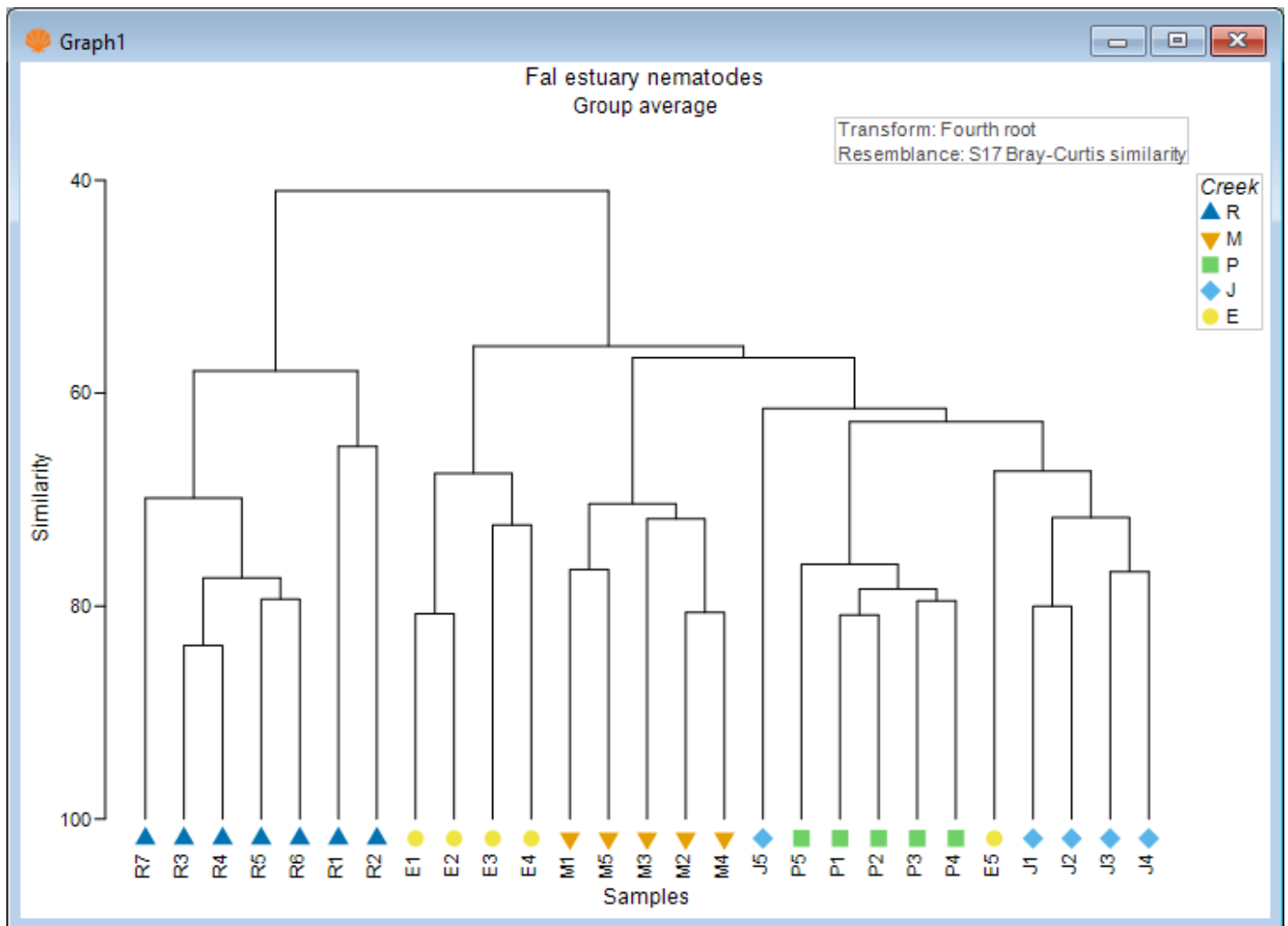
Change labels and/or symbols

- Let's add some symbols to the plot, with the 4 creeks being represented by different symbols. Click **Graph > Sample Labels & Symbols...**. In the resulting 'Graph Options' dialog, under the 'Samp. labels & symbols' tab, choose (Labels ☒ Plot) & (Symbols ☒ Plot & ☒ By factor **Creek**), then click **OK**.



Note that the 'Samp. labels & symbols' dialog will automatically offer you the option to choose labels or symbols corresponding to any factors that are associated with your resemblance matrix.
Tip: For graphing purposes, it is a good idea to use short names for samples (if you can) and (more importantly) for levels of factors, so they will be easy to display on plots.

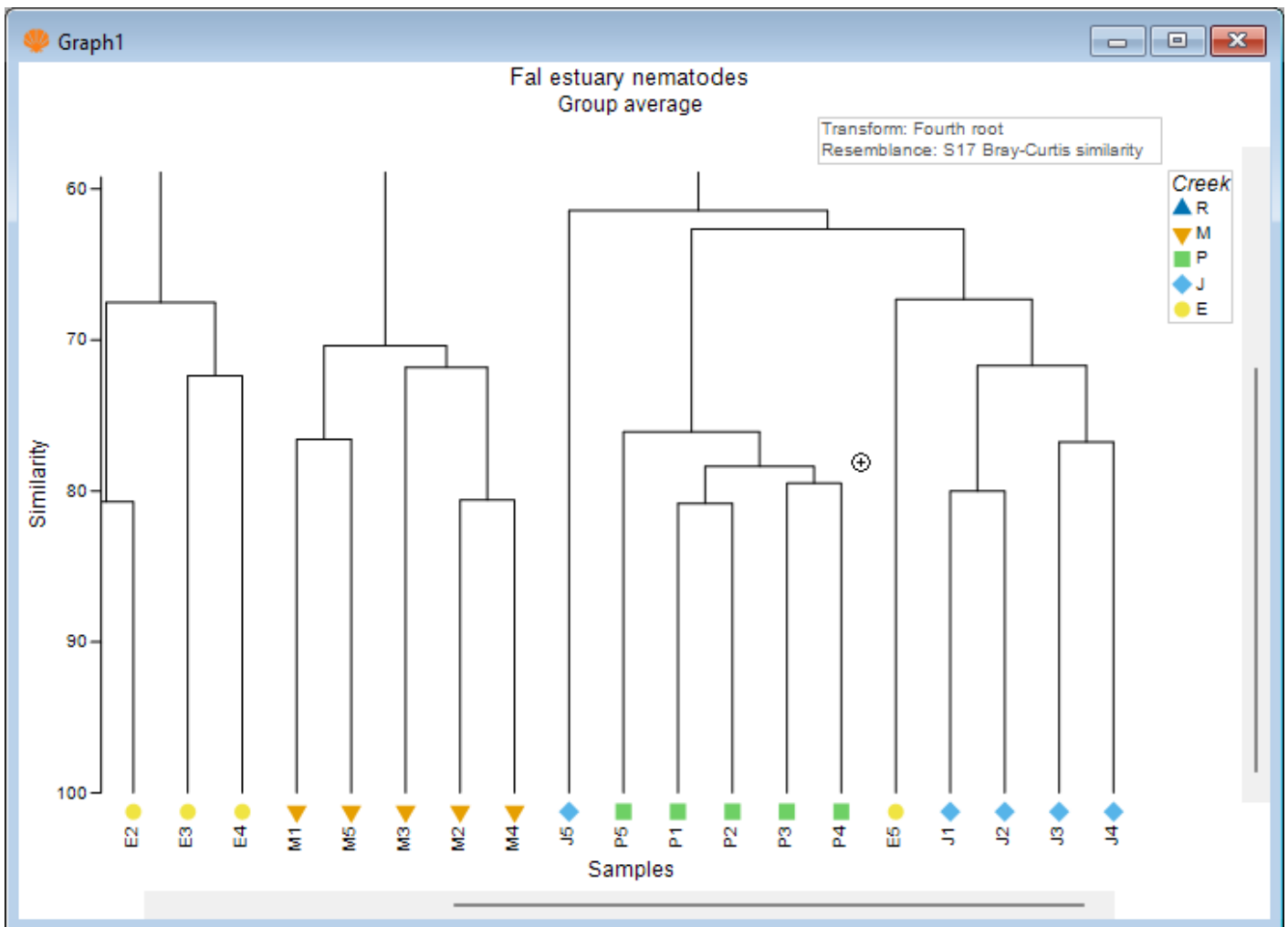
The result is a more colourful and informative dendrogram, which clearly shows a pattern of (generally) high similarity among assemblages within each of the four creeks, (similar symbols tend to be clustered together), and some clear distinctions between assemblages in different creeks, with Restronguet Creek ('R') appearing most different from the others (only about ~40% similar to other creeks, on average).



Zoom in and out

If you have a very large dendrogram with a lot of samples, it is helpful to be able to zoom in and out in order to see details present in different portions of the dendrogram.

To Zoom in, click **Graph > Zoom in** (or click on the 'zoom-in' icon  in the ribbon of tools shown under the main menu items). In zoom-in mode, your cursor will change into a little 'plus' symbol  when poised over your graphic, and with this you simply click on the dendrogram to zoom in on that area. Click multiple times to continue zooming in.

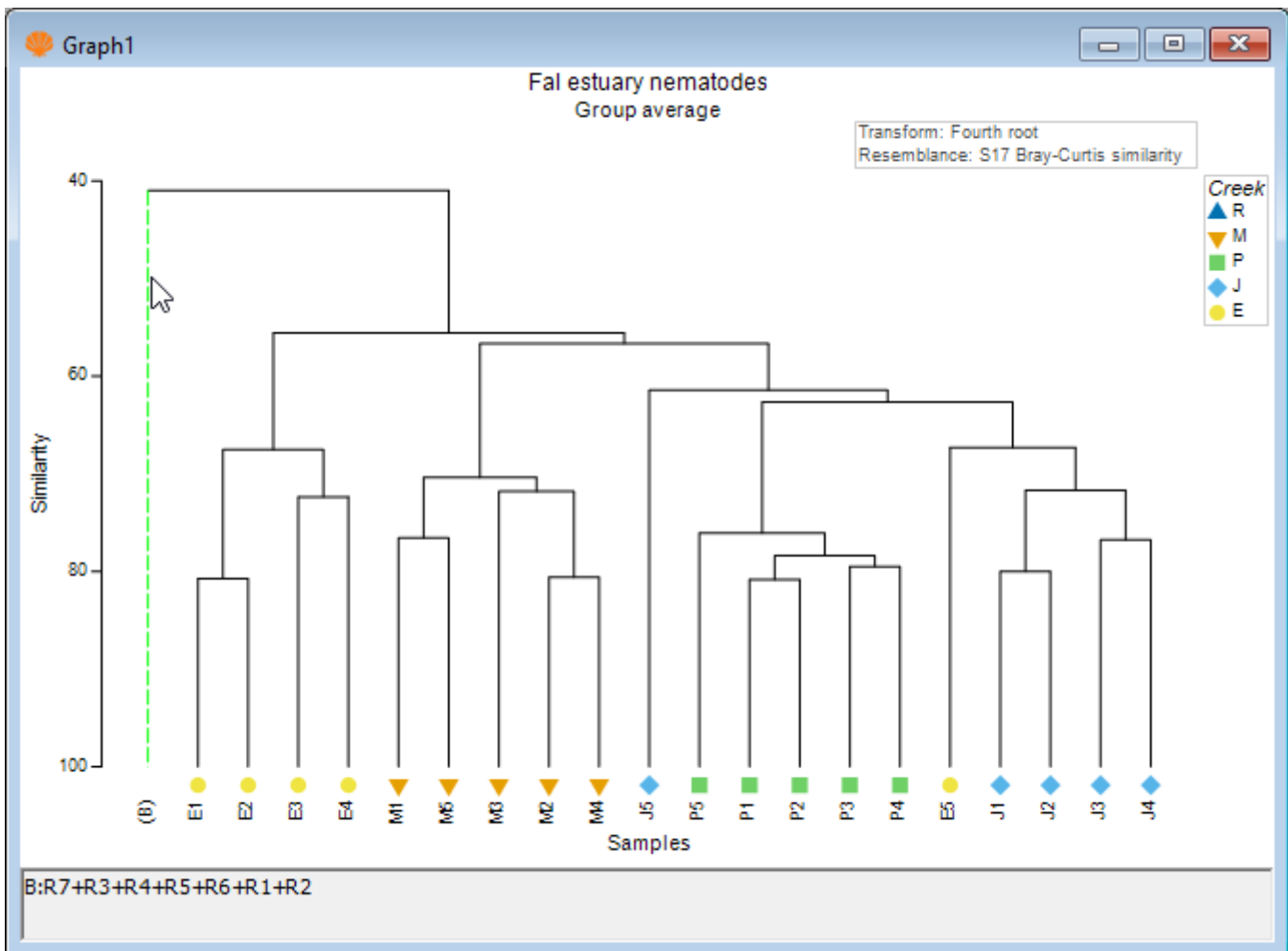


To Zoom out, click **Graph > Zoom out** (or click on the little icon with a minus symbol inside a magnifying glass ). Now the cursor changes into a little 'minus' symbol , and clicking on the graphic zooms you back out again.

To cancel the zoom entirely (to see the entire dendrogram again), click on **Graph > Cancel Zoom** (or click on the cancel-zoom icon .

Collapse nodes to simplify

Another option, if you have a very large dendrogram, is to simplify the view by clicking on any one of the vertical branches. This will collapse all of the samples under that branch into a single entity, making it easier to see patterns in what remains. For example, click on the vertical branch that leads to all of the samples from Restronguet Creek ('R'), and the dendrogram will look like this:



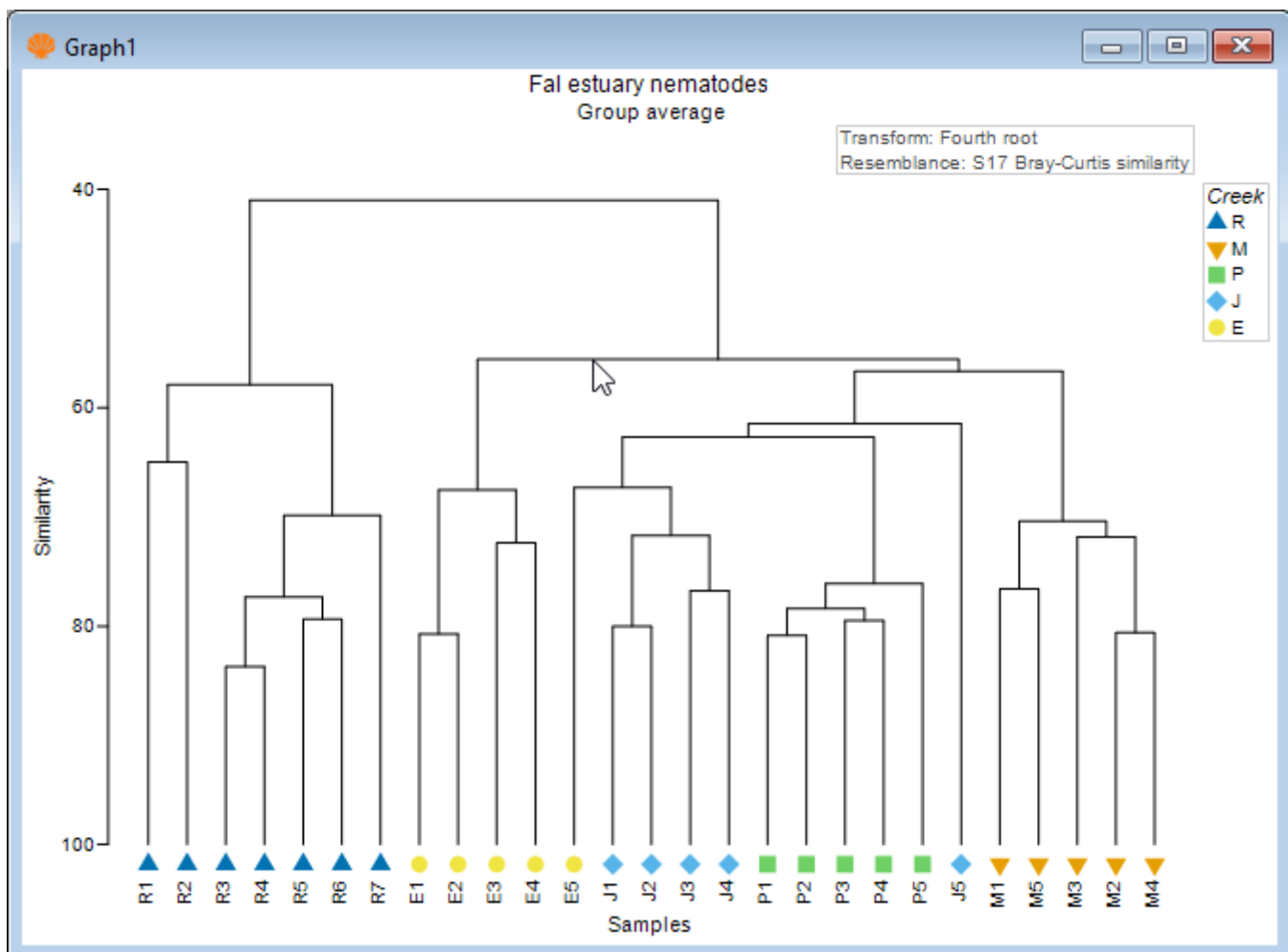
Simply click on that same vertical branch again to toggle that 'collapse' option off, and you will see all of the individual samples within that node again in the full dendrogram, as before.

Rotate leaves around the nodes

It is important to remember that the dendrogram should be thought of like a 'mobile', dangling from a ceiling. In this way, the leaves dangling from each node can be rotated around the node. Thus, the specific ordering of the samples you see in the dendrogram, while being constrained by it in some way, has a bit of flexibility: sample units can be rotated around the nodes.

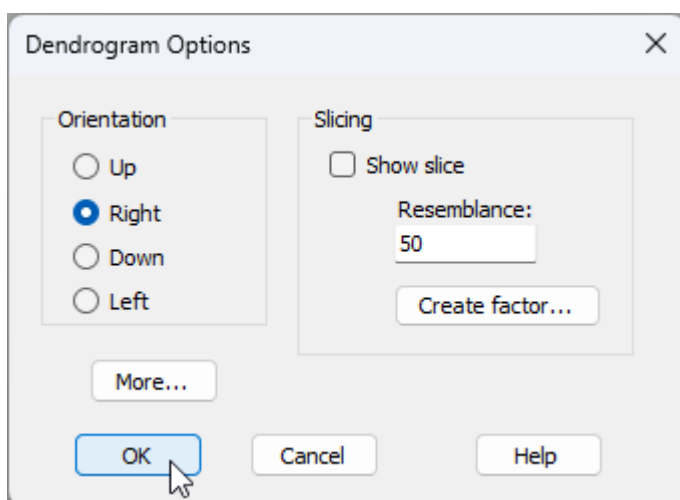
To explore this flexibility visually, click on any horizontal line of the dendrogram, and you will see its leaves will 'flip', accordingly. Clearly, quite a lot of re-arrangements of the sample ordering provided in the initial output are possible, so it is often useful to experiment with this a bit to achieve an ordering that may be most useful for interpretation of relationships/clusters.

For example, one possible re-ordering looks like this:



Change the orientation

The default output orientation for a dendrogram in PRIMER 8 is to have the hierarchical agglomeration going 'Up' (i.e., individual sampling units hang downwards); however, you can customise its orientation. For example, for the Fal estuary nematodes, click on **Graph > Special...**, then in the 'Dendrogram Options' dialog box, under 'Orientation', choose **Right**, and click **OK**.



This yields the following:

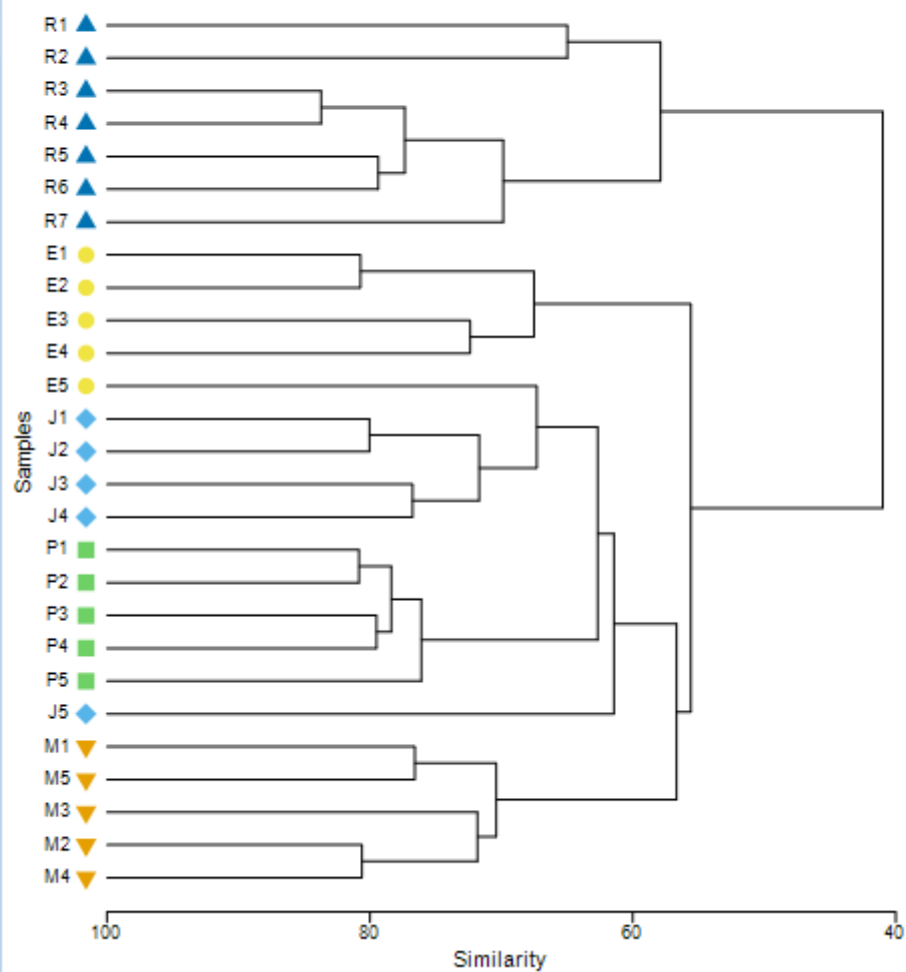
Fal estuary nematodes

Group average

Transform: Fourth root

Resemblance: S17 Bray-Curtis similarity

Creek

▲ R
▼ M
■ P
◆ J
● E

Step 4: Ordination

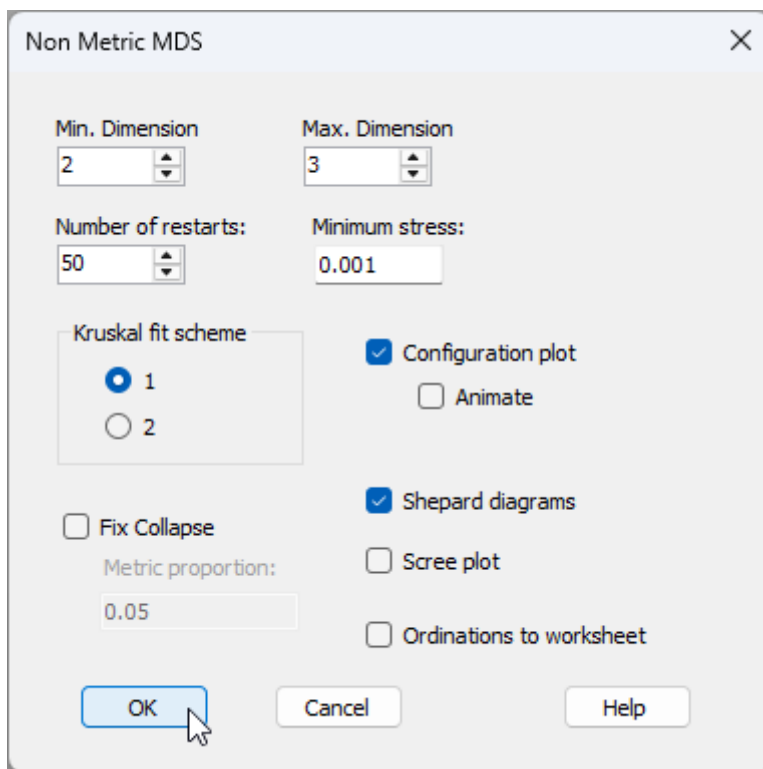
The cluster analysis goes some way towards helping us to understand potential patterns of similarity among the samples. It is particularly good at showing us clusters of samples that are highly similar. To better visualise patterns of relationships among all of the samples simultaneously, we can use methods of **ordination**. In particular, non-metric multi-dimensional scaling (nMDS) is a wonderfully robust tool for visualising high-dimensional systems on the basis of a chosen resemblance measure ([Kruskal & Wish \(1978\)](#) ; [Clarke \(1993\)](#)). Non-metric MDS essentially places points into an arbitrary Euclidean space of set dimension (typically producing a 2D or 3D plot) so as to preserve, as well as possible, the rank-order of the dissimilarities between pairs of samples. For more details on multi-dimensional scaling, see [Chapter 5](#) of '*Change in Marine Communities*'.

Create a non-metric MDS plot

From the Bray-Curtis resemblance matrix ('BC_4th-root' in this example), click **Analyse > MDS > Non-metric MDS...**, then click **OK** to run the MDS algorithm with all of the default options.

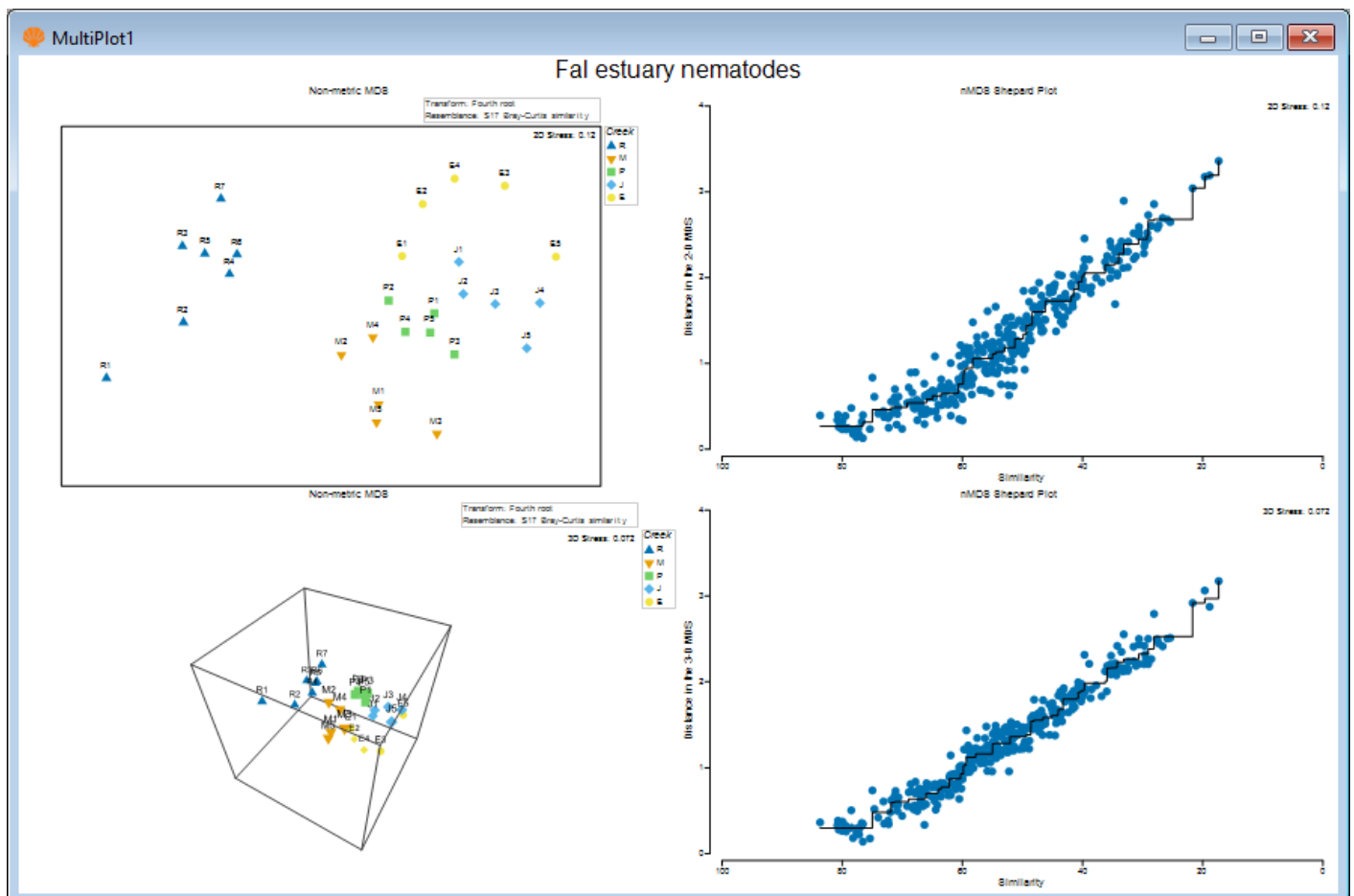
The screenshot shows the PRIMER software interface. The 'Analyse' menu is open, and the 'MDS' option is selected, which has opened a sub-menu where 'Non-metric MDS (nMDS)...' is highlighted. In the background, a table of sample dissimilarities is visible. The table has columns labeled R2, R3, R4, R5, R6, R7, and M1, and rows labeled R4, R5, R6, R7, M1, M2, M3, M4, M5, P1, P2, P3, and P4. The values represent dissimilarities between the samples.

	R2	R3	R4	R5	R6	R7	M1
R4	53.602	78.516	83.701				
R5	54.444	61.138	77.211	78.462			
R6	51.61	60.184	76.57	77.119	79.357		
R7	43.462	55.112	66.678	67.859	70.821	74.032	
M1	40.004	52.708	38.665	46.935	41.623	43.638	39.71
M2	43.811	49.494	52.47	56.677	57.561	58.843	54.027
M3	32.397	37.747	29.797	39.536	32.493	40.112	32.463
M4	36.844	46.684	47.671	56.551	53.946	59.747	53.01
M5	40.491	44.13	38.527	48.595	41.647	44.155	40.729
P1	29.593	42.889	45.549	49.238	44.34	49.873	43.439
P2	33.267	47.414	51.927	55.219	52.441	57.512	49.872
P3	31.531	40.719	39.818	44.122	40.254	46.193	43.476
P4	36.263	47.423	44.059	51.68	49.344	52.629	44.588

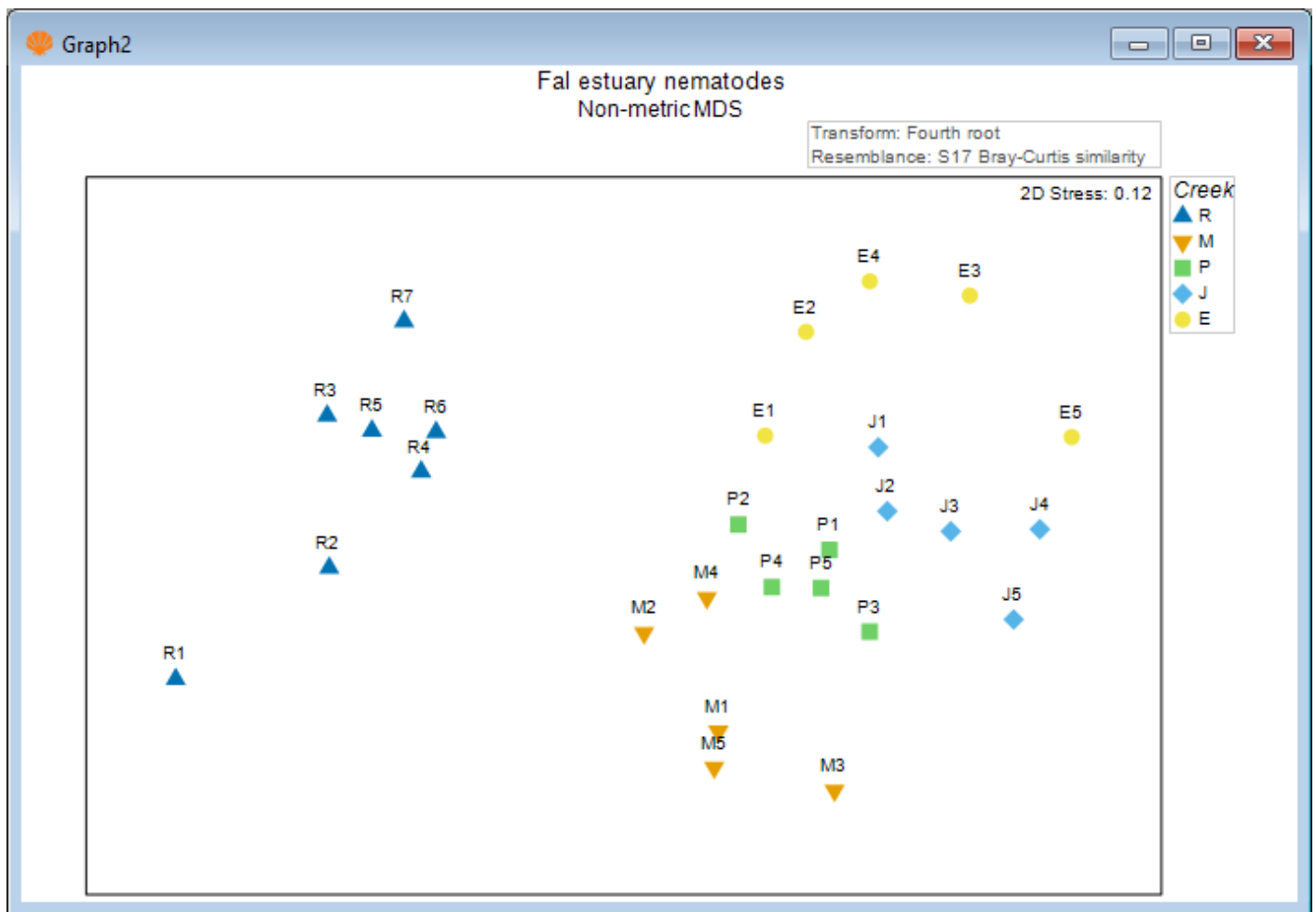


The resulting 'Multi-plot1' output graphic gives you 4 individual plots (in a 2-by-2 array), as follows:

- the best 2D solution (from 50 random starts)
- the Shepard diagram for the best 2D solution
- the best 3D solution (from 50 random starts); and
- the Shepard diagram for the best 3D solution



Clicking on the '+' symbol next to the word 'MultiPlot1' in the Explorer tree reveals the four graphics that comprise this multi-plot object ('Graph2', 'Graph3', 'Graph4', 'Graph5'). If you then click on 'Graph2' in the Explorer tree, or if you click in the Multi-plot itself anywhere on the plot in the upper left-hand corner, you will see a single graphic now; namely the best (lowest-stress) 2-dimensional non-metric MDS plot for the Fal estuary nematodes.



Interpretation

The plot does not have x and y axes, as the scale here is arbitrary. The non-metric MDS algorithm attempts only to preserve rank-order dissimilarities, so interpretations are restricted to relative distances among points in the plot. For example, we can make statements like: *"The dissimilarity between sample R7 and M3 appears to be larger than the dissimilarity between sample M3 and M2"*, but we cannot tell (directly from the plot) what the specific sizes of those particular dissimilarities are.

It is clear from this output that Restronguet Creek ('R') has assemblages of nematodes that are quite dissimilar from those found in the other four creeks. That's consistent with what we saw in the cluster analysis of these data. We can, however, get quite a few more insights from the nMDS plot than were apparent in the cluster dendrogram. For example, relationships among the assemblages from the other four creeks seem to be ordered (from the bottom-left towards the top of the plot) as follows: Percuil → St. Just → Pill → Mylor. It is also apparent that the samples from Pill Creek ('P') form a tighter cluster hence are not as variable as the samples from (say) Restronguet Creek ('R'), which are more spread out. Being able to see potential gradients of change in

assemblages, the degree of between-group differences, as well as the relative sizes of within-group dispersions, is all very useful.

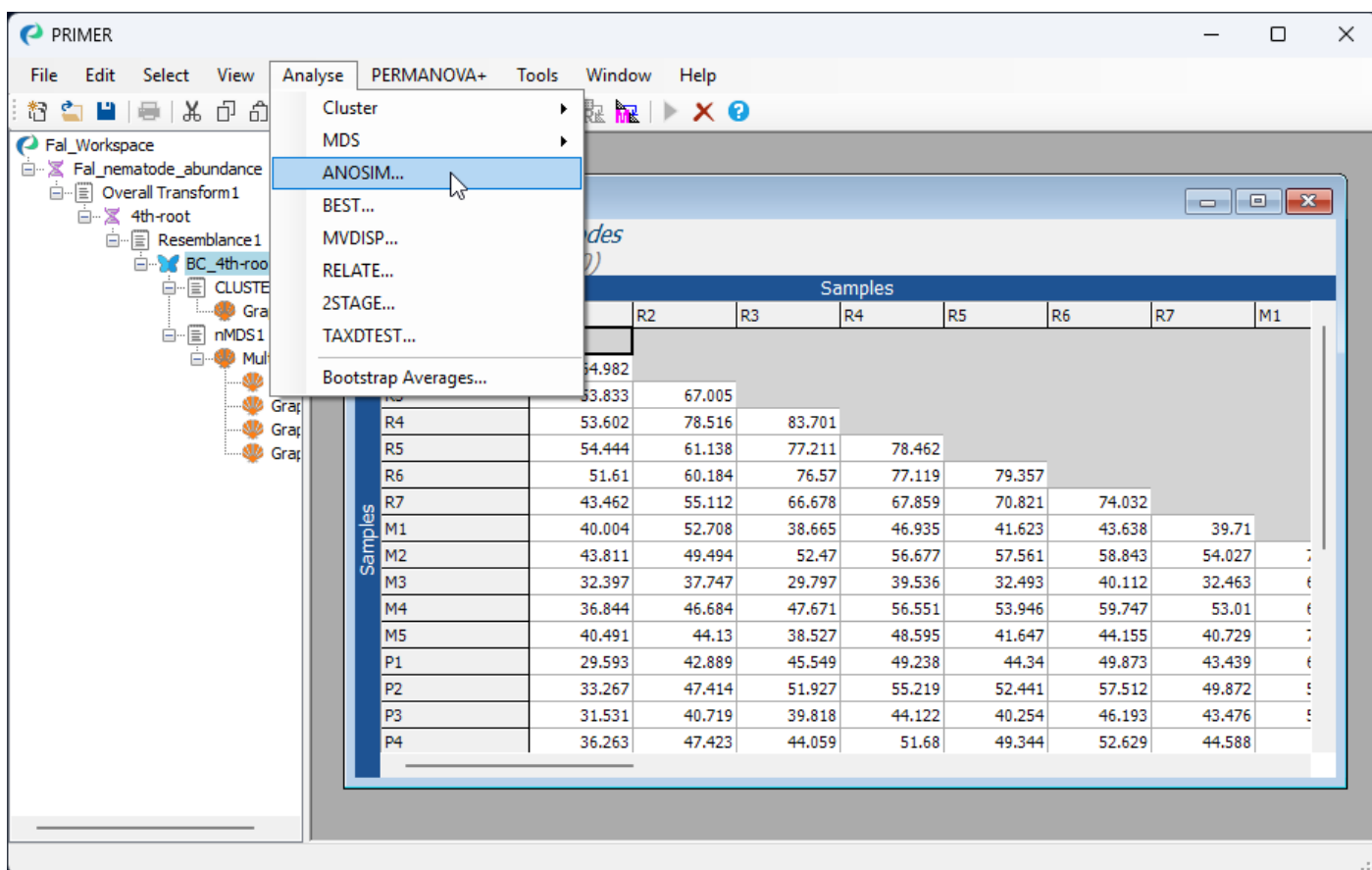
3. Example analysis pathway (CLUSTER, MDS, ANOSIM)

Step 5: ANOSIM test

The ordination plot gives us a visualisation of the rank-order relationships among the samples, based on the dissimilarity measure. Next, we may wish to test the null hypothesis that there are no differences among the five creeks. We can use a non-parametric multivariate permutation test called analysis of similarities (ANOSIM, [Clarke \(1993\)](#)) for this test. For more details on the ANOSIM test, see also [Chapter 6](#) in 'Change in Marine Communities'. We will be doing a one-way ANOSIM test here (as there is a single factor in this example); for more information on multi-way designs with ordered and/or unordered factors, see [Somerfield et al. \(2021a\)](#), [Somerfield et al. \(2021b\)](#) and [Somerfield et al. \(2021c\)](#).

Perform the ANOSIM test

From the Bray-Curtis resemblance matrix (called 'BC_4th-root' in this example), click on **Analyse > ANOSIM...** and you will see the ANOSIM dialog.



Here, we are performing a one-way ANOSIM (there is just one factor whose levels are un-ordered), so we can leave the defaults in the dialog as they are and just click '**OK**'.

ANOSIM

Design

Model:
One-Way - A

Factors	Type	
A: Creek	Unordered	Levels...
B: Creek name	Unordered	Levels...
C: Position	Unordered	Levels...

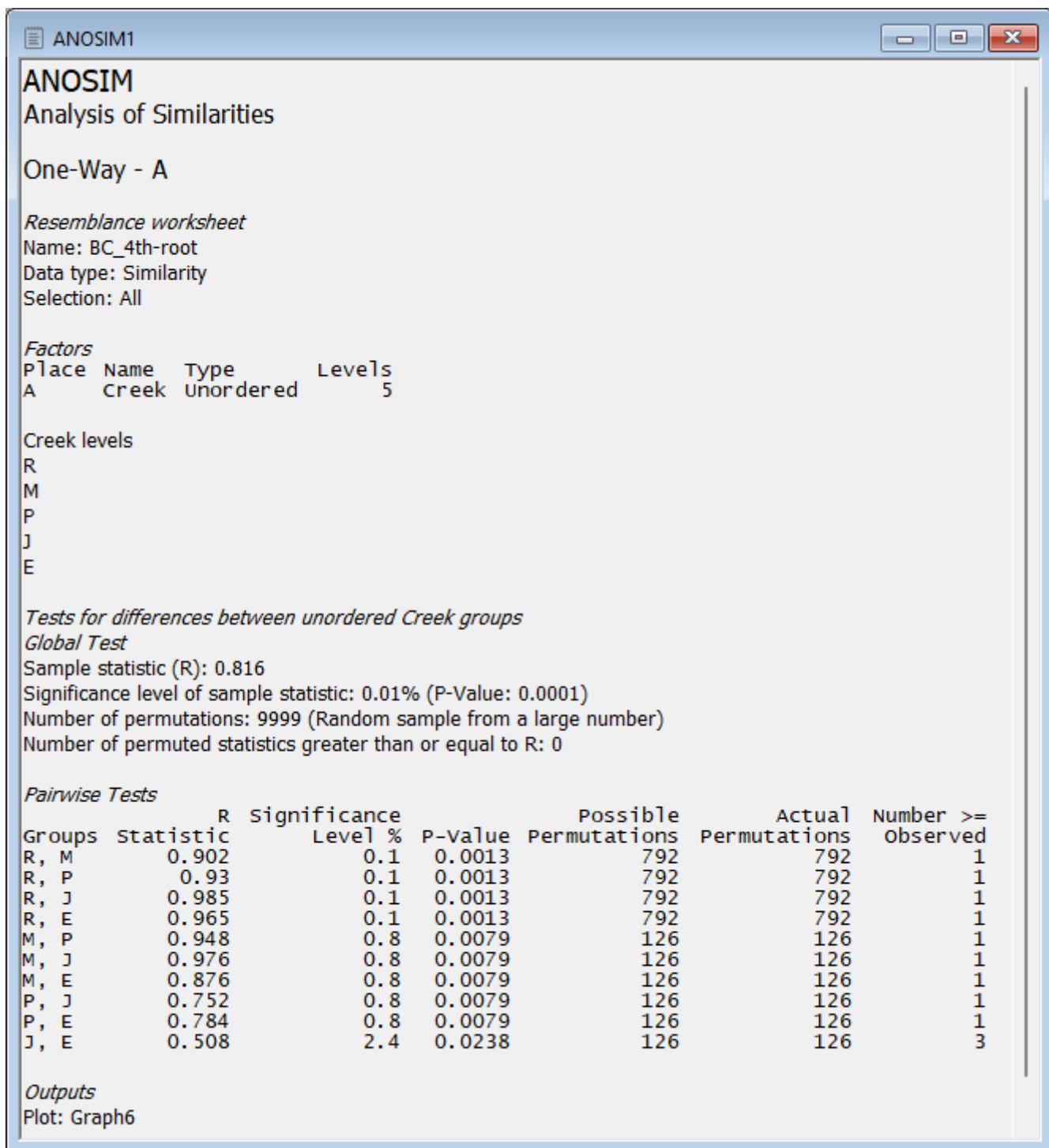
Correlation method:
(multiway tests with no replication)
Spearman rank

☐ Pairwise tests R/rho to worksheet
☐ Pairwise tests sig% to worksheet
☒ Plot histogram
☐ R values to file

Max permutations:
9999

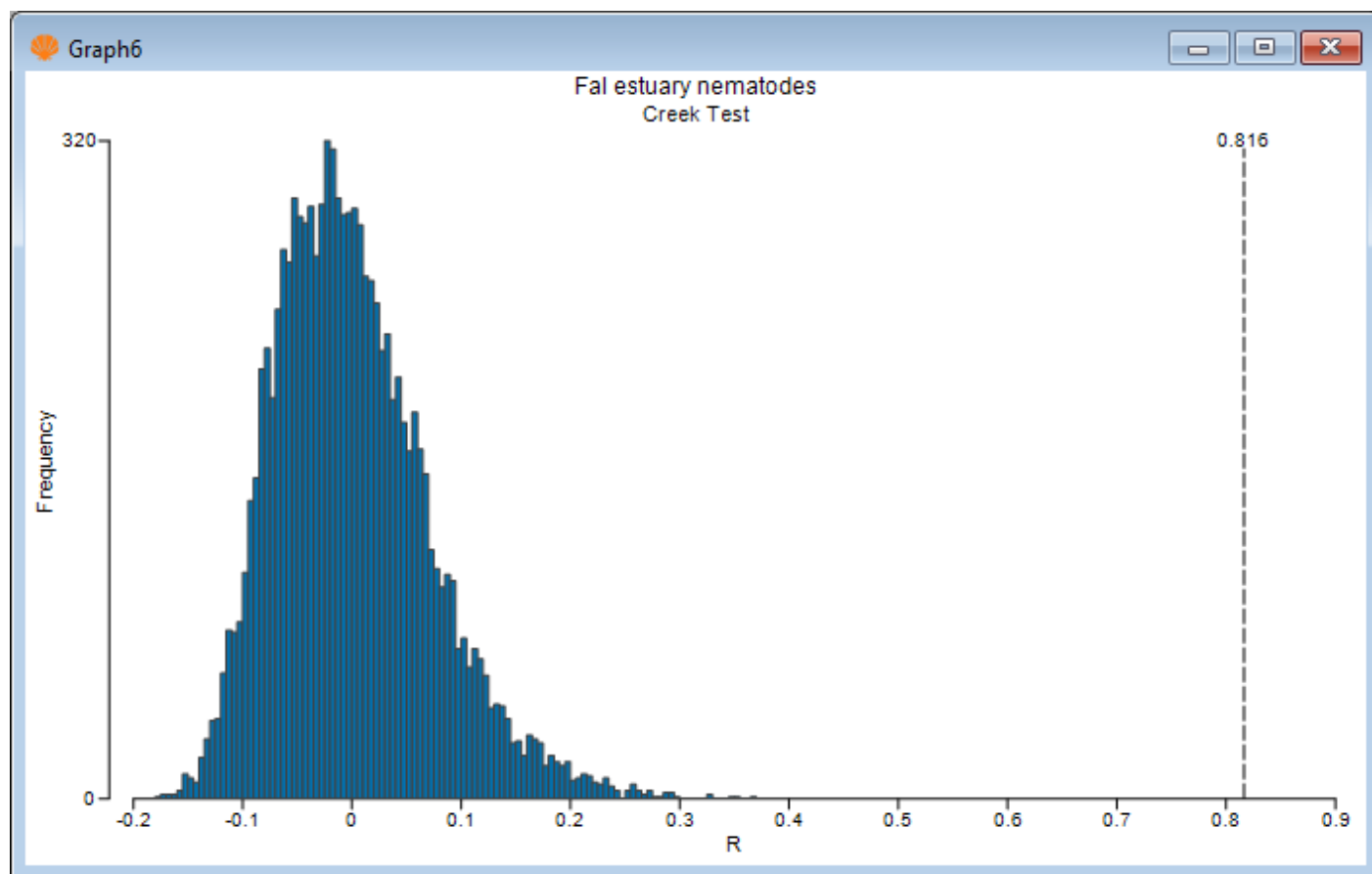
OK Cancel Help

The results, including the overall ANOSIM test for differences among all groups, as well as the ANOSIM tests of all pair-wise comparisons, are provided in the output file called 'ANOSIM1'. The overall test is clearly highly significant, with $R = 0.816$ and a significance level of 0.01%. This corresponds to a p-value of $P = 0.0001$ (the smallest possible value attainable for 9999 permutations).



The full set of pairwise tests is performed, and the strength of the differences between any pair of groups is well captured by the relative sizes of the *R* statistic - larger values indicate groups that are more different (more distinct or more easily distinguishable) from one another. The comparison of assemblages from St Just Creek (J) vs Percuil Creek (E) has the lowest *R* statistic ($R = 0.508$) of any tests done in this example, although this difference is still a statistically significant one ($P = 0.024$). Note that no 'correction' is being made here for multiple tests. The end-user is able to consider applying such a correction if desired (or may perhaps simply consider the frequency of 'significant' results obtained, given the number of tests performed), bearing in mind that each test is done using an exact permutation algorithm in order to generate the individual significance levels produced in the output, ignoring all other tests.

Also provided as output ('Graph6') is a histogram showing the distribution of values of the ANOSIM R statistic under permutation for the overall test. The observed value of the ANOSIM R statistic (0.816) is shown on the plot as a vertical dotted line. The purpose of the histogram is to show visually the departure of the observed value of R (or not) from the distribution of values of R expected under a true null hypothesis of 'no difference' among the groups.



3. Example analysis pathway (CLUSTER, MDS, ANOSIM)

Summary of the pathway

A summary of the essential routines in PRIMER 8 that were used to produce the 5-step analysis pathway described above for the nematode (biotic) data from the Fal estuary is given in the table below:

Step	To implement in PRIMER:
1. Fourth-root transformation	<i>From the original data sheet:</i> click Pre-treatment > Transform(overall) > (Transformation: Fourth root), click OK .
2. Bray-Curtis resemblance	<i>From the transformed data sheet:</i> click Analyse > Resemblance > (Measure $\bullet$ Bray-Curtis similarity) & (Analyse between $\bullet$ Samples), click OK .
3. Hierarchical group-average cluster analysis	<i>From the resemblance matrix:</i> click Analyse > Cluster > CLUSTER... > (Cluster mode $\bullet$ Group average) & ($\checkmark$ Plot dendrogram), click Finish .
4. Non-metric MDS ordination	<i>From the resemblance matrix:</i> click Analyse > MDS > Non-metric MDS (nMDS)... > (see if you are happy with the default settings), click OK .
5. Analysis of similarities (ANOSIM)	<i>From the resemblance matrix:</i> click Analyse > ANOSIM... > Design > (Model: One-Way - A) & Factors > A: Creek) & (Max permutations: 9999) & ($\checkmark$ Plot histogram), click OK .

4. Some wizards

4. Some wizards

Basic multivariate analysis

A useful analysis pathway (including the Example Analysis Pathway done above, with its five steps), can be accomplished in one fell swoop using the **Basic multivariate analysis** wizard. This will perform a suite of multivariate analyses commonly performed for either biotic or environmental data types, with options available that match the typical choices made for handling these different types of data.

Run the 'Basic multivariate analysis' wizard

Let's suppose you wanted to repeat the step-by-step analyses we did before, but this time using a different pre-treatment option. For example, you may decide you'd like to analyse the same data using presence/absence information only, so as to emphasise only the turnover in species identities (and not differences in abundance values) across sample units. It would be great to run this whole set of analyses quickly, rather than going through them again one at a time.

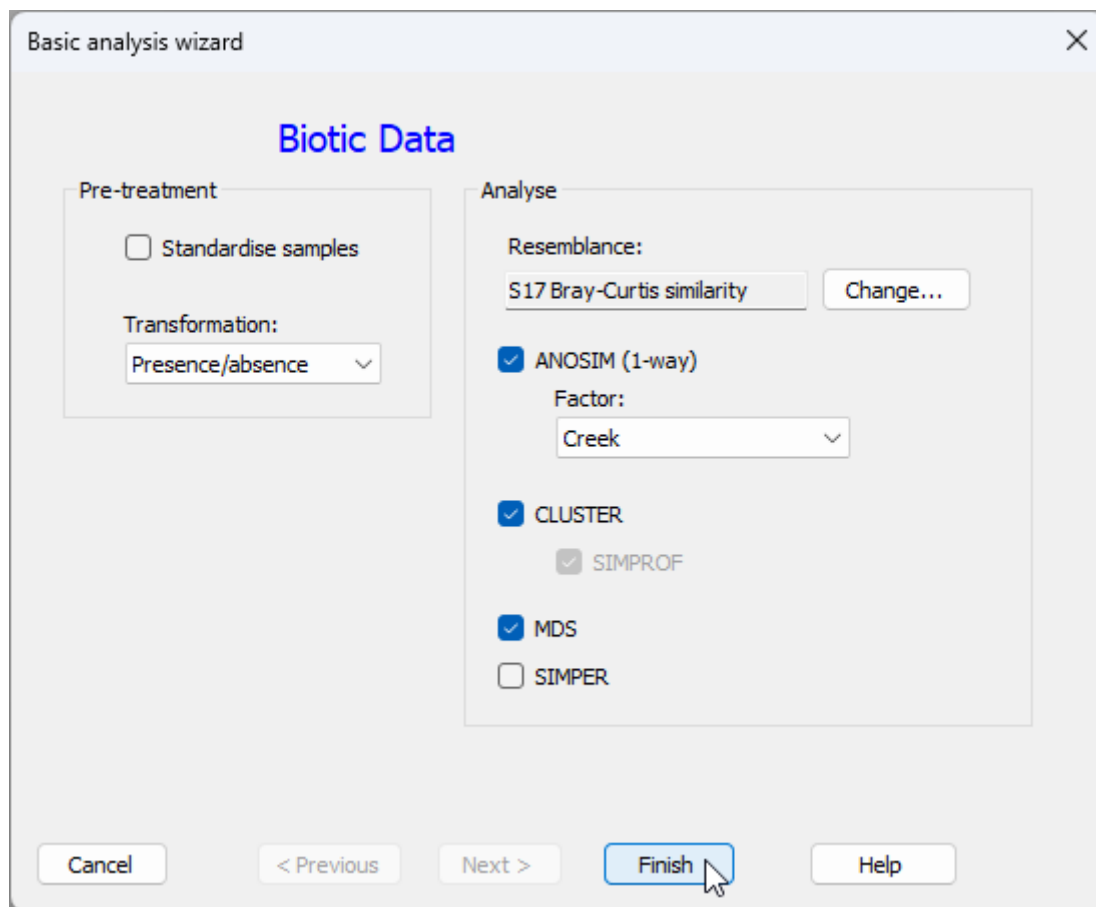
Click on the original 'Fal nematode abundance' datasheet in the Explorer tree and then click **Wizards > Basic multivariate analysis....**

The screenshot shows the PRIMER software interface. The 'Wizards' menu is open, highlighting 'Basic multivariate analysis...'. The Explorer tree on the left shows the 'Fal_Workspace' with the 'Fal_nematode_abundance' datasheet selected. The main window displays a preview of the 'Fal estuary nematodes Abundance' data table.

	R1	R2	R3	R4	R5	R6	R7	N
Anoplostoma vivip	0	0	0	0	0	0	0	0
Halalaimus gradis	0	0	0	0	0	0	0	0
Halalaimus longica	0	0	0	0	0	0	0	0
Oxystomina elonga	0	0	0	0	0	0	0	3
Viscosia viscosa	0	0	0	0	0	0	0	0
Tripyloides gradis	149	181	385	289	170	614	156	
Atrochromadora mi	0	0	0	0	0	0	0	0
Chromadora macro	0	4	29	63	263	123	67	
Chromadora nudica	0	0	0	0	0	7	19	
Chromadorella ?du	0	0	0	0	0	0	0	
Chromadorita nana	0	0	0	0	0	0	0	
Chromadorita tenta	0	0	0	0	0	0	0	
Dichromadora geop	5	0	0	0	0	7	0	
Hypodontolaimus b	40	44	25	18	5	14	3	
Ptycholaimellus po	174	424	178	99	107	123	35	

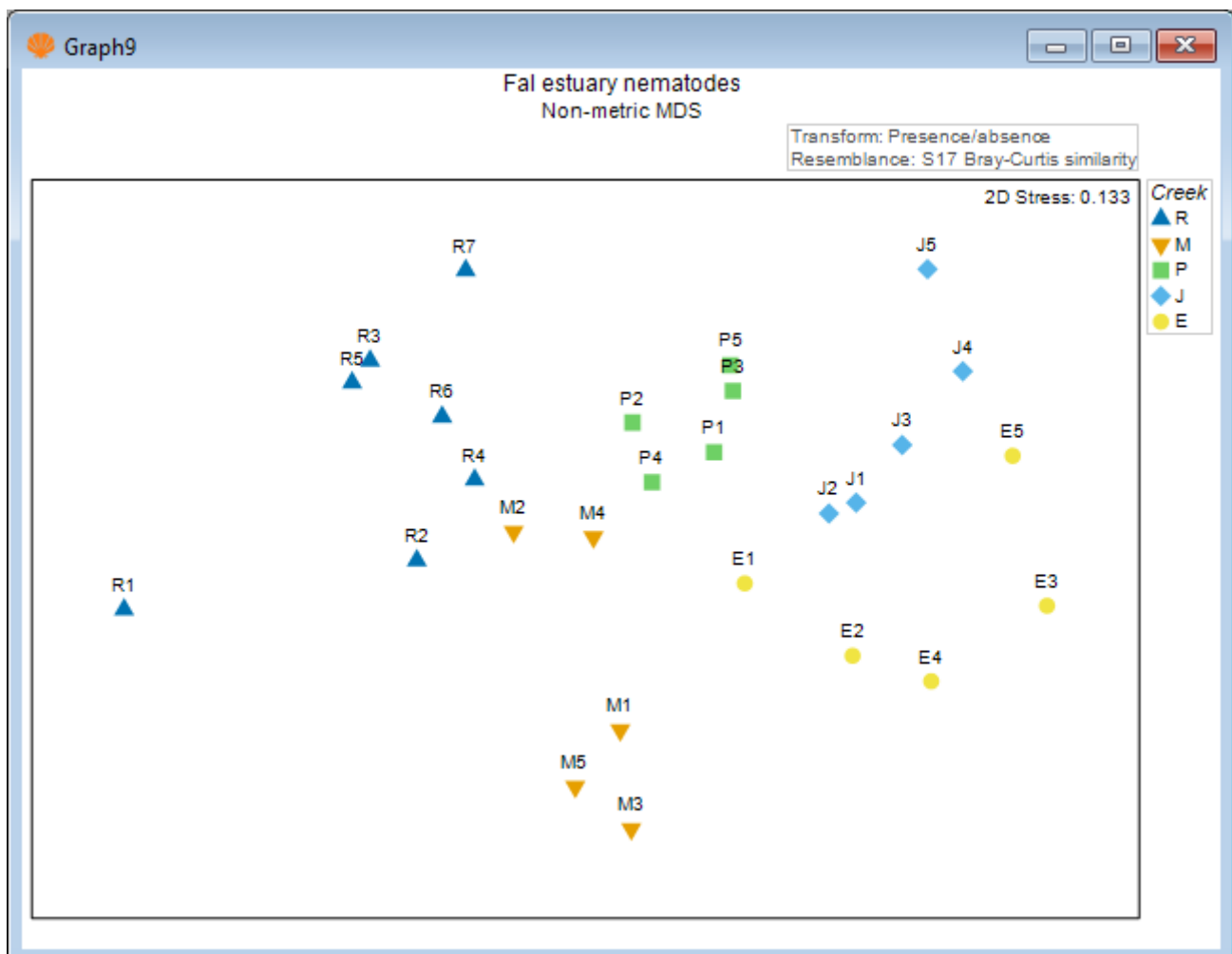
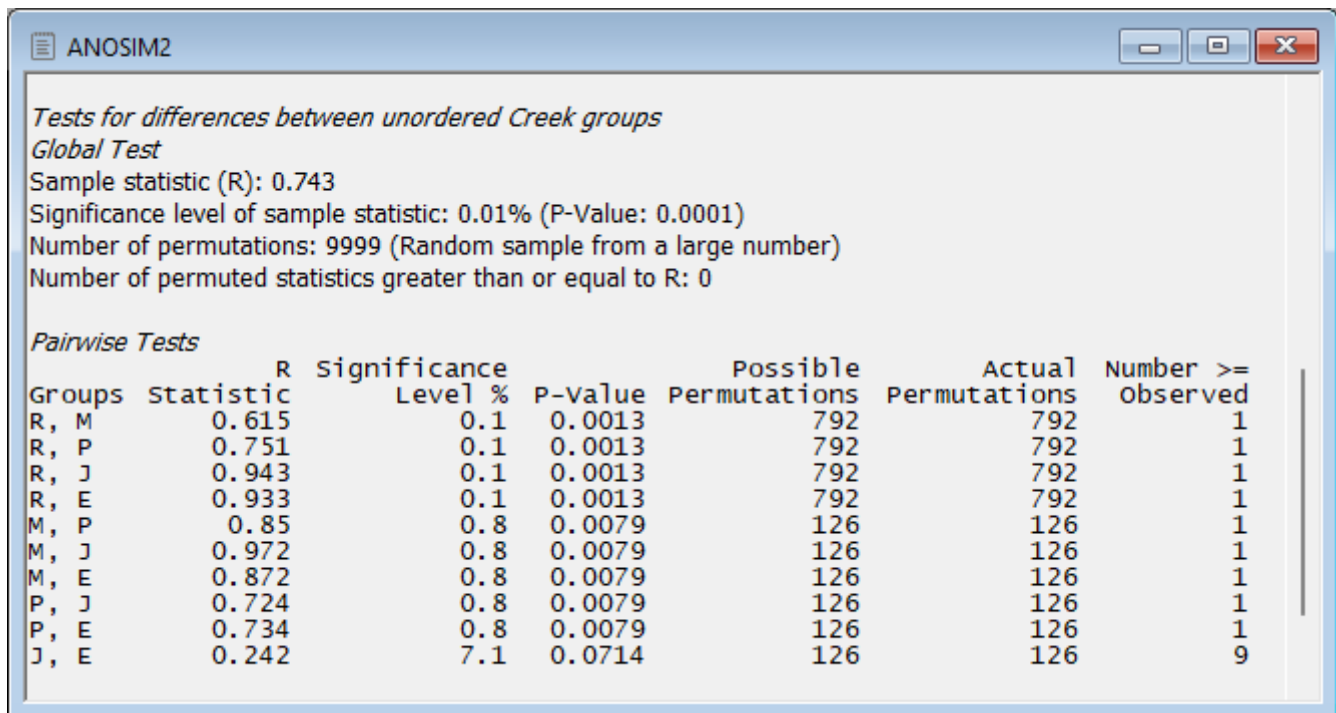
In the dialog box that follows, we can see that PRIMER is offering to perform a suite of basic multivariate analyses that are commonly performed for 'Biotic Data' (shown in bold blue font at the top). You can choose from a couple of (common) pre-treatment options, and then it will perform the analyses you choose (*via* the relevant checkboxes under 'Analyse'), using the default options for each routine. Note that different options would be shown for data of a different type (e.g., for environmental data). Recall that the 'type' of data is specified by you when you import the data into PRIMER. This can be changed for a given data sheet by clicking on **Edit > Properties** at any time.

Under 'Pre-treatment', choose 'Transformation: Presence/absence' and under 'Analyse', leave all of the default options, except you can *untick* the checkbox next to the 'SIMPER' routine (for now), then click **Finish**.



All of the requested analyses are done, and the resulting output files and graphics are shown in the Explorer tree window. Click on any of the items in the Explorer tree to see the steps that were taken, including specific data sheets, resemblance matrices, graphical outputs and results files. Although you will generally use PRIMER to perform individual analyses, one routine at a time, this wizard provides a quick way to achieve multiple analyses (provided you know *a priori* that you want to do them) with a single stroke.

The ANOSIM results ('ANOSIM2') and the 2D nMDS output (Graph9) obtained by this run of the wizard are shown below:



In this example, we can see that there are statistically significant differences in the identities of nematode species among the five creeks (ANOSIM $R = 0.743$, $P = 0.0001$ with 9999 permutations), although the pattern in the nMDS plot suggests that the difference between Restronguet Creek and the other creeks is not as great as what was observed when abundance information was included

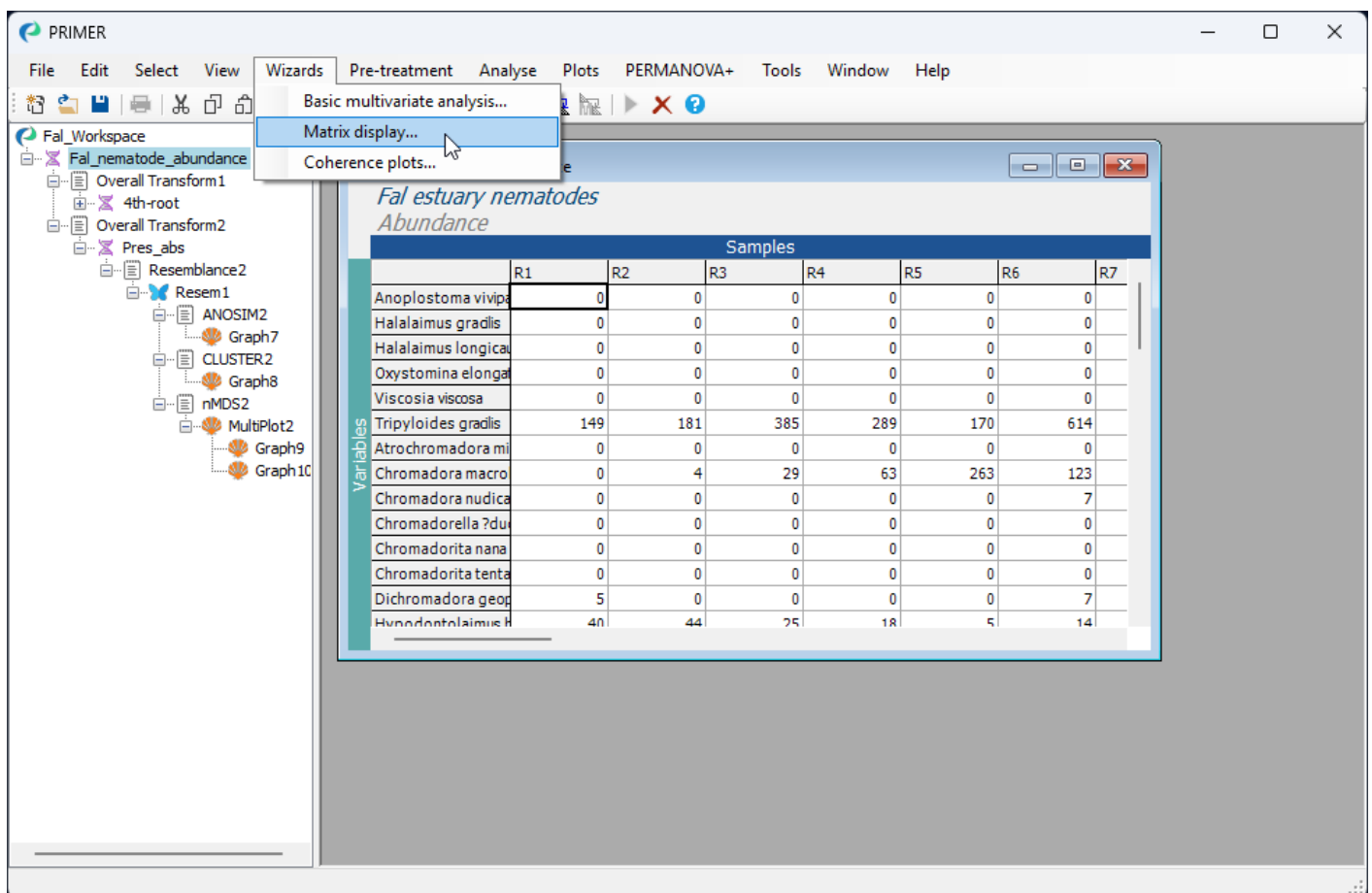
4. Some wizards

Matrix display

The **Matrix display** wizard produces a shade plot of a multivariate data matrix, with a useful ordering of its rows and columns that can help to clarify inter-sample and inter-species relationships, as well as gradients in turnover based on a resemblance matrix of choice.

Create a Matrix display

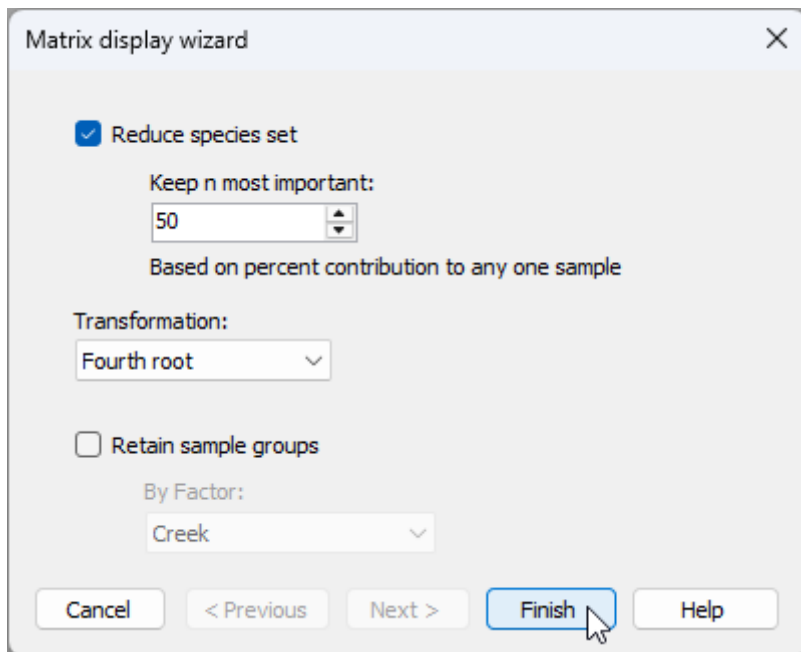
Click on the original data sheet ('Fal nematode abundance', at the top of the Explorer tree window), then click on **Wizards > Matrix display....**



The screenshot shows the PRIMER software interface. The 'Wizards' menu is open, and 'Matrix display...' is selected. The 'Matrix display...' dialog box is titled 'Fal estuary nematodes Abundance' and displays a table of abundance data for various nematode species across seven samples (R1 to R7). The table is titled 'Fal estuary nematodes Abundance' and has columns for 'Samples' (R1, R2, R3, R4, R5, R6, R7) and rows for 'Variables' (Anoplostoma vivip, Halalaimus gradis, Halalaimus longica, Oxystomina elonga, Viscosia viscosa, Tripyloides gradis, Atrochromadora mi, Chromadora macro, Chromadora nudica, Chromadorella ?du, Chromadorita nana, Chromadorita tenta, Dichromadora geop, Hymenodentolaimus H). The data values are as follows:

	R1	R2	R3	R4	R5	R6	R7
Anoplostoma vivip	0	0	0	0	0	0	0
Halalaimus gradis	0	0	0	0	0	0	0
Halalaimus longica	0	0	0	0	0	0	0
Oxystomina elonga	0	0	0	0	0	0	0
Viscosia viscosa	0	0	0	0	0	0	0
Tripyloides gradis	149	181	385	289	170	614	
Atrochromadora mi	0	0	0	0	0	0	0
Chromadora macro	0	4	29	63	263	123	
Chromadora nudica	0	0	0	0	0	0	7
Chromadorella ?du	0	0	0	0	0	0	0
Chromadorita nana	0	0	0	0	0	0	0
Chromadorita tenta	0	0	0	0	0	0	0
Dichromadora geop	5	0	0	0	0	0	7
Hymenodentolaimus H	40	44	25	18	5	14	

In the 'Matrix display wizard' dialog, leave the defaults, but choose (Transformation: **Fourth root**), then click **Finish**.



Matrix display wizard

☒ Reduce species set

Keep n most important:

50

Based on percent contribution to any one sample

Transformation:

Fourth root

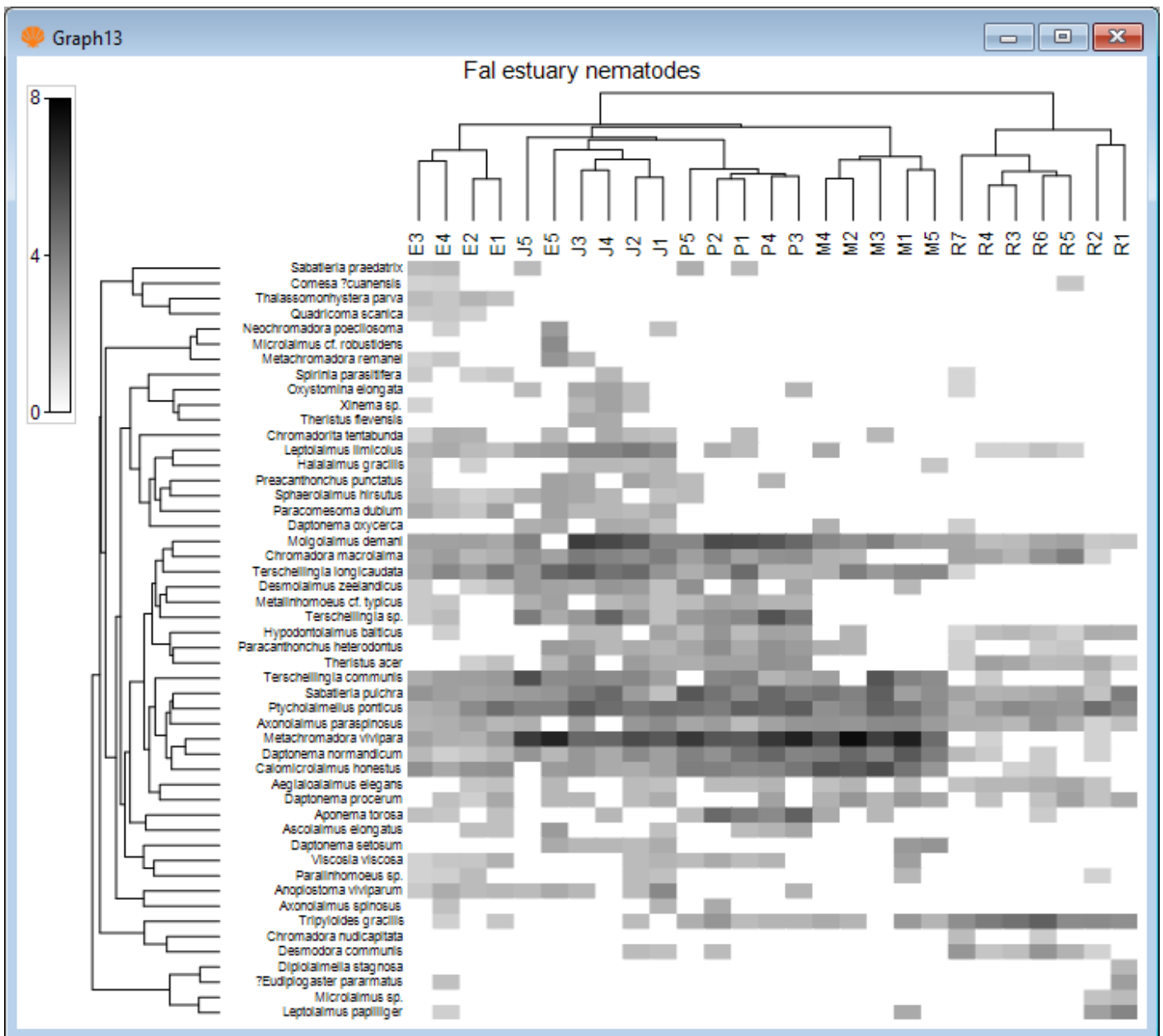
☐ Retain sample groups

By Factor:

Creek

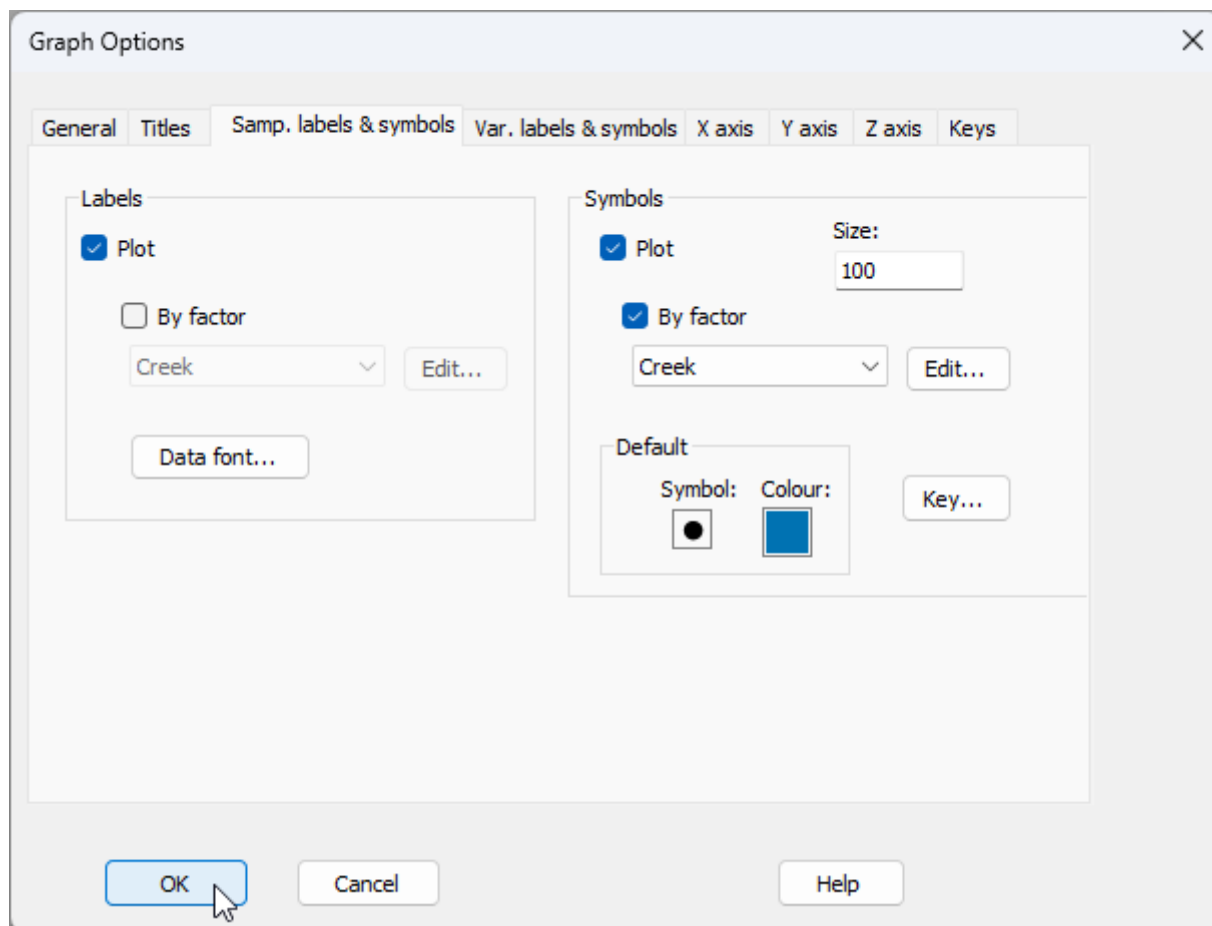
Cancel < Previous Next > Finish Help

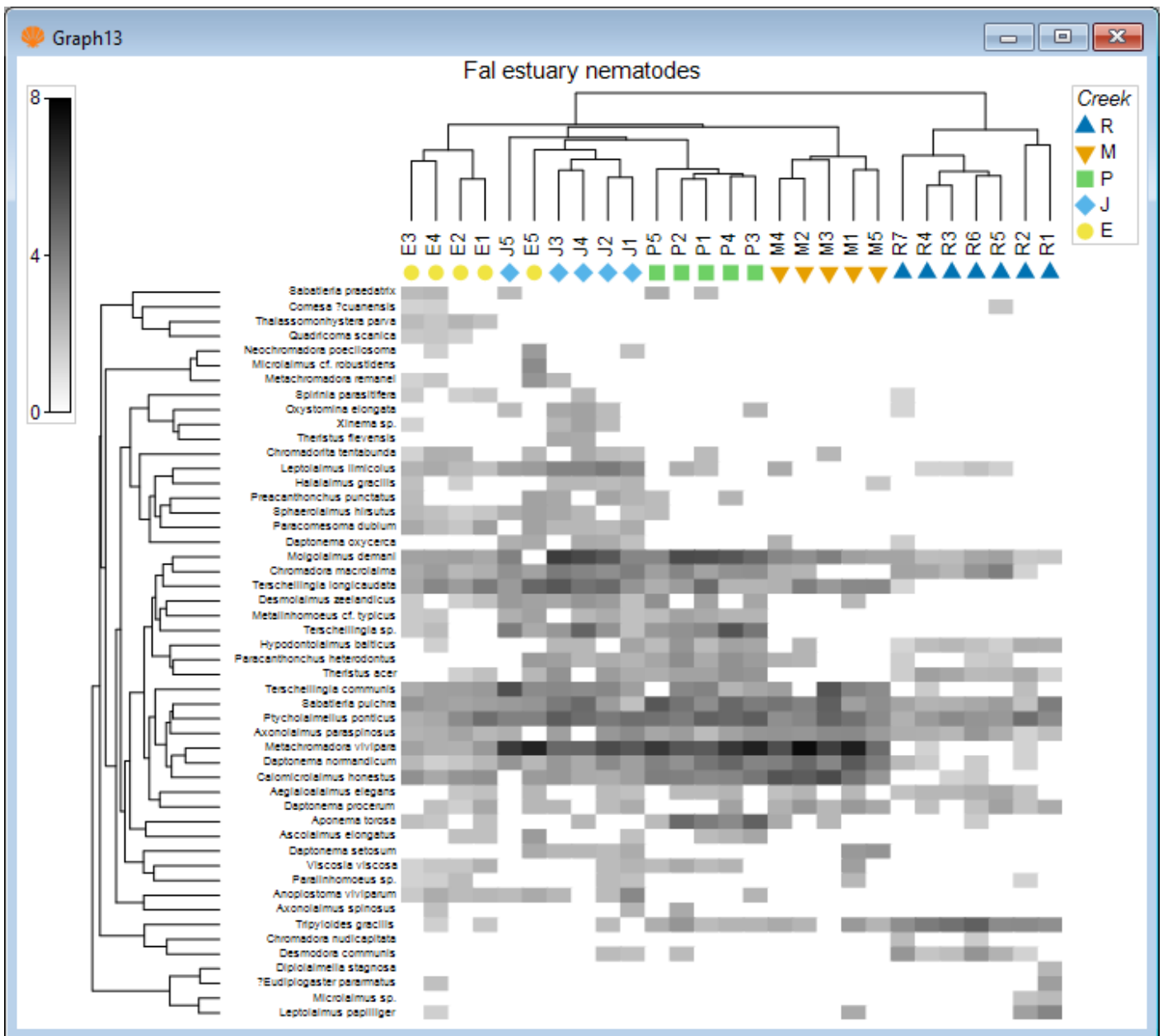
The default graphic ('Graph13') shows a shade plot of the 50 most important species (here, 'most important' is defined based on the percent(%) of the total abundance obtained by a species in *any* sample). Each value in the data matrix is represented by a shaded rectangle (from white to black, by default), where the degree of shading corresponds to the abundance values (after fourth-root transformation), as shown in the key (in this example, the transformed data values range from 0 to 8).



Here, samples are ordered so as to maximise their concordance with a model of 'seriation' (or turnover), defined on the basis of Bray-Curtis similarities calculated on fourth-root transformed data. They are (furthermore) constrained according to a dendrogram constructed from a hierarchical group-average cluster analysis of those similarities ('Graph12' in the Explorer tree). The species, in turn, are ordered according to a model of seriation based on an Index of Association calculated among species after first standardising the species variables by their total abundance. They are (similarly) further constrained by their own cluster dendrogram ('Graph11' in the Explorer tree).

It would perhaps be helpful to see the different creeks as different symbols on this plot. With Graph13 (the shade plot itself) as the active item in the Explorer tree, click **Graph > Sample Labels & Symbols**, then choose (Labels ☒ Plot) & (Symbols ☒ Plot > ☒ By factor Creek) and click **OK**.





We can see from the above image that some creeks contain a rather different set of species and/or different abundances of the same species. The ordering of whole creeks shown in the shade plot (obtained using the 'seriation' model) reflects the ordering we saw in the nMDS plot for these creeks as well (see the page 'Step 4. Ordination').

From a matrix display (or shade plot), clicking on **Graph > Special** reveals a very large number of colours and other graphical options and parameters allowing the user to alter and enhance this graphic for their purposes. These are too numerous to describe in detail here, but they include a host of methods for sorting the rows and/or columns (i.e., by clicking on the **Reorder...** button in the **Graph > Special** dialog).

5. Run a PERMANOVA

Overview

If you have purchased a subscription to *PRIMER 8 with PERMANOVA+*, then you will have access to the **PERMANOVA+** main menu item that allows you to perform a broad range of additional analyses using a suite of routines that are not available in the basic *PRIMER 8 Lite* package. Check out our website to see a [Detailed list of all features](#).

Note also that the **PERMANOVA** routine has been expanded substantially in the leap from PRIMER 7 to PRIMER 8. For a more complete guide to all of these important improvements, please see the essential resource: '[What's new in PRIMER 8](#)'. You can also consult the historical [PERMANOVA+ user manual](#) for fundamental information and references regarding some of these routines and their underlying statistical details.

Here, we will run through a quick example of how to set up and run a multi-factor PERMANOVA analysis in PRIMER 8. **Permutational multivariate analysis of variance** (PERMANOVA) partitions variation in the space of a chosen dissimilarity measure in response to one or more factors in a specified sampling protocol or experimental design ([Anderson \(2001a\)](#) , [Anderson \(2017\)](#)). Tests of individual terms in a PERMANOVA model are achieved by constructing correct (pseudo-)F ratios on the basis of expectations of mean squares (EMS), and p-values are obtained using correct permutation algorithms given the full study design. *We know of no other software package that accomplishes this.*

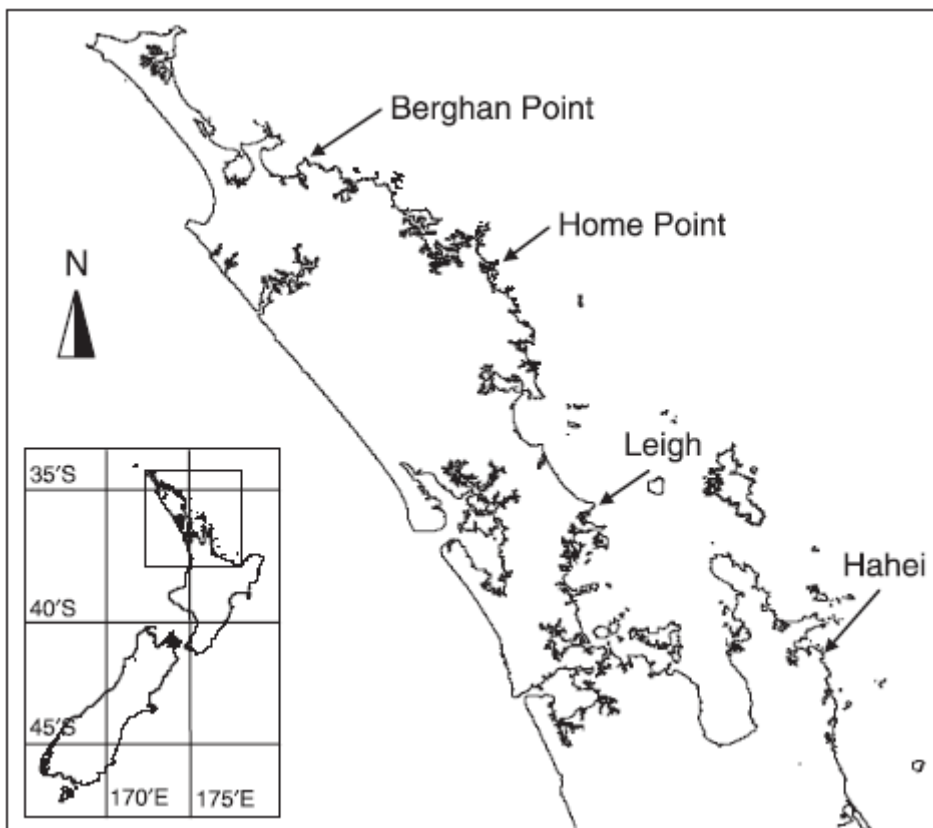
Importantly, the PERMANOVA routine in PRIMER allows the user:

- to specify whether factors are **fixed**, **random**, **finite**, or of a type called '**whole-plot/subject**' that caters to designs lacking replication at different scales,
- to account for **heterogeneity** in multivariate dispersions when testing for differences in centroids,
- to specify whether a factor is **nested** in one or more other factors,
- to test **interaction terms**,
- to include one or more quantitative **covariates** in the analysis,
- to **group covariates** (accommodating cyclical/seasonal models),
- to **re-order**, **pool** or **remove** individual terms from a model,
- to handle correctly:
 - **mixed models**
 - user-specified **contrasts**
 - **BACI designs** (before-after/control-impact),
 - **asymmetrical designs** (e.g., in environmental impact studies),
 - **randomised blocks**,
 - **split plots** and **split-split-plots**,
 - **hierarchical designs**,
 - **repeated measures**,

- **unbalanced designs** (Type I, II or III sums of squares),
- ... and more.

A three-factor hierarchical design

We will run PERMANOVA on an example dataset consisting of assemblages of molluscs collected from holdfasts of the kelp *Ecklonia radiata* in a 3-factor hierarchical experimental design. There were $n = 5$ holdfasts collected from each of 2 areas (tens of meters apart) at each of 2 sites (hundreds of meters to kilometers apart) from each of 4 locations (hundreds of kilometers apart) in rocky reef habitats along the northeastern coast of New Zealand ([Anderson et al. \(2005a\)](#) , [Anderson et al. \(2005b\)](#)).



The map above shows the four locations on the northeastern coast of New Zealand from which holdfasts were collected for the study: Berghan Point, Home Point, Leigh and Hahei (reproduced from Fig. 1 in [Anderson et al. \(2005a\)](#)).

The example data are found in the file called 'NE_NZ_holdfast_fauna.pri' in the folder named 'NE_NZ_holdfasts' inside the 'Examples_P8' folder that can be downloaded by clicking **Help > Get Examples....** There were 351 taxa (rows) from 15 different phyla quantified in this study. Here, we shall focus only on the phylum Mollusca (105 taxa).

Our interest lies in quantifying the degree of turnover in the identities of mollusc species at different spatial scales, as measured by the Jaccard resemblance measure. This a fully hierarchical sampling design with three spatial factors, as follows:

- Locations (random with 4 levels: Berghan Point, Home Point, Leigh and Hahei)

- Sites (random and nested in Locations, with 2 sites per location)
- Areas (random and nested in Sites, with 2 areas per site)

Areas are therefore also (necessarily) nested in Locations as well.

Steps in a PERMANOVA analysis

The two essential steps required to run a PERMANOVA analysis in PRIMER are always:

- *first*, **specify the design**; and
- *second*, **run the PERMANOVA analysis**, given the design, on a chosen resemblance matrix (arising from the data of interest).

Generally, we first need to get our data in to PRIMER, perform appropriate pre-treatment(s), if any, then calculate a resemblance matrix from this. A resemblance matrix will always serve as the starting point for any PERMANOVA analysis.

For this example...

An analysis of only the mollusc species from the holdfasts in accordance with the three-factor hierarchical design, based on Jaccard resemblances, can be done by following these steps:

1. **Open the example data file** in the PRIMER workspace and **select a subset of the variables** corresponding to only the mollusc species for what follows.
2. Calculate **Jaccard similarities*** among the sampling units (individual holdfasts).
3. **Specify the experimental/sampling design** (creating a design file).
4. **Run the PERMANOVA routine** (to obtain a partitioning and p-values *via* permutation for each term in the model).
5. Do an **ordination of distances among centroids** to visualise the effects and the relative importance of the factors.

*(*Note: the Jaccard resemblance measure utilises only presence/absence information, so we do not need to perform a pre-treatment transformation step for this example.)*

5. Run a PERMANOVA

Step 1: Data selection

Open up the example data file

Launch PRIMER, then click **File > Open...** from the main menu, navigate to the folder named 'NE_NZ_holdfasts' in the 'Examples_P8' directory, and select 'NE_NZ_holdfast_fauna.pri'. Click **Open** to display the species matrix.

PRIMER

File Edit Select View Wizards Pre-treatment Analyse Plots PERMANOVA+ Tools Window Help

Workspace

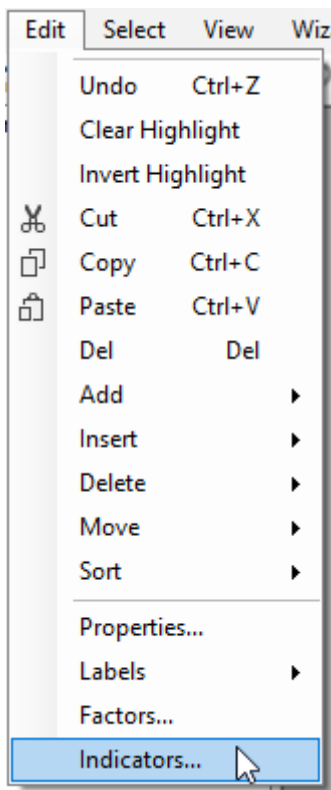
NE_NZ_holdfast_fauna

NE_NZ kelp holdfast fauna

Abundance

	Samples							
	Berghan11a	Berghan11b	Berghan11c	Berghan11d	Berghan11e	Berghan14a	Berghan14	
Acervulina inhaeren	8	4	4	2	0	0		
Trochulina dimidiata	0	0	0	0	0	0		
Sycon sp2	0	3	0	2	5	0		
Sycon sp1	0	0	0	2	0	2		
Sycon sp3	0	0	0	0	0	0		
Leucandra sp1	0	0	0	2	0	0		
Grantia sp2	0	0	0	0	0	0		
Amphiute sp1	0	0	0	0	0	0		
Leucosolenia sp1	0	0	1	0	0	1		
Leucetia sp1	0	0	0	0	0	0		
ClathrinidaeF sp1	1	1	0	2	2	0		
Clathrina sp1	0	0	0	0	0	0		
Sigmadocia flagellif	0	0	0	0	0	0		
Callyspongia sp1	0	0	0	0	0	0		

Click on **Edit > Indicators...**



You will see that there are indicators for the variables in this datasheet that show the taxonomic groups in which each species (or taxon) variable belongs, with different levels of the taxonomic hierarchy being provided as different indicators (i.e., 'Genus', 'Family', 'Order', 'Class' and 'Phylum'). There is also an indicator to identify whether individual taxa were counted (enumerated) or quantified on an ordinal scale ('Ordinal' = 'N' or 'Y', respectively).

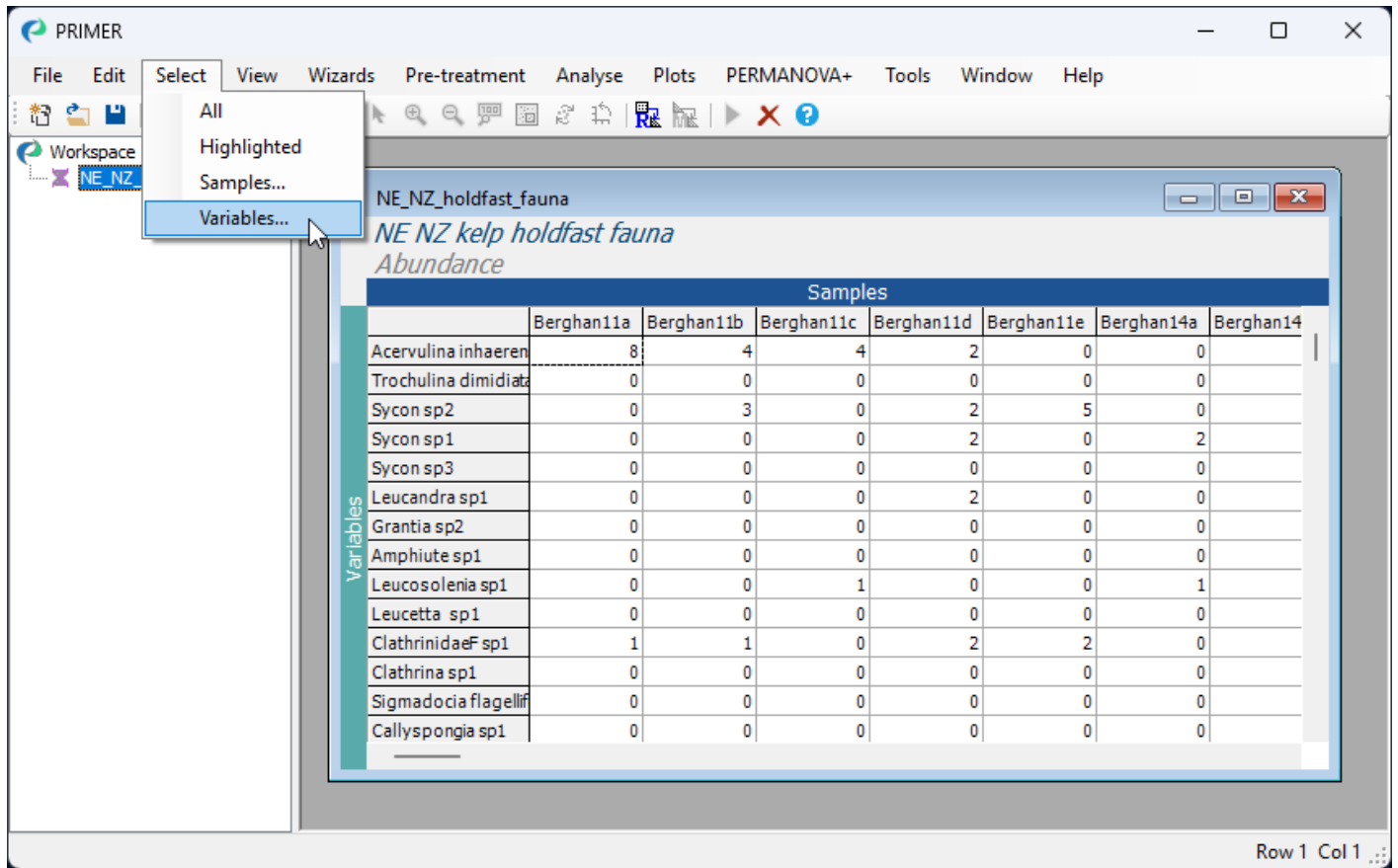
Indicators								
Edit Fill								
Add...	Label	Phylum	Class	Order	Family	Genus	Species	Ordinal
Duplicate	Acervulina inhaerens	Rhizopoda	Reticularia	Foraminiferida	Acervulinidae	Acervulina	Acervulina inha	N
Combine...	Trochulina dimidiata	Rhizopoda	Reticularia	Foraminiferida	Discorbidae	Trochulina	Trochulina dimi	N
Rename...	Sycon sp2	Porifera	Calcarea	Leucosolenida	Sycettidae	Sycon	Sycon sp2	N
Rename Levels...	Sycon sp1	Porifera	Calcarea	Leucosolenida	Sycettidae	Sycon	Sycon sp1	N
Reorder...	Sycon sp3	Porifera	Calcarea	Leucosolenida	Sycettidae	Sycon	Sycon sp3	N
Delete...	Leucandra sp1	Porifera	Calcarea	Leucosolenida	Grantiidae	Leucandra	Leucandra sp1	N
Key...	Grantia sp2	Porifera	Calcarea	Leucosolenida	Grantiidae	Grantia	Grantia sp2	N
Import...	Amphiute sp1	Porifera	Calcarea	Leucosolenida	Grantiidae	Amphiute	Amphiute sp1	N
OK	Leucosolenia sp1	Porifera	Calcarea	Leucosolenida	Leucosoleniidae	Leucosolenia	Leucosolenia sp	N
Cancel	Leucetia sp1	Porifera	Calcarea	Clathrinida	Clathrinidae	Leucetia	Leucetia sp1	Y
Help	ClathrinidaeF sp1	Porifera	Calcarea	Clathrinida	Clathrinidae	indet.	ClathrinidaeF sp	Y
	Clathrina sp1	Porifera	Calcarea	Clathrinida	Clathrinidae	Clathrina	Clathrina sp1	Y
	Sigmatocia flagellifer	Porifera	Demospongiae	Haplosclerida	Chalinidae	Sigmatocia	Sigmatocia fla	Y
	Callyspongia sp1	Porifera	Demospongiae	Haplosclerida	Callyspongiidae	Callyspongia	Callyspongia sp	Y
	Dysidea sp1	Porifera	Demospongiae	Dictyoceratida	Dysideidae	Dysidea	Dysidea sp1	Y
	HalichondridaF sp1	Porifera	Demospongiae	Dictyoceratida	Halichondrida	indet.	HalichondridaF	Y
	PhoriospongiidaeF1	Porifera	Demospongiae	Poecilosclerida	Phoriospongiidae	indet.	Phoriospongiid	Y
	PhoriospongiidaeF2	Porifera	Demospongiae	Poecilosclerida	Phoriospongiidae	indet.	Phoriospongiid	Y
	PhoriospongiidaeF3	Porifera	Demospongiae	Poecilosclerida	Phoriospongiidae	indet.	Phoriospongiid	Y
	Crella incrustans	Porifera	Demospongiae	Poecilosclerida	Crellidae	Crella	Crella incrustan	Y
	Penares sp1	Porifera	Demospongiae	Astrophorida	Ancorinidae	Penares	Penares sp1	Y
	DemospongiaeC sp1	Porifera	Demospongiae	indet.	indet.	indet.	DemospongiaeC	Y
	Sertularia simplex	Cnidaria	Hydrozoa	Conica (subClas	Sertulariidae	Sertularia	Sertularia sim	Y
	Sertularia sp	Cnidaria	Hydrozoa	Conica (subClas	Sertulariidae	Sertularia	Sertularia sp	Y
	Plumularia setacea	Cnidaria	Hydrozoa	Conica (subClas	Plumulariidae	Plumularia	Plumularia seta	Y

Click **OK** on the 'Indicators' dialog, so the data matrix is the active item in the workspace again.

Select a subset of variables, using an indicator

We wish to select just the mollusc species for the analysis.

1. From the 'NE_NZ_holdfast_fauna.pri' data sheet, click **Select > Variables...**



The screenshot shows the PRIMER software interface. The 'Select' menu is open, and 'Variables...' is highlighted. The data sheet 'NE_NZ_holdfast_fauna' is displayed, showing the abundance of various mollusc species across seven samples. The table is titled 'NE NZ kelp holdfast fauna Abundance' and has columns for 'Samples' (Berghan11a, Berghan11b, Berghan11c, Berghan11d, Berghan11e, Berghan14a, Berghan14b) and rows for species names.

	Berghan11a	Berghan11b	Berghan11c	Berghan11d	Berghan11e	Berghan14a	Berghan14b
Acervulina inhaerens	8	4	4	2	0	0	
Trochulina dimidiata	0	0	0	0	0	0	
Sycon sp2	0	3	0	2	5	0	
Sycon sp1	0	0	0	2	0	2	
Sycon sp3	0	0	0	0	0	0	
Leucandra sp1	0	0	0	2	0	0	
Grantia sp2	0	0	0	0	0	0	
Amphiute sp1	0	0	0	0	0	0	
Leucosolenia sp1	0	0	1	0	0	1	
Leucetta sp1	0	0	0	0	0	0	
ClathrinidaeF sp1	1	1	0	2	2	0	
Clathrina sp1	0	0	0	0	0	0	
Sigmatocia flagellifera	0	0	0	0	0	0	
Callyspongia sp1	0	0	0	0	0	0	

2. In the 'Select Variables' dialog, choose (• Indicator levels > **Phylum**) and click **Levels....**

Select Variables

☐ Variable names
Select Variable Names...

☐ Variable numbers
1

☒ Indicator levels
Phylum Levels...

☐ Exclude variables that:

☐ Have 10 % or more zero value(s)

☒ Have 1 or more missing value(s)

☐ Have 10 % or more missing values

☐ Top 1 variable(s) based on:

☒ Frequency of occurrence

☐ Total abundance (sum)

☐ Percent contribution to any one sample

☐ Select variables that:

☒ Occur in at least 3 sample(s)

☐ Occur in at least 5 % of sample(s)

☐ Contribute at least 10 % to total abundance (sum)

☐ Output selection to new worksheet

OK Cancel Help

3. In the 'Selection' dialog, click on the word 'Mollusca' in the list of 'Available' phylum categories, and click on the single-right-arrow button: >. This will move the word 'Mollusca' over to the 'Include:' box on the right-hand side of the 'Selection' dialog window. Click **OK**.

Selection

Select indicator groups

Available:

- Rhizopoda
- Porifera
- Cnidaria
- Platyhelminthes
- Nemertea
- Nematoda
- Sipuncula
- Mollusca**
- Annelida
- Arthropoda
- Brachiopoda
- Bryozoa
- Echinodermata
- Chordata
- Craniata

Filter:

Include:

Mollusca

Filter:

OK Cancel Help

→

Selection

Select indicator groups

Available:

- Rhizopoda
- Porifera
- Cnidaria
- Platyhelminthes
- Nemertea
- Nematoda
- Sipuncula
- Annelida
- Arthropoda
- Brachiopoda
- Bryozoa
- Echinodermata
- Chordata
- Craniata

Filter:

Include:

Mollusca

Filter:

OK Cancel Help

4. Now that we have selected the Molluscs, and we are back in the 'Select Variables' dialog window, we *also* want to choose to '☒ Output selection to new worksheet', as shown below. This will kick out just the variables that we have selected into a new worksheet for ensuing analysis. We can now click **OK** on the 'Select Variables' dialog.

Select Variables

☐ Variable names
Select Variable Names...

☐ Variable numbers
1

☒ Indicator levels
Phylum Levels...

☐ Exclude variables that:

☐ Have 10 % or more zero value(s)

☒ Have 1 or more missing value(s)

☐ Have 10 % or more missing values

☐ Top 1 variable(s) based on:

☒ Frequency of occurrence

☐ Total abundance (sum)

☐ Percent contribution to any one sample

☐ Select variables that:

☒ Occur in at least 3 sample(s)

☐ Occur in at least 5 % of sample(s)

☐ Contribute at least 10 % to total abundance (sum)

☒ Output selection to new worksheet

OK Cancel Help

Voila! Now you have just the mollusc species alone in a new worksheet called 'Data1'.

PRIMER

File Edit Select View Wizards Pre-treatment Analyse Plots PERMANOVA+ Tools Window Help

Workspace

- NE_NZ_holdfast_fauna
 - Select Variables1
 - Data1

NE_NZ_holdfast_fauna

Select Variables1

Data1

NE NZ kelp holdfast fauna

Abundance

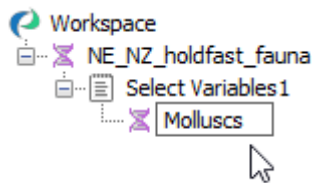
	Samples							
	Berghan11a	Berghan11b	Berghan11c	Berghan11d	Berghan11e	Berghan14a	Berghan14b	Berghan14c
Onithochiton neglectus	2	0	0	2	0	0	3	0
Rhyssoplax sp.	0	0	0	0	0	0	0	0
Notoplax violacea	2	2	1	1	1	0	1	0
Asteracmea suteri	1	0	0	1	1	0	0	0
Incisura rosea	1	1	1	1	2	0	0	0
Scissurella prendre	1	0	1	4	1	0	0	0
Haliotis sp.	0	0	0	0	0	0	1	0
Emarginula striatula	0	0	0	0	0	0	0	0
Tugali suteri	0	0	0	0	0	0	0	0
Cantharidus purpur	0	0	2	0	1	0	0	0
Herpetopoma bella	0	0	0	0	0	0	0	0
Herpetopoma larochei	0	0	0	0	0	1	0	0
Liotella polypleura	0	0	0	0	0	0	0	0

Row 1 Col 1

Rename the selected subset of data

From the subsetting data matrix ('Data1'), if you click on **Edit > Properties...**, you will see that this new sheet now only has 105 variables ('Number of rows:').

To change the name of the data sheet, click, hover and click again on the name 'Data1' in the Explorer tree window (or click **File > Rename Data** or hit the 'F2' key) and type in a new name for the subsetting data sheet: Molluscs.



At this point, you might like to save your workspace. Click **File > Save Workspace As... >** (Filename: NZ_holdfast_molluscs.pwk).

¶ An aside: Note that whenever you have selected a subset of data (this might be a subset of variables, as done here, or a subset of samples, or both), then the original data matrix on which the operation has been done will have a blue background colour to indicate that you have done this, like so:

NE_NZ_holdfast_fauna

NE NZ kelp holdfast fauna

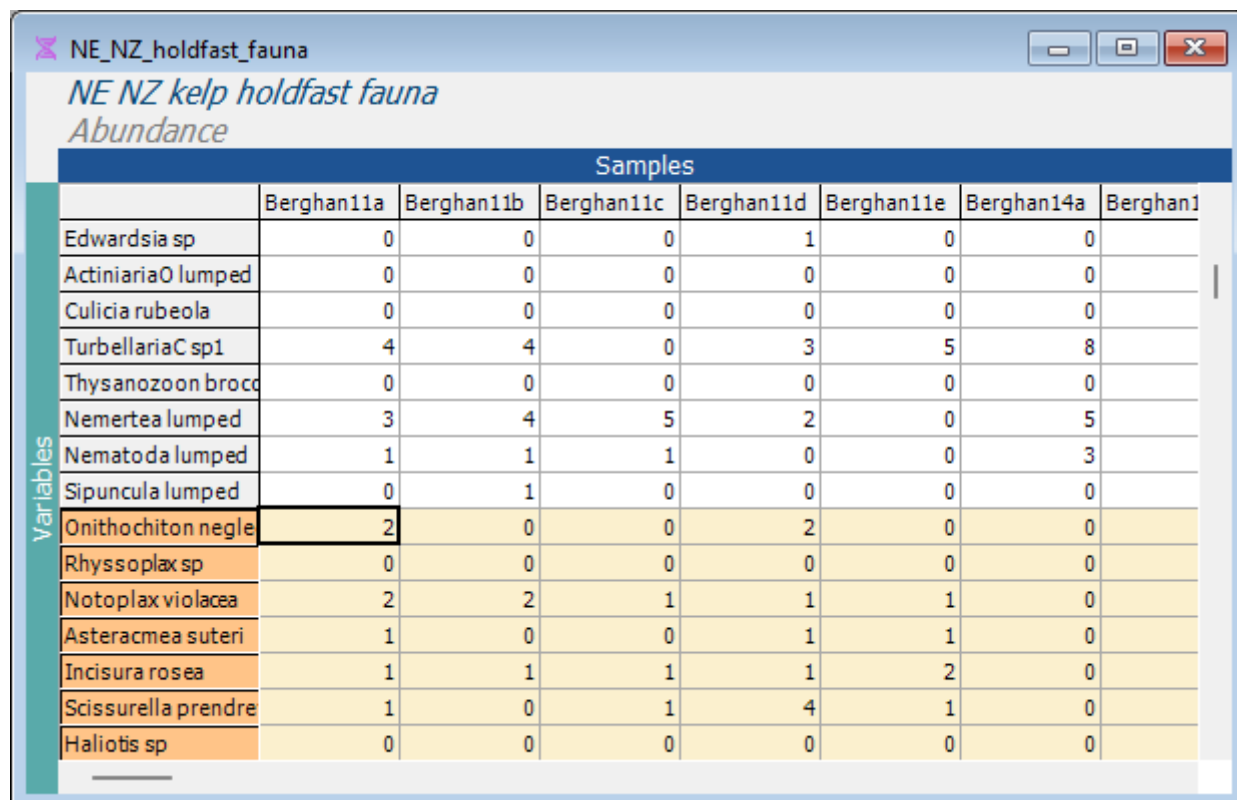
Abundance

Samples

	Berghan11a	Berghan11b	Berghan11c	Berghan11d	Berghan11e	Berghan14a	Berghan1
Onithochiton negle	2	0	0	2	0	0	
Rhyssoplax sp	0	0	0	0	0	0	
Notoplax violacea	2	2	1	1	1	0	
Asteracmea suteri	1	0	0	1	1	0	
Incisura rosea	1	1	1	1	2	0	
Scissurella prendre	1	0	1	4	1	0	
Haliotis sp	0	0	0	0	0	0	
Emarginula striatula	0	0	0	0	0	0	
Tugali suteri	0	0	0	0	0	0	
Cantharidus purpur	0	0	2	0	1	0	
Herpetopoma bella	0	0	0	0	0	0	
Herpetopoma laroc	0	0	0	0	0	1	
Liotella polypheura	0	0	0	0	0	0	
Trochus viridus	0	0	0	0	0	0	
Trochus sp	0	0	0	0	0	0	

Any analyses done on a selected subset of data will only be performed on that subset. It is usually a good idea to duplicate and rename a selected subset of data, so as to keep any analysis done on that subset of data clear and separate from the (full) original dataset. To duplicate an item in the Explorer tree, click **Tools > Duplicate**. To rename the duplicated item, hit F2 (or click **File > Rename Data**).

Note that subsetting does not affect the original full data matrix of information in any way, which is still always there. You can clear any subset selection (of variables and/or samples) by clicking on *Select > All* to return to the full data matrix in its entirety. The blue background colour will go away when you do that, and the formerly selected data will yet be highlighted in yellow-orange hues.



	Samples						
	Berghan11a	Berghan11b	Berghan11c	Berghan11d	Berghan11e	Berghan14a	Berghan1
Edwardsia sp	0	0	0	1	0	0	
ActiniariaO lumped	0	0	0	0	0	0	
Culicia rubeola	0	0	0	0	0	0	
TurbellariaC sp1	4	4	0	3	5	8	
Thysanozoon brocc	0	0	0	0	0	0	
Nemertea lumped	3	4	5	2	0	5	
Nematoda lumped	1	1	1	0	0	3	
Sipuncula lumped	0	1	0	0	0	0	
Onithochiton negle	2	0	0	2	0	0	
Rhyssoplax sp	0	0	0	0	0	0	
Notoplax violacea	2	2	1	1	1	0	
Asteracmea suteri	1	0	0	1	1	0	
Incisura rosea	1	1	1	1	2	0	
Scissurella prendre	1	0	1	4	1	0	
Haliotis sp	0	0	0	0	0	0	

Clicking on *Select > Highlighted* can then be used to re-instate the selection from before, if desired.

5. Run a PERMANOVA

Step 2: Jaccard resemblance

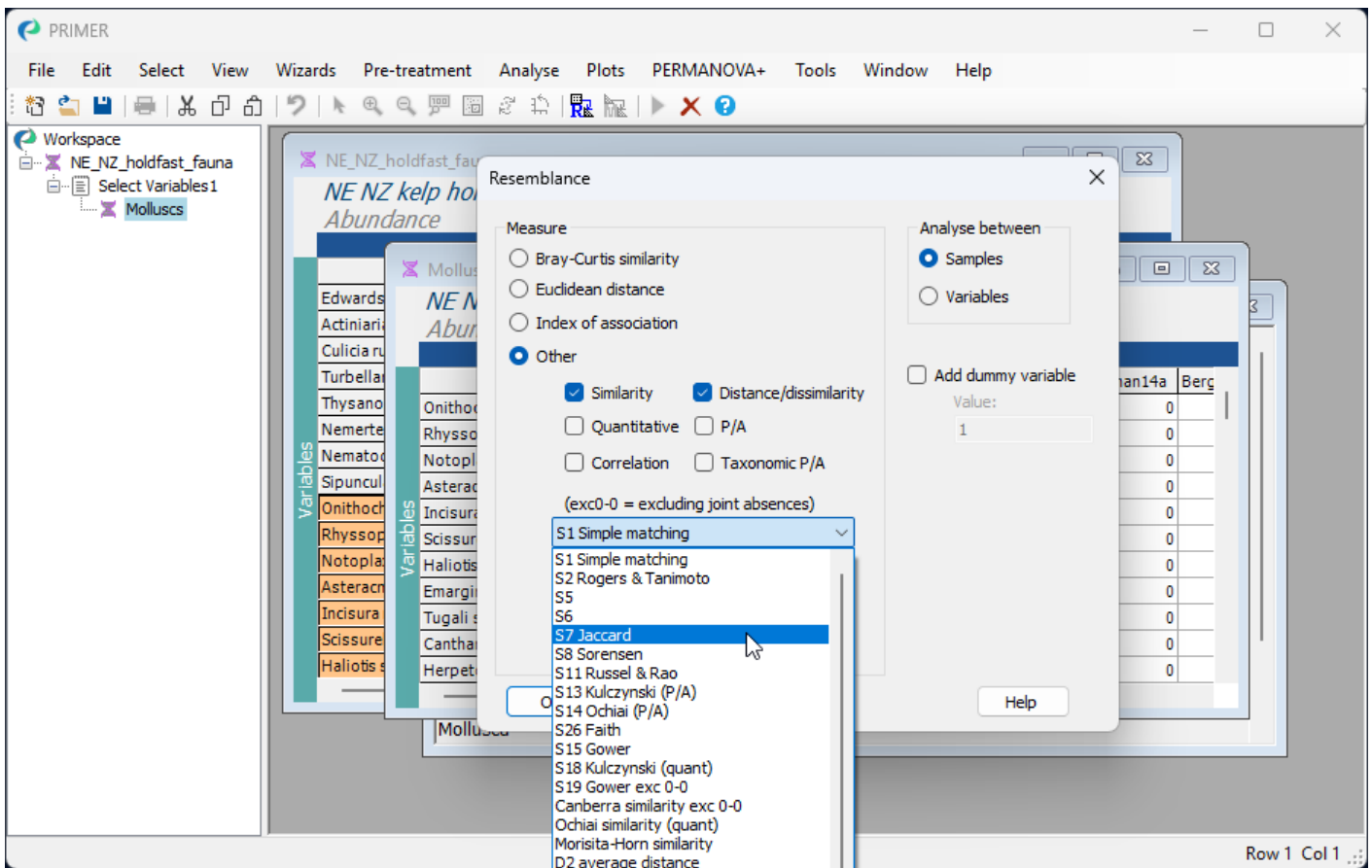
Calculate the Jaccard resemblance

From the 'Molluscs' data sheet, click **Analyse > Resemblance....**

The screenshot shows the PRIMER software interface. The 'Analyse' menu is open, and 'Resemblance...' is selected. The background displays a data table for 'Molluscs' abundance across various samples.

Variables	Berghan11a	Berghan11b	Berghan11c	Berghan11d	Berghan11e	Berghan14a	Berg
Onithochiton negle	2	0	0	2	0	0	
Rhyssoplax sp	0	0	0	0	0	0	
Notoplax violacea	2	2	1	1	1	0	
Asteracmea suteri	1	0	0	1	1	0	
Incisura rosea	1	1	1	1	2	0	
Scissurella prendre	1	0	1	4	1	0	
Haliotis sp	0	0	0	0	0	0	
Emarginula striatula	0	0	0	0	0	0	
Tugali suteri	0	0	0	0	0	0	
Cantharidus purpur	0	0	2	0	1	0	
Herpetopoma bella	0	0	0	0	0	0	

In the 'Resemblance' dialog, choose (\$\bullet\$Other) and then click on the drop-down menu to find 'S7 Jaccard', then click **OK**. (Note: 'S7' refers to the nomenclature used by [Legendre & Legendre \(2012\)](#) in their Chapter 7 on measures of ecological resemblance).



This produces a Jaccard similarity matrix among all pairs of holdfasts, called 'Resem1' by default. A nice feature of the Jaccard similarity measure is that it is directly interpretable as the percentage of shared species. For example, the first two holdfasts in the data matrix have a Jaccard similarity of 33.33%, so approximately one-third of the species that occur in either one or the other holdfast occur in both of them (i.e., are jointly present).

You can re-name this 'Jaccard' (click on the name 'Resem1' in the Explorer tree and hit the 'F2' key to re-name it), for clarity in what follows.

PRIMER

File Edit Select View Analyse PERMANOVA+ Tools Window Help

Workspace

- NE_NZ_holdfast_fauna
 - Select Variables1
 - Molluscs
 - Resemblance1
 - Jaccard

NE NZ holdfast fauna

Select Variables1

Molluscs

Resemblance1

Jaccard

NE NZ kelp holdfast fauna

Similarity (0 to 100)

	Samples						
	Berghan11a	Berghan11b	Berghan11c	Berghan11d	Berghan11e	Berghan14a	Berghan14b
Berghan11a							
Berghan11b	33.333						
Berghan11c	36.364	50					
Berghan11d	50	33.333	40				
Berghan11e	44.828	33.333	53.125	53.333			
Berghan14a	26.667	23.077	36.364	27.273	31.25		
Berghan14b	30.769	27.273	28.125	31.034	31.034	30.769	
Berghan14c	14.286	5.8824	10.714	12	12	14.286	
Berghan14d	30.769	21.739	32.258	35.714	26.667	25.926	
Berghan14e	31.707	23.077	38.636	41.463	41.463	28.571	
Berghan33a	20	14.286	27.586	21.429	21.429	11.111	

Row 1 Col 1

5. Run a PERMANOVA

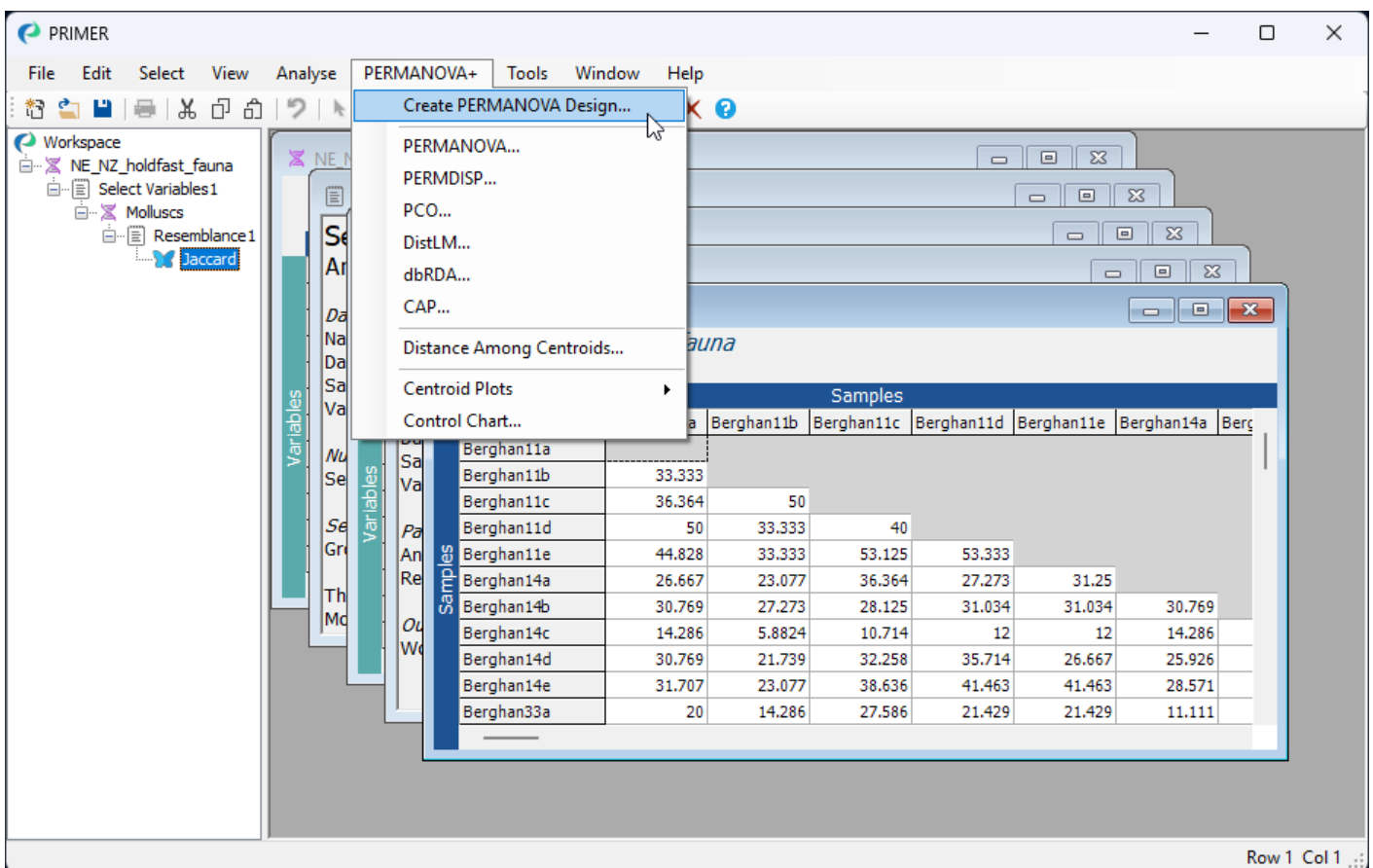
Step 3: Specify the design

PERMANOVA requires a **design file** to run.

You can see the Factors associated with the holdfast data matrix (or its resemblance matrix) by clicking on **Edit > Factors....** These factors will be 'visible' to the PERMANOVA dialog that we will use to create our design file.

For this study, we want to create a design file that has all of the information that PERMANOVA will need to construct the correct partitioning, the correct pseudo-F ratios and the correct permutation algorithms to test every term in the model that is implied by the design. For this example, we have a fully hierarchical (nested) study design with three random factors: Locations, Sites (within Locations) and Areas (within Sites).

1. From the resemblance matrix ('Jaccard'), click **PERMANOVA+ > Create PERMANOVA Design....**



2. You will see a design file, which always is indicated in the Explorer tree by the symbol , which you can use to specify all pertinent aspects of the sampling / experimental design[‡]. Initially, it will only be shown with a single row. There need to be as many rows as there are factors in your design. So, here, we will add two more rows by clicking on the 'Add row' button twice:

Design1

NE NZ kelp holdfast fauna

Add row Remove row

Factor	Nested in	Type	Contrasts
		Fixed	

Sources of variation

Terms... Pool...

☐ Use short names

Sums of Squares

☒ Type I (sequential)

☐ Type II (conditional)

☐ Type III (partial)

Covariables

Worksheet:

☐ Group covariables (indicator)

Make interaction terms available:

☒ Factor-by-covariable

☐ Covariable-by-covariable

Dispersions

☐ Allow for heterogeneity

Groups...

Term identifying groups with different dispersions:

This yields:

Design1

NE NZ kelp holdfast fauna

Add row Remove row

Factor	Nested in	Type	Contrasts
		Fixed	
		Fixed	
		Fixed	

Sources of variation

Terms... Pool...

☐ Use short names

Sums of Squares

☒ Type I (sequential)

☐ Type II (conditional)

☐ Type III (partial)

Covariables

Worksheet:

☐ Group covariables (indicator)

Make interaction terms available:

☒ Factor-by-covariable

☐ Covariable-by-covariable

Dispersions

☐ Allow for heterogeneity

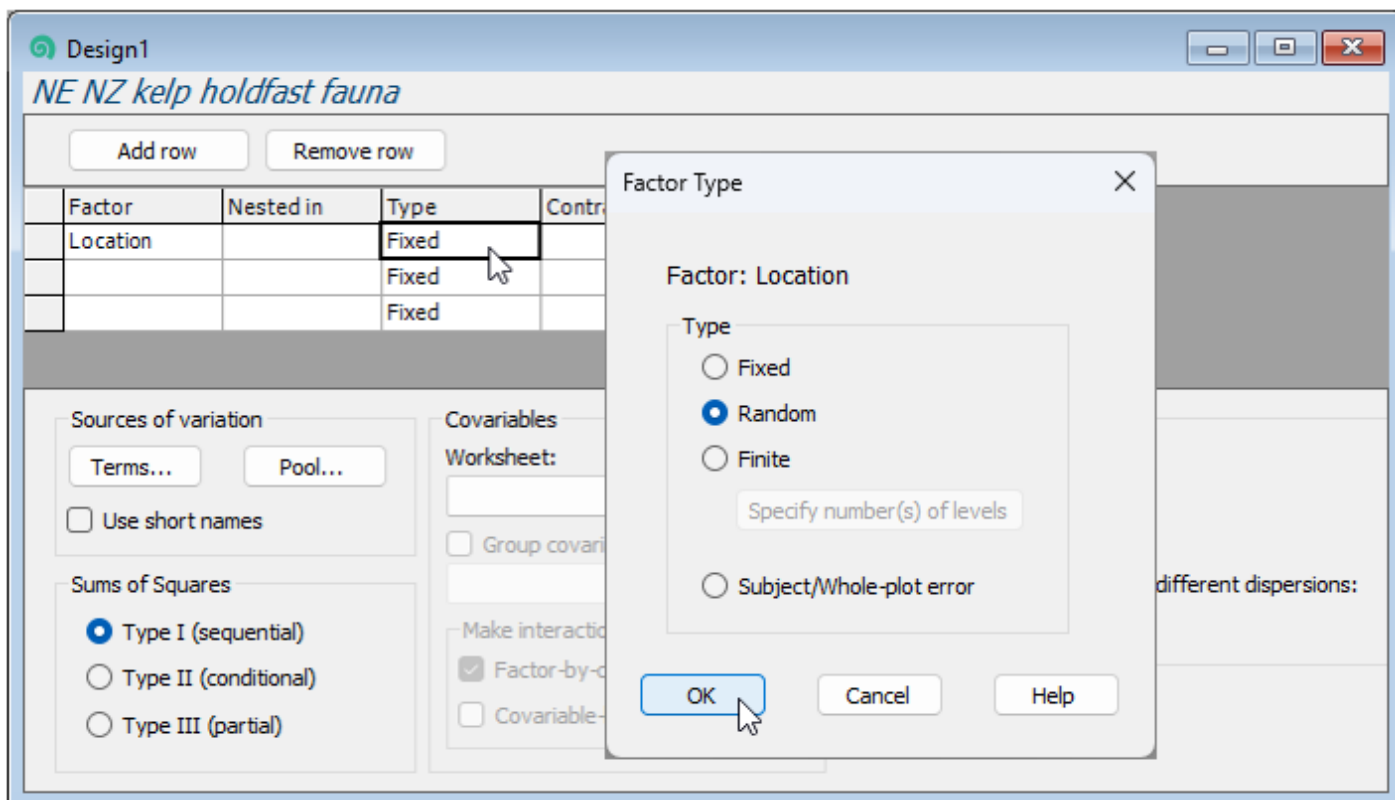
Groups...

Term identifying groups with different dispersions:

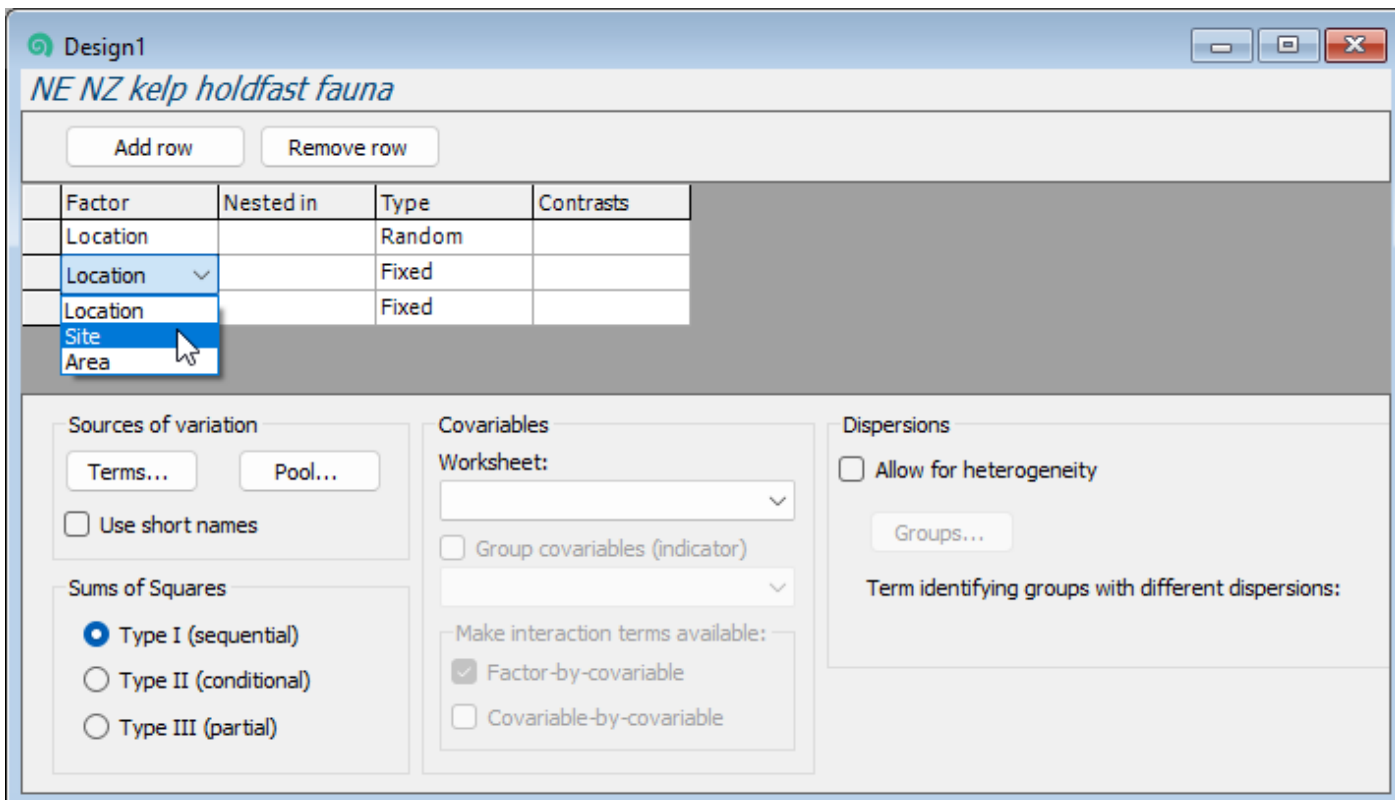
Next, you will need to specify, in turn, the name and properties of each factor for the analysis in its own row.

- Location:** First, in row 1, double-click in the blank cell under the word 'Factor', and you will see a drop-down menu listing all of the factors associated with the resemblance matrix from which this design file was created. Choose 'Location' to fill this cell. Location

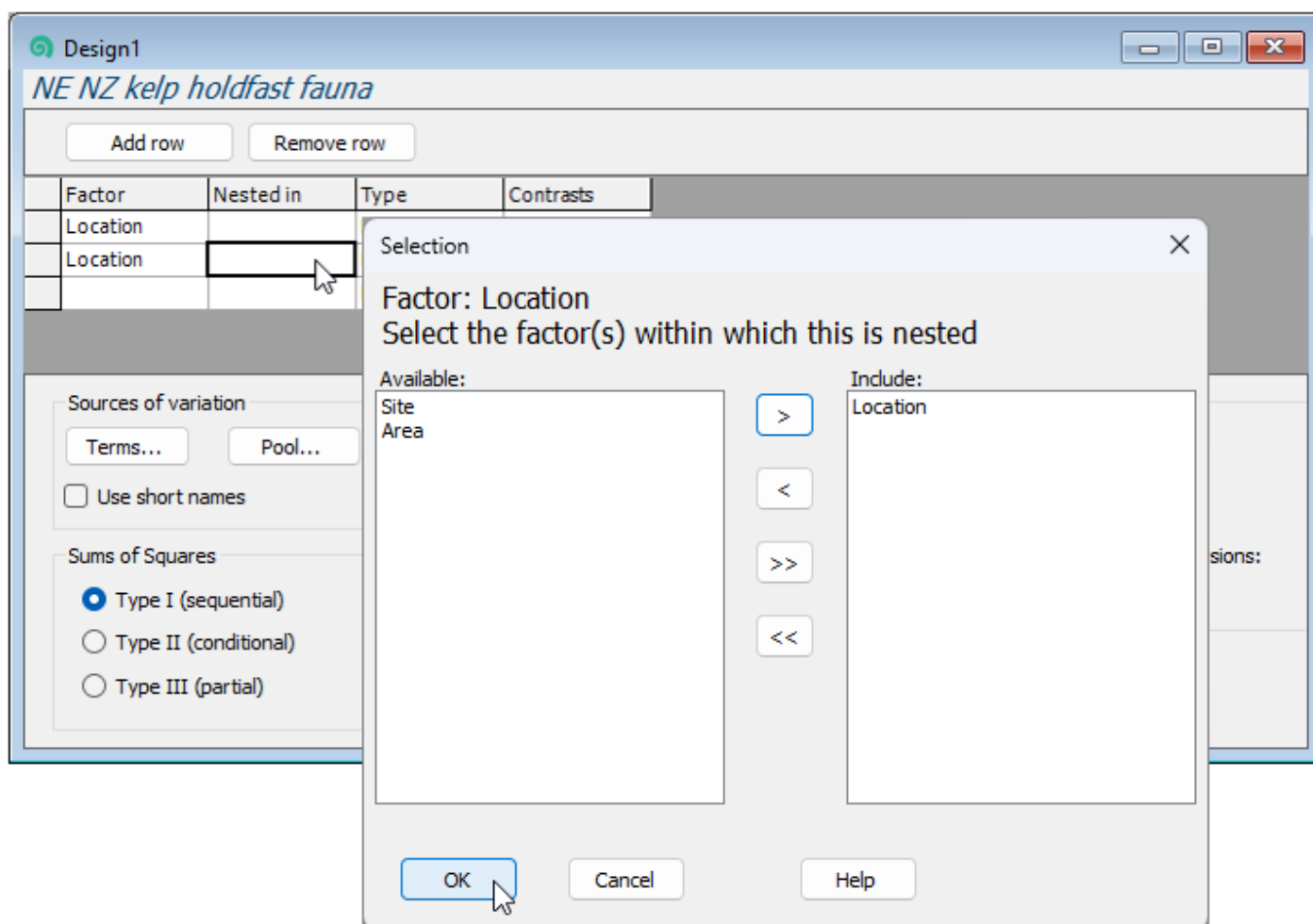
is not nested in anything, so we leave the second cell in row 1 blank. In the third cell of row 1, we have to specify that Location is a **random** factor, so double-click on the word 'Fixed' and in the 'Factor Type' dialog window under 'Type', click on $\bullet$ Random, then click **OK**.



5. **Site:** Next, we need to specify the second factor in the design in row 2 of the design file. Double-click on the cell in row 2 of the 'Factor' column (column 1) and choose 'Site'.



6. Sites are **nested in Locations**, so we have to specify that in column two ('Nested in'), accordingly. Double-click on the cell in row 2 in column 2 and a 'Selection' dialog will pop up to allow you to choose the factors within which 'Site' is nested. You will need to click on the word 'Location' (in the 'Available:' box on the left), then on the single-right-arrow button () to move it over into the 'Include:' box (on the right), then click **OK**, like so:



7. Make sure that the factor 'Site' is also specified in column 3 as 'Random'. (This happens automatically after specifying a nested term in column 2, because nested terms are, almost always, random factors.) Your design file should now look like this:

Design1

NE NZ kelp holdfast fauna

Add row Remove row

Factor	Nested in	Type	Contrasts
Location		Random	
Site	Location	Random	
		Fixed	

Sources of variation

Terms... Pool...

☐ Use short names

Sums of Squares

☒ Type I (sequential)

☐ Type II (conditional)

☐ Type III (partial)

Covariables

Worksheet:

☐ Group covariables (indicator)

Make interaction terms available:

☒ Factor-by-covariable

☐ Covariable-by-covariable

Dispersions

☐ Allow for heterogeneity

Groups...

Term identifying groups with different dispersions:

8. **Area:** Finally, we need to specify the third factor (Areas) correctly in row 3 of the design file. Under 'Factor' choose 'Area'. Under 'Nested in', specify 'Site', and make sure that column three has the word 'Random'. The final three-way nested design file should look like this:

Design1

NE NZ kelp holdfast fauna

Add row Remove row

Factor	Nested in	Type	Contrasts
Location		Random	
Site	Location	Random	
Area	Site	Random	

Sources of variation

Terms... Pool...

☐ Use short names

Sums of Squares

☒ Type I (sequential)

☐ Type II (conditional)

☐ Type III (partial)

Covariables

Worksheet:

☐ Group covariables (indicator)

Make interaction terms available:

☒ Factor-by-covariable

☐ Covariable-by-covariable

Dispersions

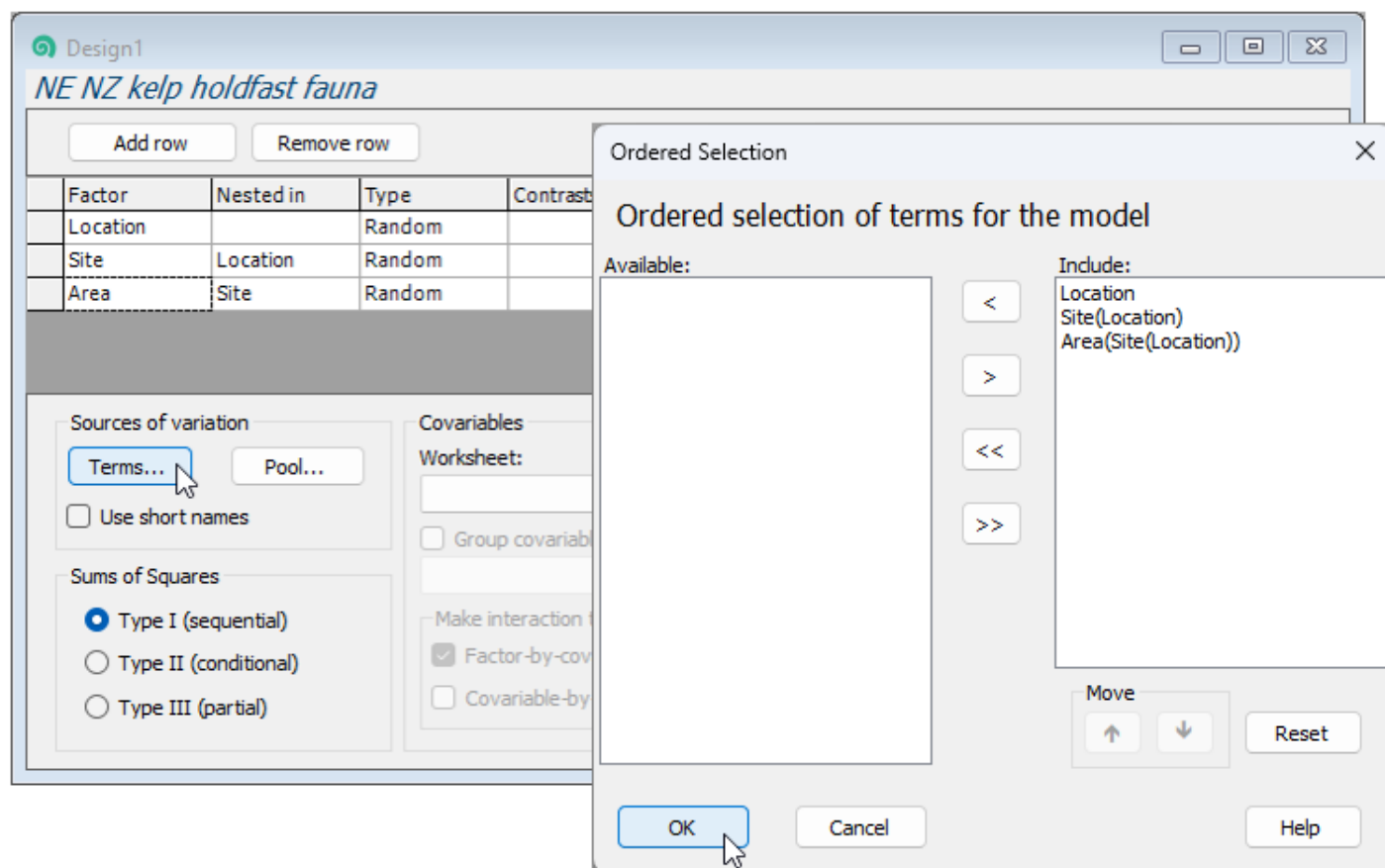
☐ Allow for heterogeneity

Groups...

Term identifying groups with different dispersions:

Now our design file has all of the relevant information to run the PERMANOVA analysis.

You can see all of the terms in your PERMANOVA model by clicking on the 'Terms' button, as shown below:



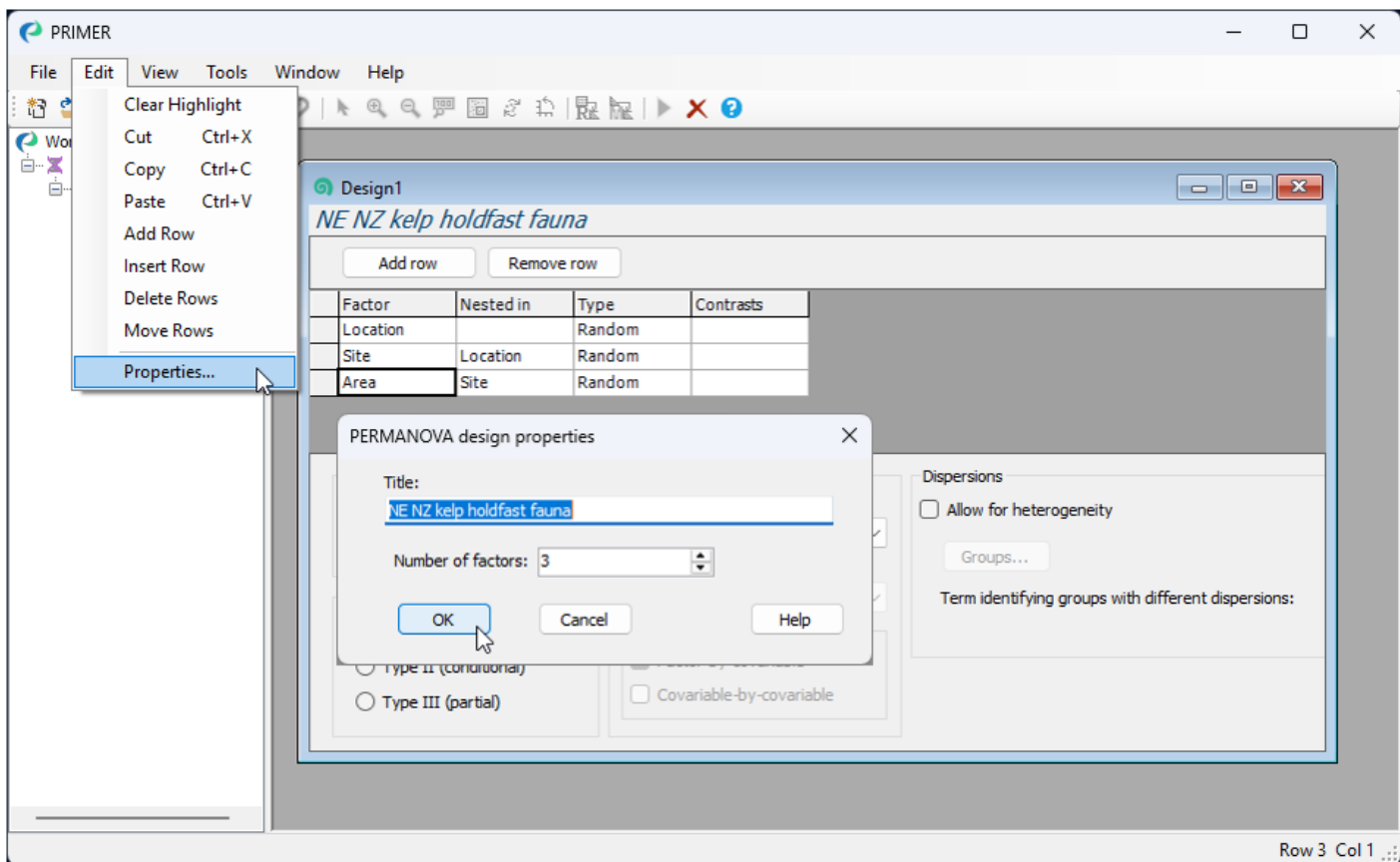
(Nb. There are no interaction terms in this particular fully nested hierarchical design, but if interactions are implied by your design e.g., in crossed designs, then all such interaction terms, by default, will be included in the model.)

You will note that there are other options here, including the ability to:

- change the model by removing or pooling terms, or change their order,
- include tests of specific contrasts,
- choose the 'Type' of sums of squares (e.g. for unbalanced cases),
- add quantitative covariates; and/or
- accommodate heterogeneity of dispersions, etc.

Here, we shall leave all of these as their defaults. We are now ready to run the PERMANOVA analysis itself according to these design specifications.

‡ Note that a Design-file item in PRIMER 8 has its own properties, which you can see by clicking on Edit > Properties. Doing this shows a title for the design (which can be modified if desired) and you can also specify the number of factors (rows) in this way (rather than using the 'Add row' and 'Remove row' buttons):

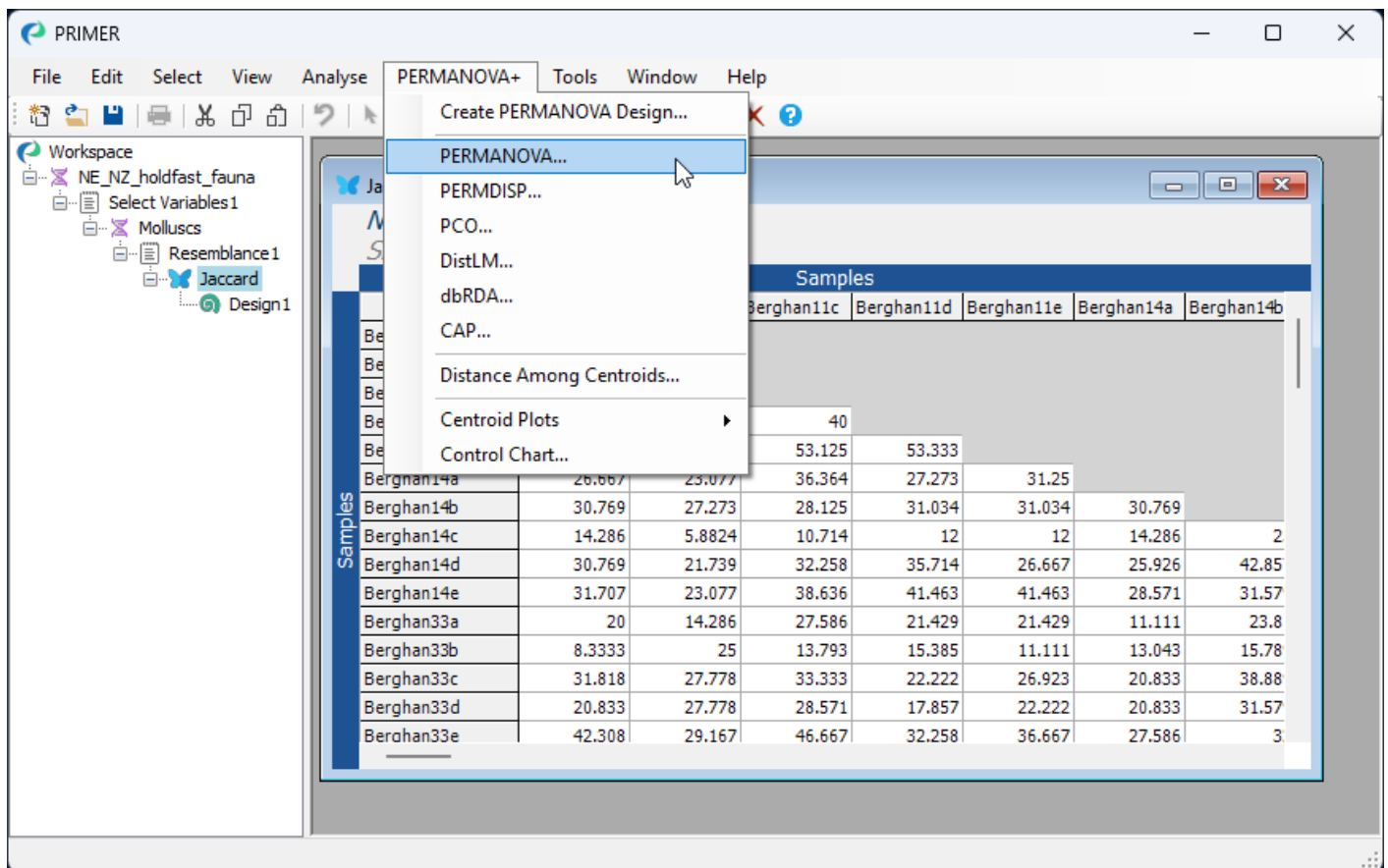


5. Run a PERMANOVA

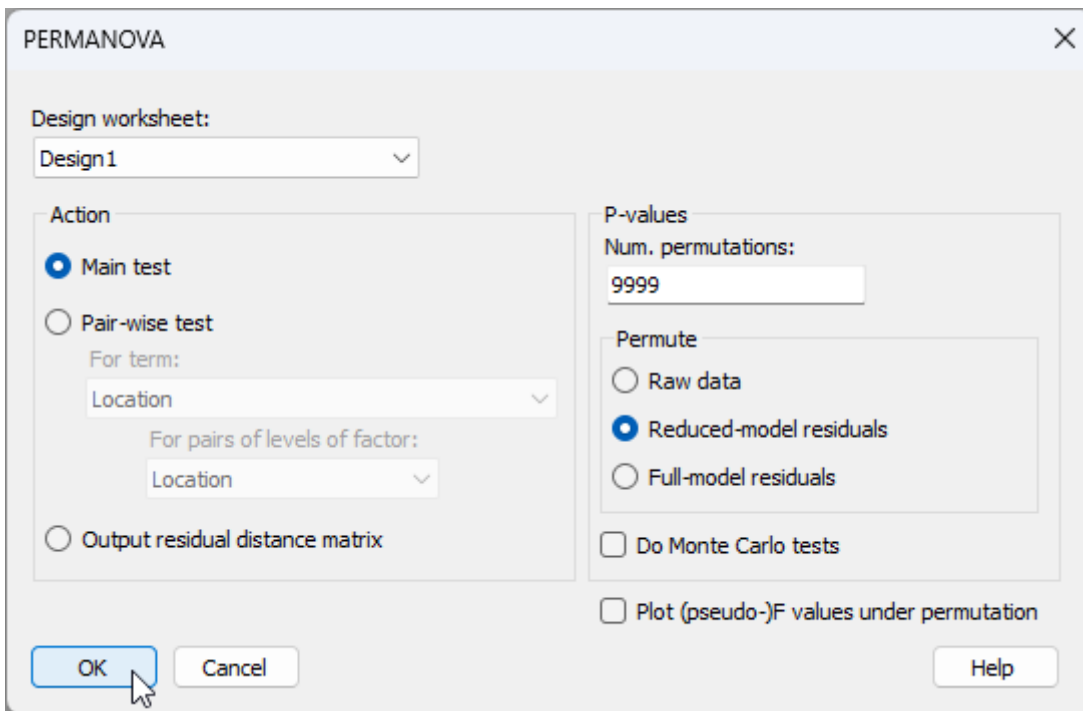
Step 4: Run PERMANOVA

Once the design file is created, we are ready to go ahead with the PERMANOVA analysis.

1. Click on the 'Jaccard' resemblance matrix in the Explorer tree so that it is the active item in the workspace, then click **PERMANOVA+ > PERMANOVA...**



2. Check to see that the 'Design worksheet:' is **Design1**; this is the design file we created in the previous step that contains the three-way nested design. For the rest, we will keep the defaults in the PERMANOVA dialog, as shown below, then click **OK**.



3. This produces an output file (called 'PERMANOVA1') in the Explorer tree. It shows:

- the details of the choices you made in the PERMANOVA dialog to run the analysis;
- the details of your experimental design; and
- the PERMANOVA table of results

as follows:

PERMANOVA1

PERMANOVA

Permutational MANOVA

Resemblance worksheet

Name: Jaccard

Data type: Similarity

Selection: All

Resemblance: S7 Jaccard

PERMANOVA design worksheet

Name: Design1

Title: NE NZ kelp holdfast fauna

Sums of squares type: Type I (sequential)

Permutation method: Permutation of residuals under a reduced model

Number of permutations: 9999

Factors

Name	Type	Levels
Location	Random	4
Site	Random	8
Area	Random	16

Excluded terms

No terms excluded

Inestimable terms

No inestimable terms.

PERMANOVA table of results

Source	df	SS	MS	Pseudo-F	P(perm)	Unique perms
Location	3	35564	11855	2.8086	0.0106	105
Site(Location)	4	16883	4220.8	1.3564	0.0331	9859
Area(Site(Location))	8	24895	3111.8	1.232	0.0077	9709
Res	64	1.6165E+05	2525.7			
Total	79	2.3899E+05				

Interpretation

These results show that there is statistically significant variability in the identities of molluscs among holdfasts at each of the three spatial scales in the experimental design: **Areas** ($F_{8,64} = 1.23$, $P < 0.01$), **Sites** ($F_{4,8} = 1.36$, $P < 0.05$) and **Locations** ($F_{3,4} = 2.81$, $P < 0.05$). Note that the p-value for Locations is somewhat limited by the number of unique values of the pseudo-F statistic under permutation that are available here. Specifically, when we permute 2 samples per group (i.e., the 8 Sites) randomly across 4 groups (the Locations), there are just 105 unique values of pseudo-F that can be obtained, so the minimum possible p-value here is $1/105 = 0.0095$).

5. Run a PERMANOVA

Step 4 (continued): Key additional details about PERMANOVA in PRIMER

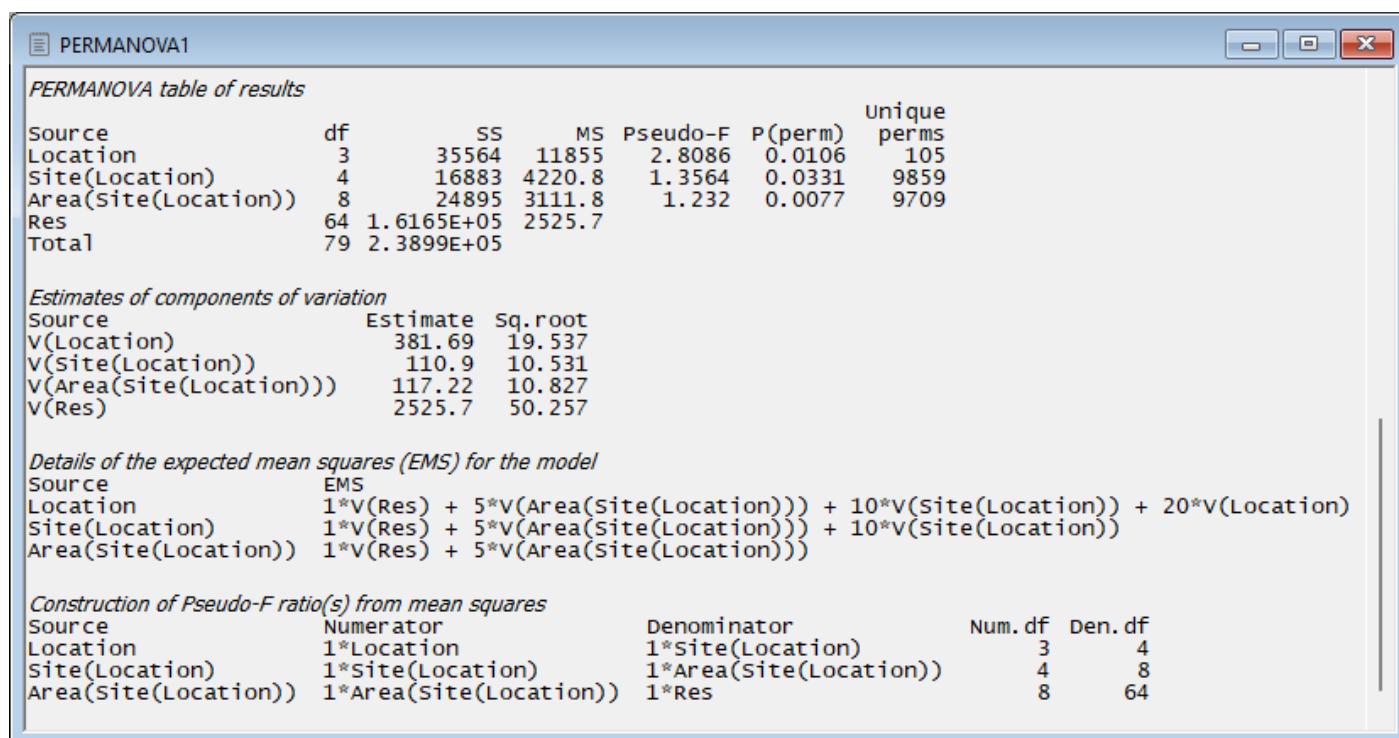
Following the PERMANOVA table of results, a suite of key additional details regarding the analysis can be seen in the PERMANOVA output file.

(Note: It is not necessary to fully unpack all of these details to continue on with the analysis and interpretation of results, but some information is provided here to highlight what makes the implementation of PERMANOVA in PRIMER so unique, surpassing all other software tools in its robust handling of multi-factorial experimental designs).

Additional details in the PERMANOVA output file include:

- estimates of the **components of variation** for each term in the model in the space of the resemblance measure.
- details of the **expected mean squares (EMS)** for each term in the model;
- the **construction of the pseudo-F ratios** for each term in the model from the appropriate mean squares (along with the associated numerator and denominator degrees of freedom); and

For example, scrolling down further in the output file from the PERMANOVA done on the holdfast data, we see:



PERMANOVA table of results

Source	df	SS	MS	Pseudo-F	P(perm)	Unique perms
Location	3	35564	11855	2.8086	0.0106	105
Site(Location)	4	16883	4220.8	1.3564	0.0331	9859
Area(Site(Location))	8	24895	3111.8	1.232	0.0077	9709
Res	64	1.6165E+05	2525.7			
Total	79	2.3899E+05				

Estimates of components of variation

Source	Estimate	Sq.root
V(Location)	381.69	19.537
V(Site(Location))	110.9	10.531
V(Area(Site(Location)))	117.22	10.827
V(Res)	2525.7	50.257

Details of the expected mean squares (EMS) for the model

Source	EMS
Location	1*V(Res) + 5*V(Area(Site(Location))) + 10*V(Site(Location)) + 20*V(Location)
Site(Location)	1*V(Res) + 5*V(Area(Site(Location))) + 10*V(Site(Location))
Area(Site(Location))	1*V(Res) + 5*V(Area(Site(Location)))

Construction of Pseudo-F ratio(s) from mean squares

Source	Numerator	Denominator	Num. df	Den. df
Location	1*Location	1*Site(Location)	3	4
Site(Location)	1*Site(Location)	1*Area(Site(Location))	4	8
Area(Site(Location))	1*Area(Site(Location))	1*Res	8	64

The implementation of PERMANOVA in PRIMER always uses expected mean squares (EMS) to construct correct tests for every term in the model; specifically:

- to construct the correct pseudo-F ratio; and
- to implement a permutation algorithm that
 - identifies the correct permutable units; and
 - accounts correctly for other terms in the model.

Of course, each term will require careful construction of its own test-statistic and its own permutation algorithm, both of which will depend on whether terms are fixed or random, whether there are other nested terms, covariates or interactions, etc. For unbalanced cases, the 'Type' of sum of squares is also important for the partitioning and correct construction of individual tests. All of these things can affect the EMS.

The EMS are also used to estimate the **components of variation** attributable to different sources of variation. These are *not* the same thing as raw R^2 values (as are typically used to compare the relative importance of individual predictor variables in regression models).[†] In PERMANOVA, these are calculated in a directly analogous way to the unbiased univariate ANOVA estimators. The column labeled 'Estimate' expresses these components in squared dissimilarity units, and its square root ('Sq.root', interpretable as a sort of standard deviation in the space of the resemblance measure) is also provided.

In the present example, we can see that the greatest variation occurs at the smallest spatial scale - from holdfast to holdfast within a given area (i.e., the square root of 'V(Res)' is 50.257 in Jaccard % dissimilarity units). The sources of variation (in order of importance, as quantified by the PERMANOVA model) are: Residual > Location > Area > Site.

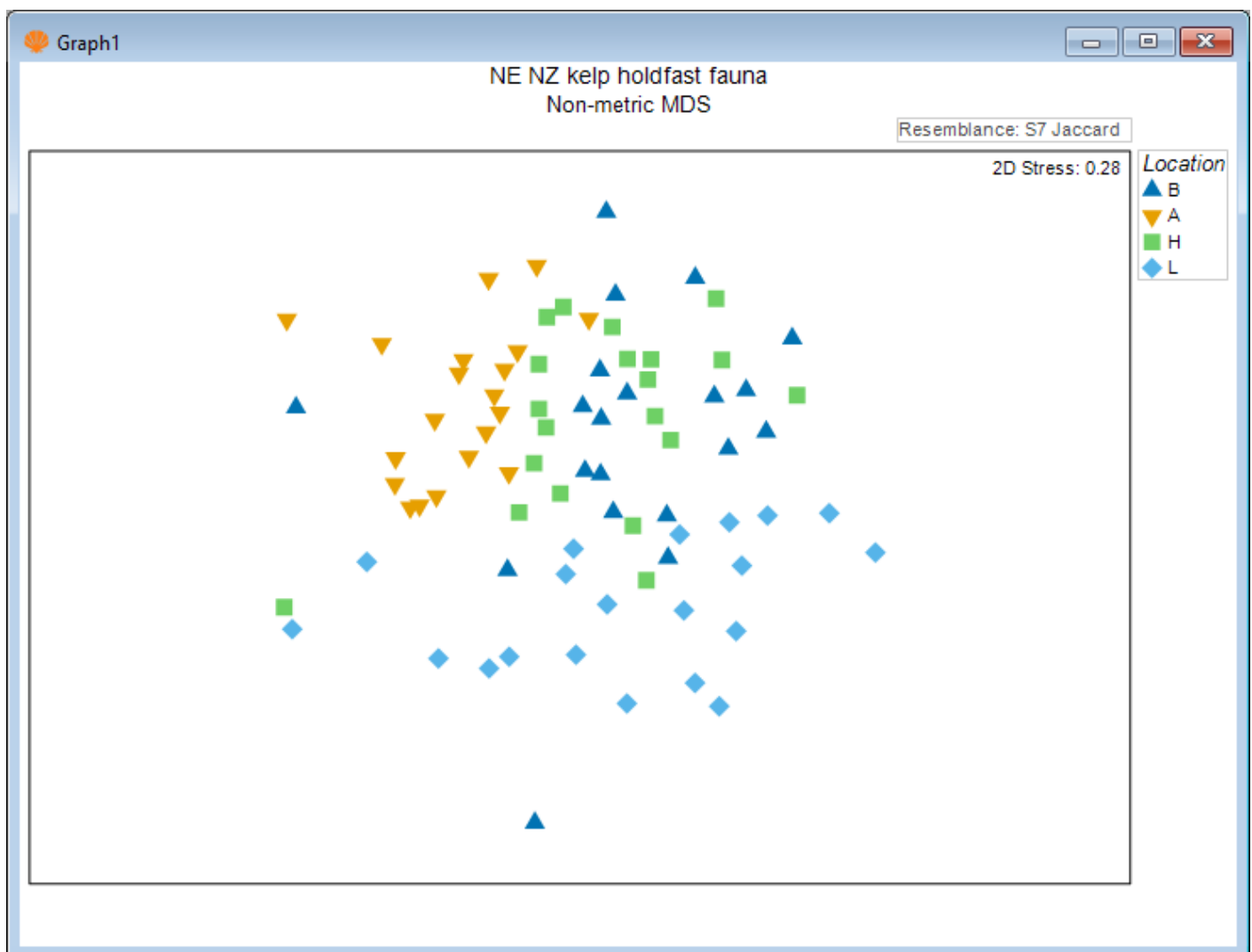
For more information about all of these key additional details provided in the PERMANOVA output file, please consult the [PERMANOVA+ manual](#). For a discussion of why using R^2 values to compare the relative importance of terms in ANOVA models is inappropriate, see the section entitled '[What about \$R^2\$ values?](#)' in the book: '[Should I use PRIMER or R?](#)'.

Step 5: Ordination of centroids

Having seen the results of a PERMANOVA analysis, it is natural to wish to see a **visualisation** of the patterns among centroids belonging to different groups or combinations of levels of different factors in the study design. In many cases, particularly if there is a large number of replicates in a complex study design, if one were simply to create an ordination of all individual sampling units, there would just be a lot of noise (the plot may look very messy), due to high residual variation. This tends to obscure salient patterns and important effects.

Ordination of sampling units (all replicates)

We can produce a non-metric MDS ordination of the holdfast data by starting from the Jaccard resemblance matrix and clicking on **Analyse > MDS > Non-metric MDS...**, taking all of the defaults, then clicking **OK**. The resulting configurations are very unsatisfactory with quite high stress, either in 2 dimensions (stress = 0.28, shown below), or 3 dimensions (stress = 0.20).

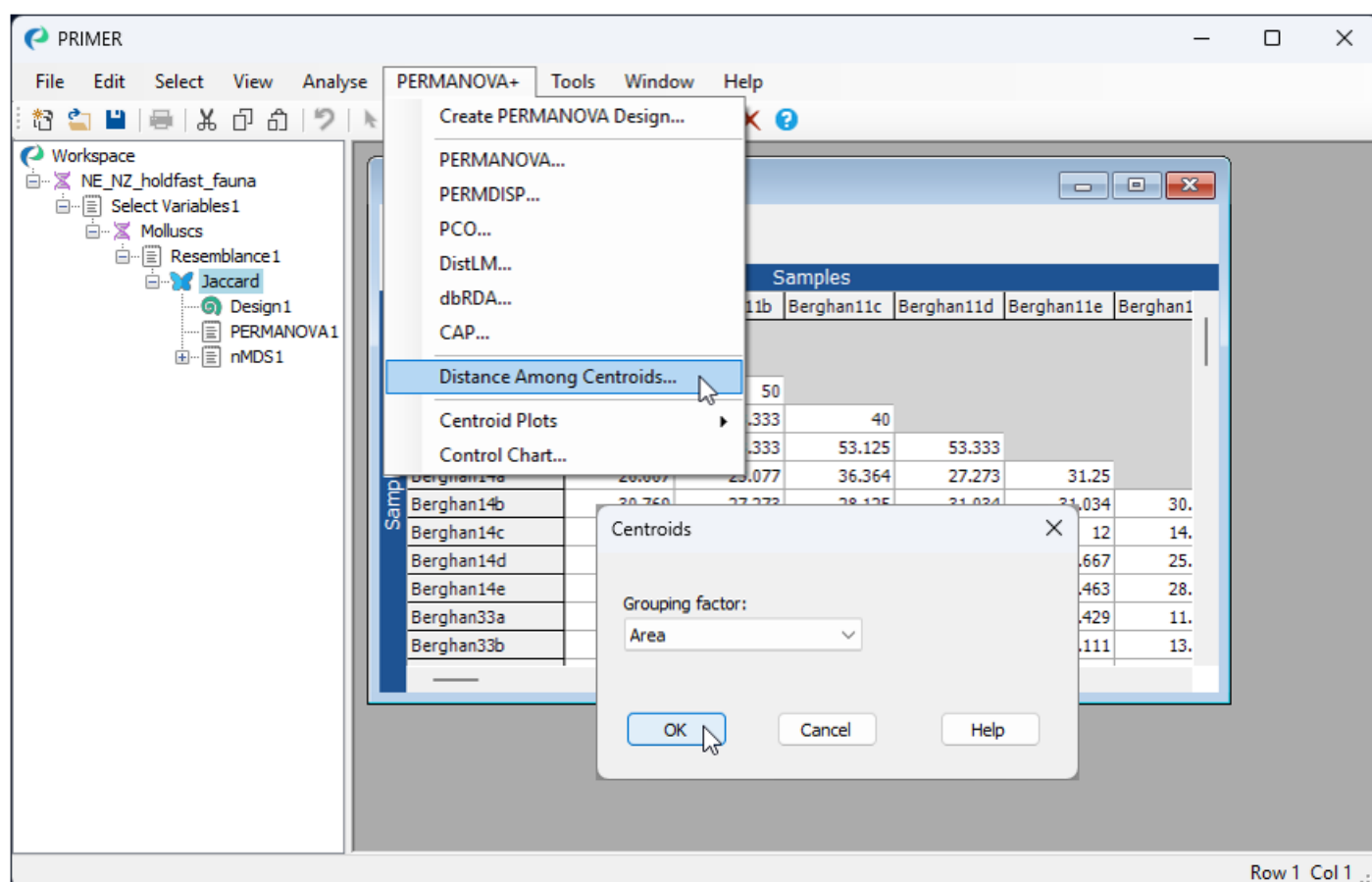


Seeing a bit of a 'mess' when we plot replicates like this is actually not too surprising in many cases, and particularly in this case, considering the very high variation (high turnover) in the identities of molluscs among holdfasts at small spatial scales (within areas), as seen on the previous page (recall that the residuals contributed by far the greatest source of variation to this system).

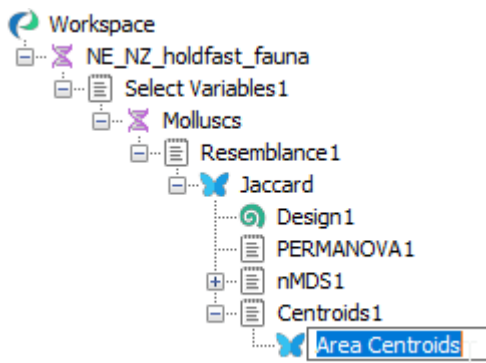
Ordination of distances among centroids

We can instead examine an ordination plot of the **centroids** in the space of the resemblance measure. In multi-factor designs when there is more than one factor and these factors are crossed with one another, we may need to create a single factor that consists of combinations of levels of the crossed factors (using **Edit > Factors... > Combine...**), but in the case of a fully nested design, and where nested factors all utilise unique labels (such as we have here), we can proceed without such a step.

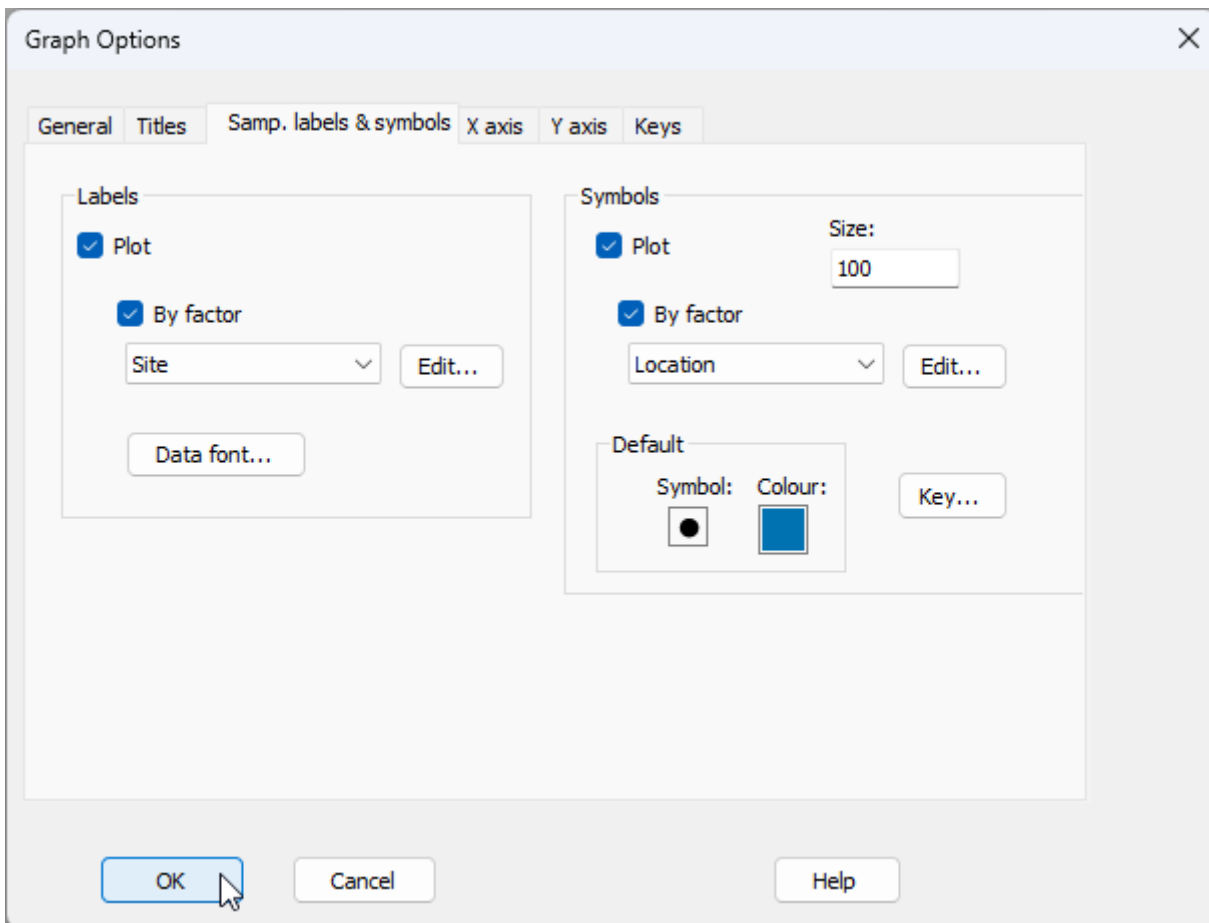
1. From the 'Jaccard' resemblance matrix, click **PERMANOVA+ > Distances Among Centroids**, then in the 'Centroids' dialog, choose (Grouping factor: **Area**), and click **OK**.



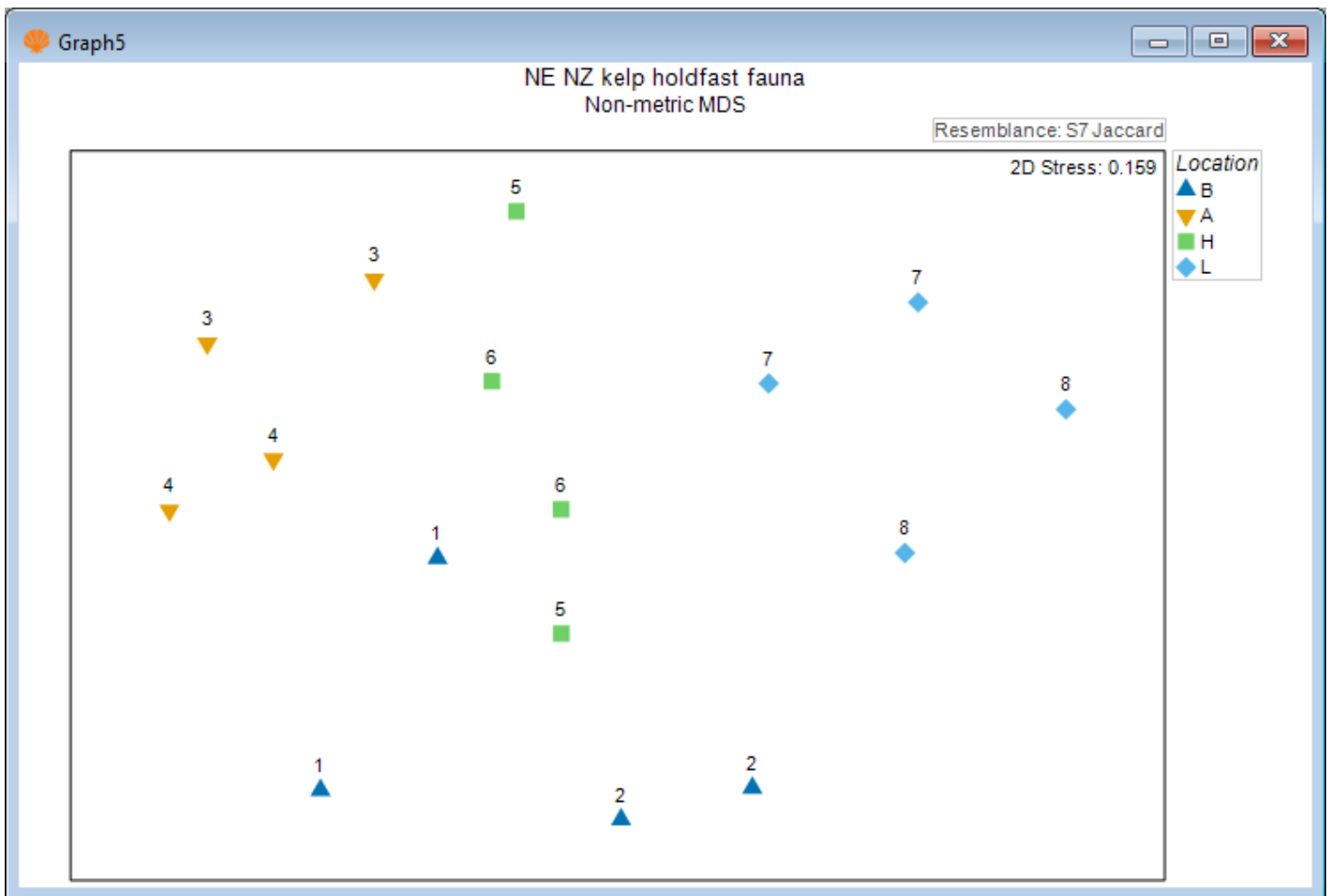
This will give you a matrix of Jaccard dissimilarities among the 16 Area centroids; each centroid being comprised of $n = 5$ replicate holdfasts within a given area. These are constructed in the space of the dissimilarity measure, which is not (quite) the same thing as calculating the arithmetic averages in the space of the original variables (i.e., that corresponds to a centroid in Euclidean space). In other words, these are the centroids just 'as PERMANOVA sees them' in *Jaccard* space, when doing the partitioning. You can re-name this matrix (from 'Resem1') to 'Area centroids'; i.e.,



2. From the 'Area centroids' resemblance matrix, click **Analyse > MDS > Non-metric MDS...**, take all the defaults and click **OK**.
3. From the 2D nMDS (probably called 'Graph5' in the Explorer tree), click on **Graph > Sample Labels & Symbols...**, then choose (Labels > ☒ Plot > ☒ By factor > Site) & (Symbols > ☒ Plot > ☒ By factor > Location), then click **OK**.



The result is a much more interpretable ordination plot, with far lower stress, viz:



Each point now represents the centroid (in Jaccard space) for $n = 5$ holdfasts in a given area. The numbers identify the 8 different sites, and the colours correspond to the 4 different locations. The patterns we see here are consistent with what was learned from the PERMANOVA analysis. More specifically, we can see that, within any particular location, the variation from one area to the next (any 2 centroids having the same symbol and number) is fairly similar to the variation between the two sites (any 2 centroids having the same colour, but a different number), and also that variation among locations (different colours) exceeds this - all four locations (colours) are clearly distinguishable from one another on the plot.

5. Run a PERMANOVA

Summary of the PERMANOVA analysis

A summary of the essential steps associated with performing this PERMANOVA analysis of the holdfast data according to the 3-factor hierarchical experimental design is given in the table below:

Step	To implement in PRIMER:
1. Select variable subset	<i>From the original data sheet:</i> click Select > Variables... > (\$\bullet\$Indicator levels > Indicator name: Phylum), click Levels... > (Selection > Include: Mollusca), click OK .
2. Jaccard resemblance	<i>From the subset-selected data sheet:</i> click Analyse > Resemblance > (Measure \$\bullet\$Other > S7 Jaccard) & (Analyse between \$\bullet\$Samples), click OK .
3. Specify the design	• <i>From the resemblance matrix:</i> click PERMANOVA+ > Create PERMANOVA design... • <i>Fill in the appropriate details of the design in the design file, including types of factors and any nested structures.</i>
4. Run PERMANOVA	<i>From the resemblance matrix:</i> click PERMANOVA+ > PERMANOVA... > (Design Worksheet: Design1) & (default settings for the rest), click OK .
5. Ordination of centroids	• <i>From the resemblance matrix:</i> click PERMANOVA+ > Distance Among Centroids... > (Grouping factor: Area), click OK. • <i>From the resulting resemblance matrix among area centroids:</i> click Analyse > MDS > Non-metric MDS..., take all the defaults and click OK. • <i>To specify labels and symbols:</i> click Graph > Sample Labels & Symbols...

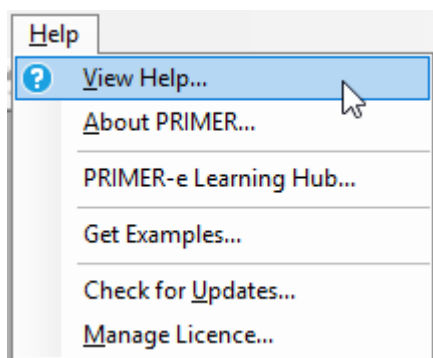
6. Help & additional resources

Help & support

You can get help with using PRIMER 8 software and diving deeper into your analyses in a variety of ways:

Learn the software

- Get **context-specific help** inside the software itself (click **Help** > **View Help...**).



- Go to PRIMER-e's **Learning Hub** (click **Help** > **PRIMER-e Learning Hub...**) to access extensive manuals and books (like this one!), that include a host of worked examples.
- Attend one of our **international courses** in multivariate analysis. Details of **upcoming courses**, offered both online and in person around the world, are always published on our website: <https://www.primer-e.com>.

Get advice on your analysis

- Contact our **statistical consulting team** to ask a specific question or get tailored advice, using PRIMER with PERMANOVA+ and other statistical software tools.

Contact us

- If you have questions about licensing, subscriptions, or any other general enquiries, check out the [Support page](#) on our website. If you can't find the answer to your question there, contact our admin team at: primer@primer-e.com.
- Please report any bugs or technical problems to: tech@primer-e.com.

Further reading

This guide is fairly succinct and is designed to get you rather quickly up and running with PRIMER 8 software. For detailed information about all of the *new features in PRIMER 8 with PERMANOVA+*, check out the following key resource:

- Anderson, M.J. (2026). **"What's New in PRIMER 8."** *PRIMER-e Learning Hub*. PRIMER-e, Auckland, New Zealand.

For additional details regarding the methods and tools available in the software produced by PRIMER-e, please consult the following *historical books and manuals*, available in the [PRIMER-e Learning Hub](#):

- Clarke KR, Gorley RN, Somerfield PJ & Warwick RM. (2014). *"Change in Marine Communities, 3rd edition."* PRIMER-E Ltd: Plymouth, UK.
- Clarke KR & Gorley RN. (2015). *"PRIMER v7: User Manual / Tutorial."* PRIMER-E Ltd: Plymouth, UK.
- Anderson MJ, Gorley RN & Clarke KR. (2008). *"PERMANOVA+ for PRIMER: Guide to Software and Statistical Methods."* PRIMER-E Ltd: Plymouth, UK.

References

[Anderson \(2001a\)](#)

Anderson, M.J. (2001a) A new method for non-parametric multivariate analysis of variance. *Austral Ecology*, **26**, 32-46.

[Anderson \(2001b\)](#)

Anderson, M.J. (2001b) Permutation tests for univariate or multivariate analysis of variance and regression. *Canadian Journal of Fisheries and Aquatic Sciences*, **58**, 626-639.

[Anderson \(2017\)](#)

Anderson, M.J. (2017) Permutational multivariate analysis of variance (PERMANOVA). *Wiley StatsRef: Statistics Reference Online*, **stat07941**, 1-15.

[Anderson et al. \(2005a\)](#)

Anderson, M.J., Connell, S.D., Gillanders, B.M., Diebel, C.E., Blom, W.M., Saunders, J.E. & Landers, T.J. (2005a) Relationships between taxonomic resolution and spatial scales of multivariate variation. *Journal of Animal Ecology*, **74**, 636-646.

[Anderson et al. \(2011\)](#)

Anderson, M.J., Crist, T.O., Chase, J.M., Vellend, M., Inouye, B.D., Freestone, A.L., Sanders, N.J., Cornell, H.V., Comita, L.S., Davies, K.F., Harrison, S.P., Kraft, N.J.B., Stegen, J.C. & Swenson N.G. (2011) Navigating the multiple meanings of β diversity: a roadmap for the practicing ecologist. *Ecology Letters*, **14**, 19-28.

[Anderson et al. \(2005b\)](#)

Anderson, M.J., Diebel, C.E., Blom, W.M. & Landers, T.J. (2005b) Consistency and variation in kelp holdfast assemblages: spatial patterns of biodiversity for the major phyla at different taxonomic resolutions. *Journal of Experimental Marine Biology and Ecology*, **320**, 35-56.

[Anderson & ter Braak \(2003\)](#)

Anderson, M.J. & ter Braak, C.J.F. (2003) Permutation tests for multi-factorial analysis of variance. *Journal of Statistical Computation and Simulation*, **75**, 85-113.

[Clarke \(1993\)](#)

Clarke, K.R. (1993) Non-parametric multivariate analyses of changes in community structure. *Australian Journal of Ecology*, **18**, 117-143.

[Clarke et al. \(2006\)](#)

Clarke, K.R., Somerfield, P.J. & Chapman, M.G. (2006) On resemblance measures for ecological studies, including taxonomic dissimilarities and a zero-adjusted Bray-Curtis coefficient for denuded assemblages. *Journal of Experimental Marine Biology and Ecology*, **330**, 55-80.

[Gray et al. \(1990\)](#)

Gray, J.S., Clarke, K.R., Warwick, R.M. & Hobbs, G. (1990) Detection of initial effects of pollution on marine benthos: an example from the Ekofisk and Eldfisk oilfields, North Sea. *Marine Ecology Progress Series*, **66**, 285-299.

[Kruskal & Wish \(1978\)](#)

Kruskal, J.B. & Wish, M. (1978) *Multidimensional scaling*. Beverly Hills: Sage Publications.

[Legendre & Legendre \(2012\)](#)

Legendre, P. & Legendre, L. (2012) *Numerical ecology, 3rd English edition*. Amsterdam: Elsevier.

[McArdle & Anderson \(2001\)](#)

McArdle, B.H. & Anderson, M.J. (2001) Fitting multivariate models to community data: a comment on distance-based redundancy analysis. *Ecology*, **82**, 290-297.

[Somerfield et al. \(1994a\)](#)

Somerfield, P.J., Gee, J.M. & Warwick, R.M. (1994a) Benthic community structure in relation to an instantaneous discharge of waste water from a tin mine. *Marine Pollution Bulletin*, **28**, 363-369.

[Somerfield et al. \(1994b\)](#)

Somerfield, P.J., Gee, J.M. & Warwick, R.M. (1994b) Soft sediment meiofaunal community structure in relation to a long-term heavy metal gradient in the Fal estuary system. *Marine Ecology Progress Series*, **105**, 79-88.

[Somerfield et al. \(2021a\)](#)

Somerfield, P.J., Clarke, K.R. & Gorley, R.N. (2021a) A generalised analysis of similarities (ANOSIM) statistic for designs with ordered factors. *Austral Ecology*, **46**, 901-910.

[Somerfield et al. \(2021b\)](#)

Somerfield, P.J., Clarke, K.R. & Gorley, R.N. (2021b) Analysis of Similarities (ANOSIM) for 2-way layouts using a generalised ANOSIM statistic, with comparative notes on Permutational Multivariate Analysis of Variance (PERMANOVA). *Austral Ecology*, **46**, 911-926.

[Somerfield et al. \(2021c\)](#)

Somerfield, P.J., Clarke, K.R. & Gorley, R.N. (2021c) Analysis of Similarities (ANOSIM) for 3-way designs. *Austral Ecology*, **46**, 927-941.
