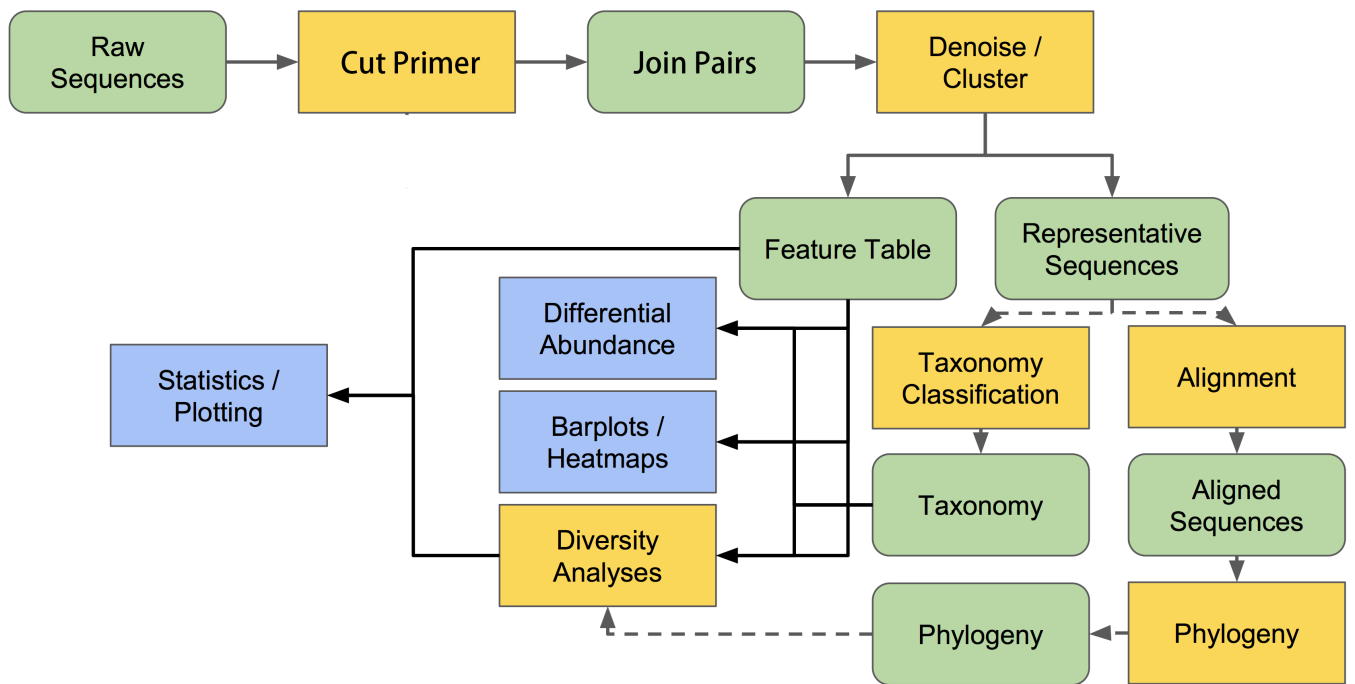


目录

- 课程简介
- 16S rRNA基因扩增子测序的基本原理
- 搭建分析环境
 - 配置zshell终端及常用命令介绍
 - 安装mamba和sankemake
 - 扩增子相关软件安装
- 分析流程拆解
 - 去引物与裁剪低质量序列(Cutadapt)
 - 双端序列拼接(Flash)
 - dada2及后续流程分析
- Snakemake分析流程自动化
 - 准备基础工作目录和修改配置文件
 - 配置VS Code
 - 运行snakemake
 - 结果文件说明
- 下游数据分析与可视化
 - 绘制物种组成丰度堆叠柱状图
 - alpha多样性可视化
 - beta多样性可视化
- 微生物共现网络分析(Hmisc, igraph)
- Mantel test环境因子与物种组成相关性分析
- CCA、RDA环境因子与物种分布相关性分析
- 导入qiime2

1. 课程简介

扩增子测序的分析流程

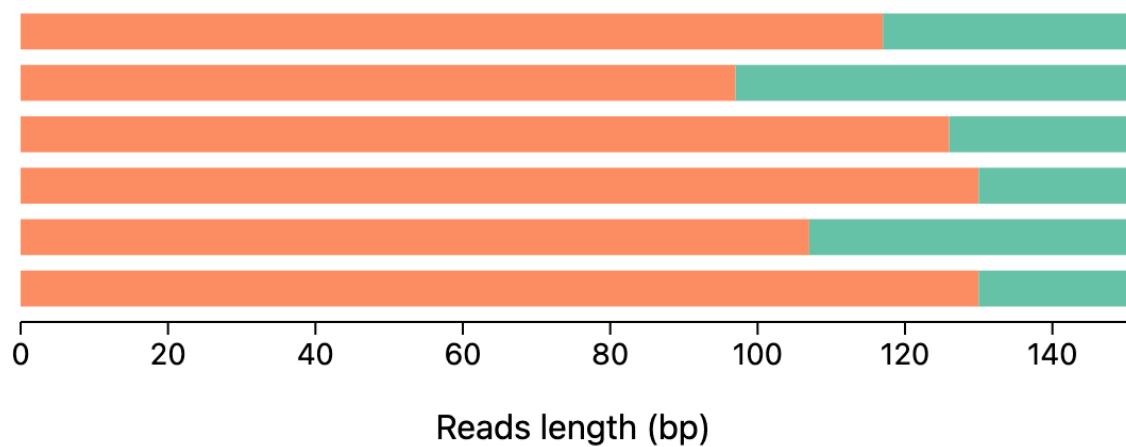


统一裁剪的弊端

[非统一裁剪文章例子1](#)

[非统一裁剪文章例子2](#)

- 每一对reads不合格的地方都是随机的，统一裁剪会导致某些合格的区域被过度裁剪，而且是影响了所有参与计算的样本。
- 导致分析过程繁琐，不同批次的分析都要人工去决定裁剪到哪个长度。
- 对于输入是指控过后的长短不一的数据，低于裁剪长度的read会被直接丢弃，但只要这对reads有足够的overlap区域就可以拼接起来，没必要丢弃。





适合人群

- 有一定的生信基础
- 会基础的linux和R语言
- 希望自己分析16S扩增子

你将学到什么？

- 扩增子测序基本原理
- 扩增子分析的基本流程和相关概念
- 一些实用的linux技巧、zshell命令
- 使用mamba搭建分析环境
- 相关软件参数的使用经验
- 相关R包的使用
- snakemake分析流程自动化

资料下载

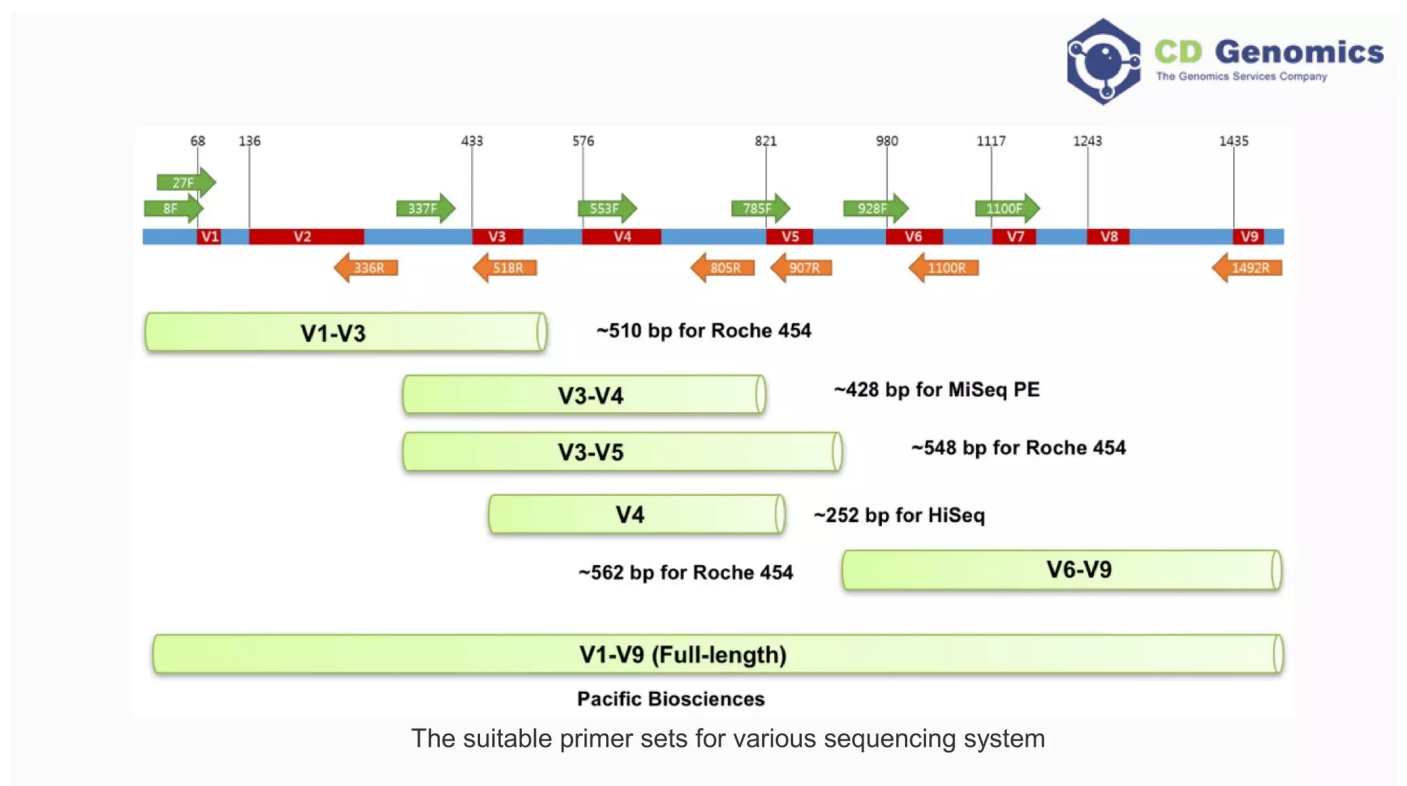
网页版讲义更好用！

- 课程资料及网页版讲义：[链接](#)
- 添加微信进交流群，请备注来自CCtalk

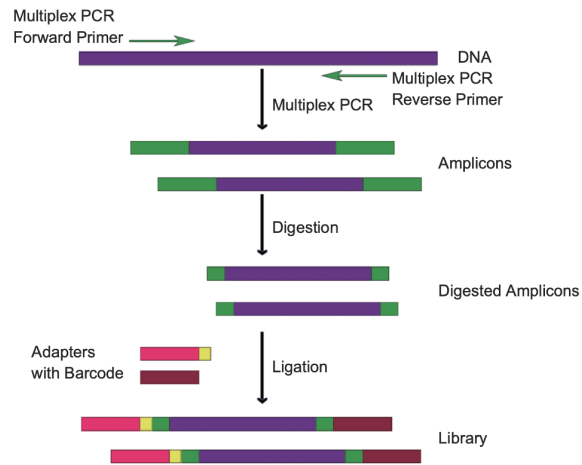


2. 16S rRNA基因扩增子测序的基本原理

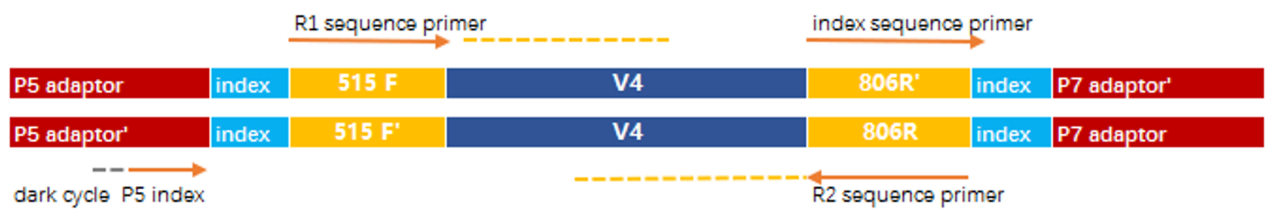
[原理ppt](#)



- 建库流程



VAHTS® AmpSeq Library Prep Kit V3

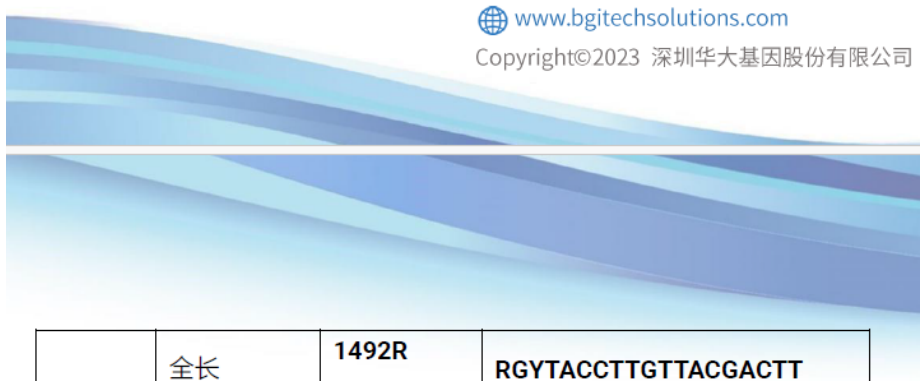


- 测序引物

3、测序策略

	扩增区域	引物名称	引物对
细菌	V4	515F	GTGCCAGCMGCCGCGGTAA
		806R	GGACTACHVGGGTWTCTAAT
	V1-V3	8F	AGAGTTTGATYMTGGCTCAG
		518R	ATTACCGCGGCTGCTGG
	V3-V4	341F	ACTCCTACGGGAGGCAGCAG
		806R	GGACTACHVGGGTWTCTAAT
真菌	V4-V5	515F	GTGCCAGCMGCCGCGG
		907R	CCGTCAATTCMTTTRAGT
	ITS1	its1	CTTGGTCAATTAGAGGAAGTAA
		its2	GCTGCGTTCTTCATCGATGC
	ITS2	its3	GCATCGATGAAGAACGCAGC
		its4	TCCTCCGCTTATTGATATGC

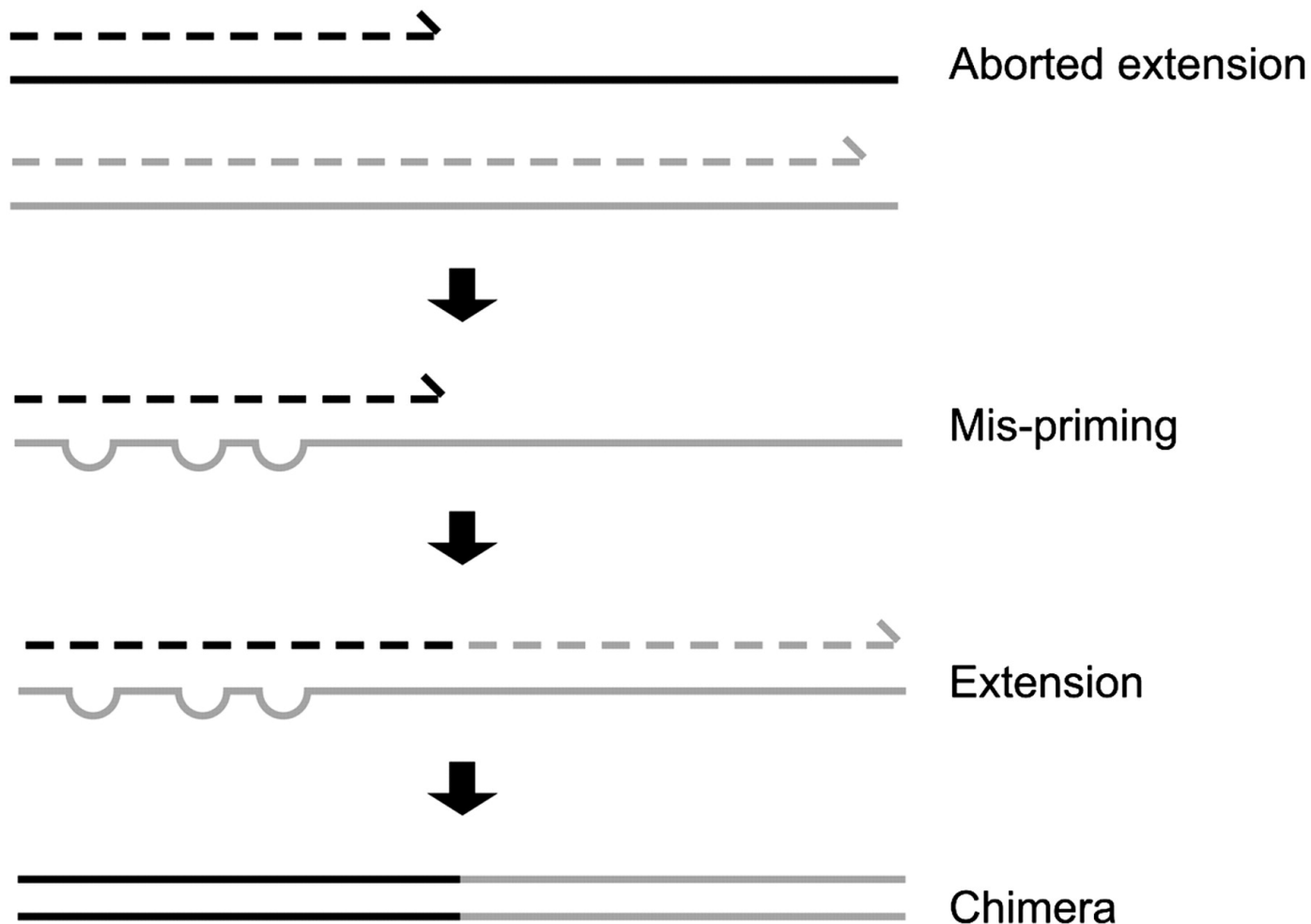
	扩增区域	引物名称	引物对
细菌	V4	515F	GTGCCAGCMGCCGCGGTAA
		806R	GGACTACHVGGGTWTCTAAT
	V3-V4	338F	ACTCCTACGGGAGGCAGCAG
		806R	GGACTACHVGGGTWTCTAAT
		27F	AGRGTTYGATYMTGGCTCAG



	全长	1492R	RGYTACCTTGTTACGACTT
真菌	ITS1	its1	CTTGGTCATTTAGAGGAAGTAA
		its2	GCTGCGTTCTTCATCGATGC
	ITS2	its3	GCATCGATGAAGAACGCAGC
		its4	TCCTCCGCTTATTGATATGC

- chimera(嵌合体)

[参考文献](#)



3. 搭建分析环境

3.1 配置一个舒适的zshell终端

- 安装zshell

```
1 | sudo apt install zsh
```

- 更改用户的shell程序

```
1 | chsh -s /bin/zsh
```

- 安装oh-my-zsh: [官网](https://github.com/ohmyzsh/ohmyzsh)

```
1 | sh -c "$(wget https://raw.githubusercontent.com/ohmyzsh/ohmyzsh/master/tools/install.sh -O -)"
```

- 更改主题

```
1 # 修改配置文件
2 vim ~/.zshrc
3 # 修改主题名称
4 ZSH_THEME="ys"
```

- 添加自动提示插件

```
1 # 把下载好的插件移动到插件目录
2 mv zsh-autosuggestions/ ~/.oh-my-zsh/custom/plugins/
3 # 修改插件配置
4 vim ~/.zshrc
5 plugins=(git last-working-dir zsh-autosuggestions)
```

- 实用zshell命令

```
1 # 批量重命名
2
3 autoload -U zmv
4 zmv 'G([123]A).fq.gz' 'G-7D-$1.fq.gz'
5 zmv 'G(*).(*).gz' 'G-$1.$2kk.gz'
6
7 # 递归搜索
8
9 ## 查找所有普通文件
10 print -l **/*(.)
11 ## 查找所有后缀是gz结尾的文件
12 print -l **/*.gz
13 ## 列出最近一天修改过内容的文件
14 print -l *(.m-1)
15 ## 列出最近一个月没有读取过的文件
16 print -l *(.aM+1)
17 ## 列出当前目录下小于 2k 的文件
18 print -l *(.Lk-2)
19
20 # 字符串截取
21
22 aa="1234"
23 ## 取第一个
24 echo $aa[1]
25 ## 取最后一个
26 echo $aa[-1]
27 ## 取中间两个
28 echo $aa[2,3]
29
30 # 字符串拼接
31
32 % str1=abc
33 % str2=def
34
35 % str2+=${str1}
```

```

36 % echo $str2
37 defabc
38
39 # 字符串替换
40
41 % str=abcabc
42
43 ## 只替换找到的第一个
44 % echo ${str/bc/ef}
45 aefabc
46
47 # for循环
48 for i (*){
49 echo $i
50 }
51 # while循环
52 while {read i} {
53 echo $i
54 } < Flavo_accession.txt

```

3.2 安装mamba和snakemake

- Miniforge的github地址: [Miniforge](https://github.com/conda-forge/miniforge)

```

1 # 下载安装脚本到本地
2 wget https://github.com/conda-forge/miniforge/releases/latest/download/Miniforge3-
Linux-x86_64.sh
3 # 添加执行权限
4 chmod +x Miniforge3-Linux-x86_64.sh
5 # 运行脚本
6 ./Miniforge3-Linux-x86_64.sh
7 # 更新
8 source ~/.zshrc

```

- snakemake

```

1 # 单独创建一个环境
2 mamba create -n snakemake -c bioconda snakemake graphviz
3 # 激活环境
4 mamba activate snakemake

```

3.3 扩增子相关软件安装

- 新建目录

```

1 mkdir 16S_rRNA_amplicon
2 cd 16S_rRNA_amplicon

```

- 创建环境依赖清单文件 `env.yaml`

```
1 channels:
2   - conda-forge
3   - bioconda
4 dependencies:
5   - bioconductor-shortread
6   - bioconductor-dada2
7   - bioconductor-phyloseq
8   - r-pheatmap
9   - r-ggplot2
10  - r-dplyr
11  - r-phangorn
12  - flash
13  - cutadapt
14  - mafft
15  - fasttree
16  - biopython
17  - pandas
```

- 创建环境并且安装上述依赖

```
1 mamba env create -n 16S_rRNA_amplicon --file env.yaml
2 # 激活环境
3 mamba activate 16S_rRNA_amplicon
```

- 安装nodejs

[nodejs官网](#)

```
1 # 下载
2 wget https://nodejs.org/dist/v20.11.0/node-v20.11.0-linux-x64.tar.xz
3 # 解压
4 tar -xf node-v20.11.0-linux-x64.tar.xz
5 # 添加到环境变量
6 vim ~/.zshrc
7 export PATH=$PATH:/home/xjm/software/node-v20.11.0-linux-x64/bin
8
9 # 更新
10 source ~/.zshrc
```

- 安装puppeteer

[github](#)

```
1 # 创建一个目录
2 mkdir puppeteer
3 cd puppeteer
4 # 初始化工程
5 npm init
```

```

6
7 📌👤 # xjm @ lmit-X12DPi-N-T-6 in /userdata2/cc/TaiwanStrait-
16S/spring_custom_pipeline/puppeteer [16:43:45]
8 📌 cat package.json
9 {
10     "name": "puppeteer",
11     "version": "1.0.0",
12     "description": "",
13     "main": "app.js",
14     "scripts": {
15         "test": "echo \"Error: no test specified\" && exit 1"
16     },
17     "author": "",
18     "license": "ISC"
19 }
20
21 # 更改module模式
22
23 vim package.json
24
25 📌👤 # xjm @ lmit-X12DPi-N-T-6 in /userdata2/cc/TaiwanStrait-
16S/spring_custom_pipeline/puppeteer [16:45:44]
26 📌 cat package.json
27 {
28     "name": "puppeteer",
29     "version": "1.0.0",
30     "description": "",
31     "type": "module",
32     "main": "app.js",
33     "scripts": {
34         "test": "echo \"Error: no test specified\" && exit 1"
35     },
36     "author": "",
37     "license": "ISC"
38 }
39
40 # 安装
41 npm i puppeteer
42 # 或者
43 pnpm i puppeteer

```

[pnpm](#)应该类似mamba之于conda

- 复制 app.js

```
1 mv app.js puppeteer/
```

4. 分析流程拆解

4.1 去引物与裁剪低质量序列

准备示例数据

- 解压

```
1 | tar -xzf example_data.tar.gz
```

- 目录结构：3个环境水体的采样站点，每个站点有3个平行，raw data

```
1 | # example_data
2 | └─ S1-1
3 |   └─ S1-1_1.fq.gz
4 |     └─ S1-1_2.fq.gz
5 | └─ S1-2
6 |   └─ S1-2_1.fq.gz
7 |     └─ S1-2_2.fq.gz
8 | └─ S1-3
9 |   └─ S1-3_1.fq.gz
10 |     └─ S1-3_2.fq.gz
11 | └─ S2-1
12 |   └─ S2-1_1.fq.gz
13 |     └─ S2-1_2.fq.gz
14 | └─ S2-2
15 |   └─ S2-2_1.fq.gz
16 |     └─ S2-2_2.fq.gz
17 | └─ S2-3
18 |   └─ S2-3_1.fq.gz
19 |     └─ S2-3_2.fq.gz
20 | └─ S3-1
21 |   └─ S3-1_1.fq.gz
22 |     └─ S3-1_2.fq.gz
23 | └─ S3-2
24 |   └─ S3-2_1.fq.gz
25 |     └─ S3-2_2.fq.gz
26 | └─ S3-3
27 |   └─ S3-3_1.fq.gz
28 |     └─ S3-3_2.fq.gz
```

- 测序区域为V3~V4区

F_primer: ACTCCTACGGGAGGCAGCAG

R_primer: GGACTACHVGGGTWTCTAAT

cutadapt

[官方文档](#)

```
1 # -g 5' adapter
2 # -G 5' adapter to be removed from R2
3 # -e 最大允许的错频率 (0 <= E < 1) 或者错配个数 (E >= 1) , 因为你给定的引物对并不是百分百匹配的
4 # --discard-untrimmed 去掉那些没有检测到adapter的序列 这些序列往往是错误序列 测序错误或者是PCR扩增错误
5 # -q 裁剪3'端低质量碱基, 没有裁剪会导致后续拼接overlap区域的错配偏高
6 # --cores 指定线程
7 cutadapt -q 20 --discard-untrimmed -e 4 --cores 32 -g ACTCCTACGGGAGGCAGCAG -G
GGACTACHVGGGTWTCTAAT -o cutadapt_output/S1-1/S1-1_1.fq.gz -p cutadapt_output/S1-1/S1-
1_2.fq.gz example_data/S1-1/S1-1_1.fq.gz example_data/S1-1/S1-1_2.fq.gz
```

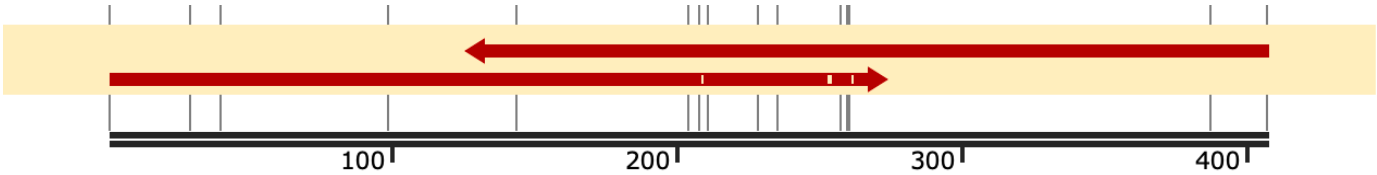
- Degenerate bases

IUPAC nucleotide code	Base
R	A or G
Y	C or T
S	G or C
W	A or T
K	G or T
M	A or C
B	C or G or T
D	A or G or T
H	A or C or T
V	A or C or G
N	any base

4.2 双端序列拼接

flash

[官网](#)



特别注意这个参数: -M 最大重叠长度, 如果重叠区大于这个长度也会正常拼接, 但是下面的错配比率的分母还是按照这个值去算而不是整个重叠区长度, 会导致拼接失败的比例升高, 非常鸡肋, 你自己自动识别不就行 [default: 65]

```

1 # -m 两个reads之间可靠重叠所需的最小重叠长度 [default: 10]
2 # -M 最大重叠长度，如果重叠区大于这个长度也会正常拼接，但是下面的错配比率的分母还是按照这个值去算而不是整个重叠区长度，会导致拼接失败的比例升高 [default: 65]
3 # -x 两个读段的错配碱基数与重叠长度之间的最大允许比率 [default: 0.25]
4 # -o 输出结果前缀
5 # -d 输出目录
6 # -t 线程，如果想按输入文件的顺序输出就把线程设置为1
7 # -p 质量分数偏移值 [Default: 33]
8 flash -m 15 -M 250 -x 0.1 -o S1-1 -d flash_output/S1-1 -t 16 cutadapt_output/S1-1/S1-1_1.fq.gz cutadapt_output/S1-1/S1-1_2.fq.gz

```

4.3 dada2及后续流程分析

4.3.1 RStudio安装相关R包

```

1 .cran_packages <- c("dplyr", "phangorn")
2 .bioc_packages <- c("dada2", "phyloseq", "ShortRead")
3 if(!requireNamespace("BiocManager")){
4   install.packages("BiocManager")
5 }
6
7 .inst <- .cran_packages %in% installed.packages()
8 if(any(!.inst)) {
9   install.packages(.cran_packages[!.inst])
10 }
11 .inst <- .bioc_packages %in% installed.packages()
12 if(any(!.inst)) {
13   BiocManager::install(.bioc_packages[!.inst])
14 }

```

4.3.2 生成去嵌合体的特征表与代表序列

- ASV (Amplicon Sequence Variants)
- 嵌合体

4.3.3 物种注释、去叶绿体和线粒体序列

自定义数据库格式

`assignTaxonomy(...)` 需要一个训练fasta文件（或压缩的fasta文件），其中每个序列对应的分类信息被编码在id行中，格式如下（第二个序列没有被分配到第6级）：

```

1 >Level1;Level2;Level3;Level4;Level5;Level6;
2 ACCTAGAAAGTCGTAGATCGAAGTTGAAGCATCGCCCGATGATCGTCTGAAGCTGTAGCATGAGTCGATTTTCACATTCAGGGATA
  CCATAGGATAC
3 >Level1;Level2;Level3;Level4;Level5;
4 CGCTAGAAAGTCGTAGAAGGCTCGGAGGTTTGAAGCATCGCCCGATGGGATCTCGTTGCTGTAGCATGAGTACGGACATTCAGGGA
  TCATAGGATAC

```

`assignSpecies(...)` 和 `addSpecies(...)` 期望训练数据以 fasta 文件的形式提供（或压缩的 fasta 文件），其中 id 行的格式如下：

```
1 >ID Genus species
2 ACCTAGAAAGTCGTAGATCGAAGTTGAAGCATCGCCCGATGATCGTCTGAAGCTGTAGCATGAGTCGATTTTCACATTCAGGGATA
  CCATAGGATAC
3 >ID Genus species
4 CGCTAGAAAGTCGTAGAAGGCTCGGAGGTTTGAAGCATCGCCCGATGGGATCTCGTTGCTGTAGCATGAGTACGGACATTCAGGGA
  TCATAGGATAC
```

4.3.4 alpha多样性计算

1. 丰富度: **Richness**

衡量一个生态系统有多少不同的物种

2. 均匀度: **Evenness**

衡量生态系统中，不同物种之间数量的差异度

- [Shannon](#)
- [Simpson's Index](#)
- [ACE和Chao1](#)

```
1 # 辛普森多样性指数是基于在一个无限大的群落中，随机抽取两个个体，它们属于同一物种的概率
2 # Simpson's Index: D
3 # 这里的Simpson 算的应该是Simpson's Index of Diversity: 1-D
4 # "Observed", "Chao1", "ACE"基本值都是一样的，因为样本里面数目为1或者2的ASV基本没有
```

4.3.5 构建ASV系统发育树

- `mafft` 多序列比对

```
1 mkdir mafft_output
2 mafft --auto dada2_single_end_output/ASV_seqs_no_chloro_mito.fasta >
  mafft_output/ASV_seqs_no_chloro_mito_align.fasta
```

- `FastTree` 最大似然法构建进化树

```
1 mkdir fasttree_output
2 FastTree -gtr -nt mafft_output/ASV_seqs_no_chloro_mito_align.fasta >
  fasttree_output/ASV_seqs_no_chloro_mito_tree.nwk
```

4.3.6 beta多样性计算

Beta多样性是生物多样性研究中的一个重要概念，用于衡量不同生态系统或样本之间的物种多样性差异程度。它通常通过计算物种组成的差异性来衡量，可以帮助我们了解不同地点或环境中物种的多样性情况。

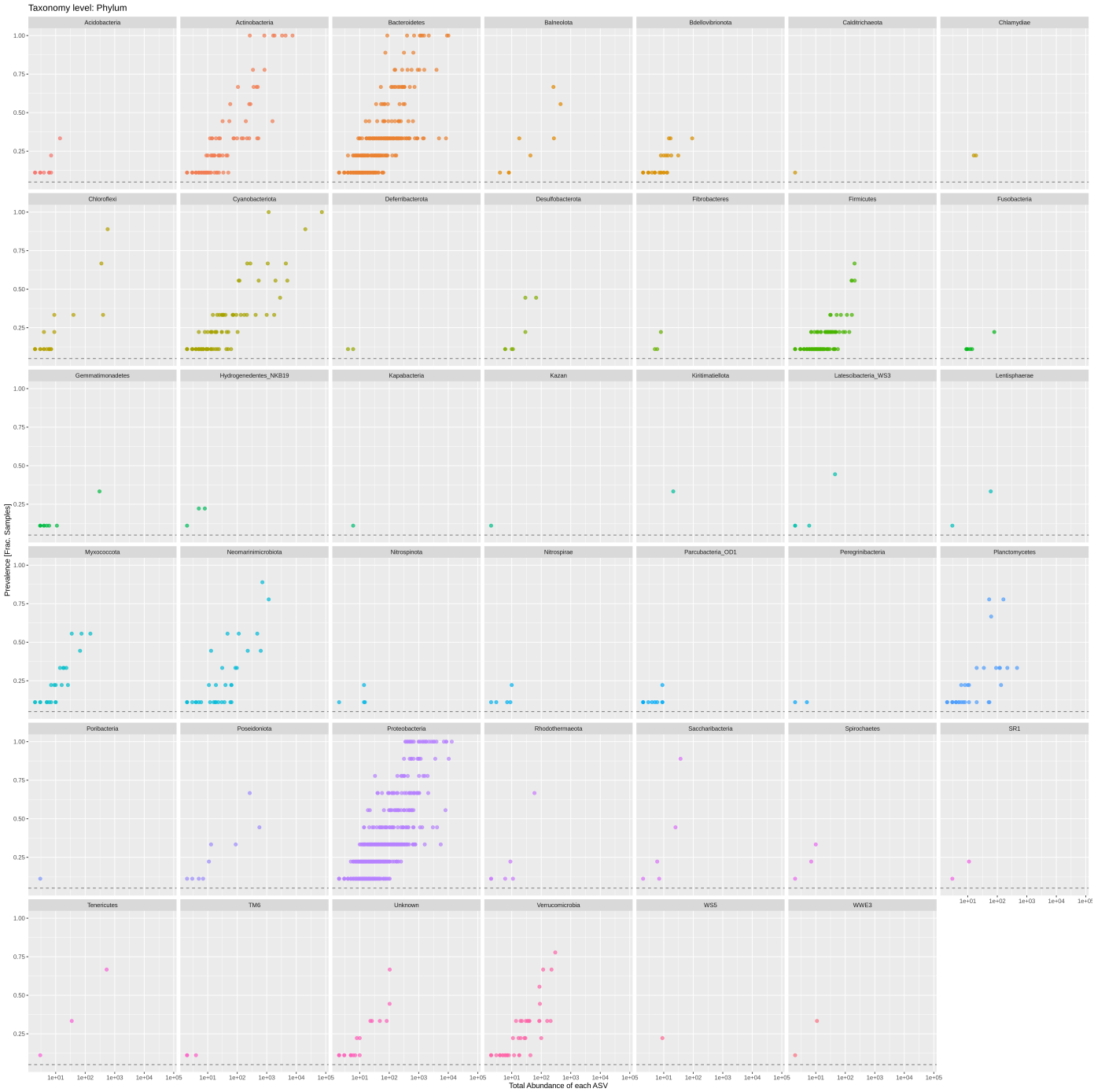
4.3.7 ASV流行度和总丰度计算

- 流行度: **Prevalence**

一个物种出现在越多的样品中，这个物种的流行度越高。

- 总丰度: **Total Abundance**

一个物种在其所有出现的样品中的丰度总和。



5. Snakemake分析流程自动化

5.1 准备基础工作目录和修改配置文件

- 复制基础工作目录

```
1 # 解压
2 tar -xzf custom_pipeline.tar.gz
```

```
1 # 整个工作目录结构
2 custom_pipeline
3 |— configs
4 |   └─ 16S_rRNA_amplicon_config.yaml
5 |— taxonomy_databases
6 |   └─ silva_nr99_v138.1_train_Mitochondria_Chloroplast_set.fa
7 |   └─ silva_nr99_v138.1_train_set.fa.gz
8 |   └─ silva_species_assignment_v138.1.fa.gz
9 |— workflow
10 |   └─ 16S_rRNA_amplicon_EZ.smk
11 |   └─ 16S_rRNA_amplicon_Silva.smk
12 |   └─ scripts
13 |       └─ python
14 |           └─ ezbiocloud.py
15 |       └─ R
16 |           └─ dada2
17 |               └─ assign_taxonomy.R
18 |               └─ dada2_single_reads_multi_samples.R
19 |               └─ downstream_analysis.R
20 |           └─ reads_length_frequency_distribution.R
```

- 修改 configs/16S_rRNA_amplicon_config.yaml

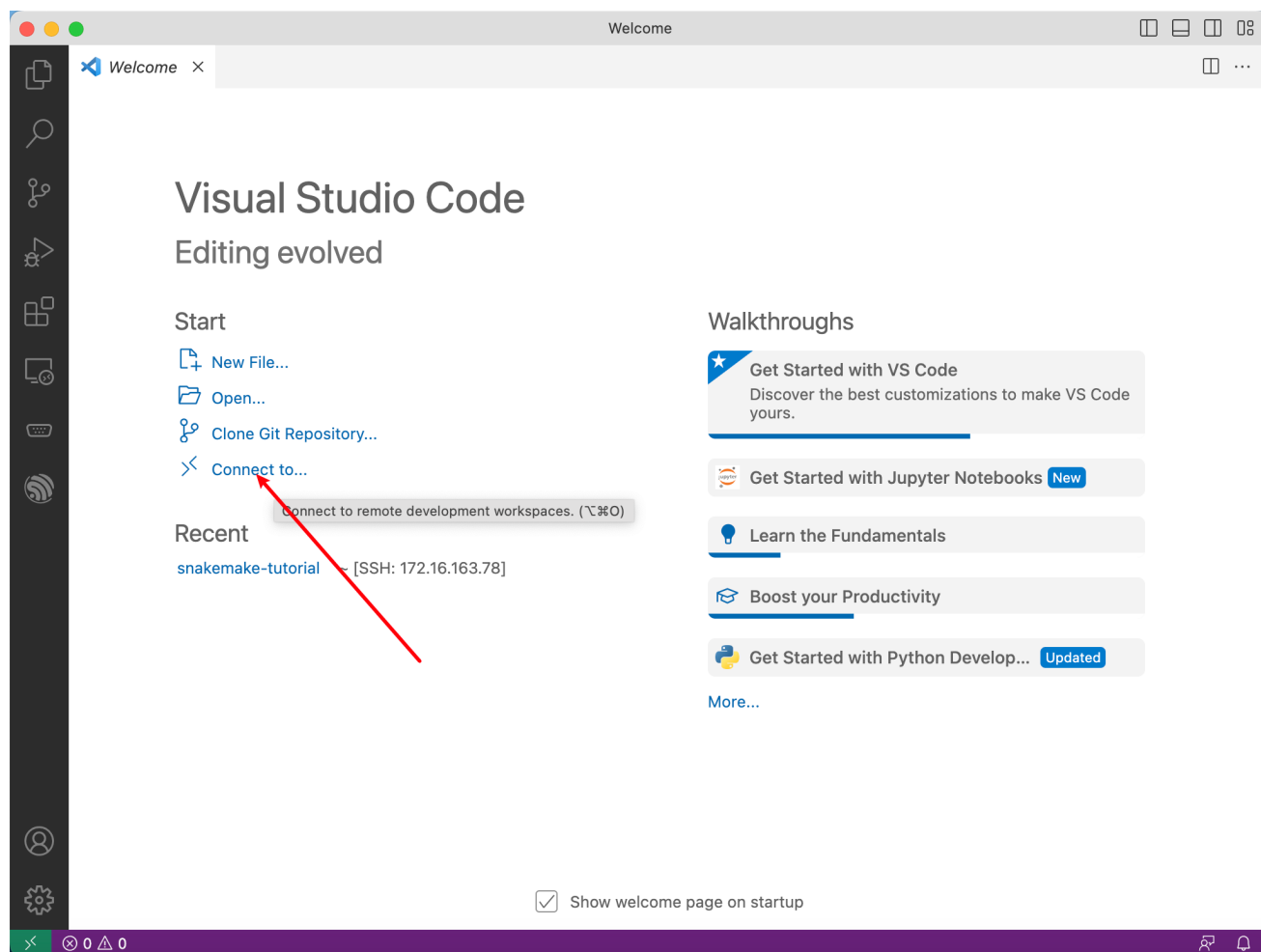
```
1 samples_dir: "samples_XJM"
2 f_suffix: "_1.fq.gz"
3 r_suffix: "_2.fq.gz"
4 env_name: "16S_rRNA_amplicon"
5 f_primer: "GTGCCAGCMGCCGCGG"
6 r_primer: "CCGTCAATTCMTTTRAGT"
7 phred_offset: 33 # 只有很早期是数据是64，现在都是33
8 minFoldParentOverAbundance: 1.5 # 嵌合体父母本丰度倍数
9 max_error_num: 4 # for cutadapt
10 error_rate_learning_nbases: 500000000 # 错误率学习的采样数量，会影响降噪准确度，1e8 大概能采样70000reads X 4样品
11 is_raw_data: false # 是否是raw data
12 taxon_mito_chloro:
13     "taxonomy_databases/silva_nr99_v138.1_train_Mitochondria_Chloroplast_set.fa"
14 db_taxon_file: "taxonomy_database/silva_nr99_v138.1_train_set.fa.gz"
15 db_spec_file: "taxonomy_database/silva_species_assignment_v138.1.fa.gz"
16 puppeteer_js: "puppeteer/app.js"
17 ez_account:
18     username: "xxxxx@xxxx"
19     password: "xxxxxxxxx"
```

5.2 配置VS Code

软件下载

- VS Code 官方网站: [链接](#)

连接远程服务器



安装插件

- SVG查看器

EXTENSIONS: MARKETPLACE

svg

**SVG**
SVG Coding, Minify, Pretty, Preview All-In-One
jock

992K ★ 4.5

Install

**Svg Preview**
Preview for Svg files
Simon Siefke

535K ★ 5

Install

**SVG Previewer**
Show SVG preview to the side panel
Vitalii Mazurenko

132K ★ 4.5

Install

**SVG Editor**
Visual and literal SVG editor for VSCode.
henoc

108K ★ 3.5

Install

**SVG Snippets**
A basic SVG snippet generator.
sidthesloth

32K ★ 5

Install

**SVG Gallery**
View SVG images in gallery
Yongjian Huang

33K ★ 5

Install

**SVG Icons**
Snippets for popular SVG icons, including Octicon...
idleberg

18K ★ 5

Install

**svg-fold**
Fold SVG tag
misbah ansori

1K ★ 5

Install

**SVG dev**
SVG visual editor
kontrail

9K ★ 4.5

Install

Extension: Svg Preview ×

**Svg Preview** v2.8.3
Simon Siefke | 535,276 | ★★★★★
Preview for Svg files
Install ⚙️

DETAILS FEATURE CONTRIBUTIONS CHANGELOG

build failing renovate enabled

Svg Preview for VSCode

Features

- Live editing of svg files and svg's inside files
- Panning and zooming of the preview (up to 32767%)

Commands

Command	Keybinding
Svg Preview: Open	ctrl+alt+⬆️
Preview to the Side	

Categories

Other

Extension Resources

[Marketplace](#)
[Repository](#)
[License](#)
[Simon Siefke](#)

More Info

Published 2019-4-11, 00:37:24
Last released 2020-2-14, 22:01:19
Identifier simonsiefke.svg-preview

SSH: 172.16.163.78 0 0 0

- 表格查看器

EXTENSIONS: MARKETPLACE

tsv

**VSCode TSV**
Automatically set the tabsize to align...
chitenb

20K ★ 4.5

Install

**Rainbow CSV**
Highlight CSV and TSV files, Run S...
mechatroner

5.4M ★ 5

Install

**Markdown table from ...**
Convert a CSV or TSV file to a Mark...
jojoco

8K ★ 4.5

Install

**TSV Paste JSON**
Convert the TSV text in the clipboar...
cmiz

912 ★ 5

Install

**Excel Viewer**
Edit Excel spreadsheets and CSV fil...
GrapeCity

4.3M ★ 3.5

Install

**Edit csv**
extension to edit csv files with a tab...
janisdd

1.2M ★ 5

Install

**ts-vue**
Add Typescript Vue Snippets
Marco Mantovani

316

Install

**CSV to Table**
Convert a CSV/TSV/PSV file to an A...
Andrew Armstrong

391K ★ 5

Install

**json2ts-vsc**
convert json to ts type
cylly

418 ★ 5

Install

**SandDance for VSC...**
Visually explore, understand, and pr...
Microsoft DevLabs

106K ★ 4.5

Install

**Data Preview**
Data Preview ?? extension for impor...
Random Fractals Inc.

520K ★ 4

Install

**CSV to Markdown Table**

19K ★ 5

Install

Extension: Excel Viewer ×

**Excel Viewer** v4.2.58
GrapeCity  grapecity.com | 4,337,458 | ★★★★★ (118)
Edit Excel spreadsheets and CSV files in Visual Studio Code and VS Code for the Web.
Install ⚙️

DETAILS FEATURE CONTRIBUTIONS CHANGELOG

Excel Viewer

Powered by [Wijmo](#), this extension provides custom editors and previews for CSV files and Excel spreadsheets in Visual Studio Code and [Visual Studio Code for the Web](#).

Version 4.2.58 fixes many CSV editing issues that occurred in files containing multiline cells.

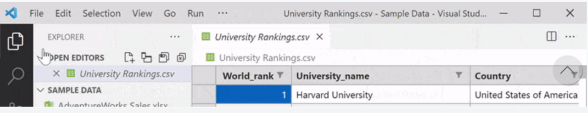
Version 4.2 now supports first-class **custom editors** that implement operations such as save, undo, redo, and hot exit. For Excel files, this is the default, and clicking the name of an Excel file in explorer view opens the custom editor directly. For CSV files, this is optional, and executing the **Open With** command on the context menu prompts for the built-in or custom editor to be opened. The **Open Preview** command is still supported for both file types.

Version 4.2 also supports **Visual Studio Code for the Web**. To get started, visit <https://vscode.dev> in your browser.

This extension requires Visual Studio Code 1.63.0 or greater.

CSV Usage

For files with a .csv, .tsv, or .tab extension, use the explorer context menu or editor title menu to invoke the **Open Preview** command. The contents of the file will be displayed in a **FlexGrid** control, which supports sorting and filtering via its column headers. You can also use the **Open With** command on the explorer context menu to open a custom editor, as shown here:



World_rank	University_name	Country
1	Harvard University	United States of America

Categories: Other

Extension Resources: [Marketplace](#), [Issues](#), [Repository](#), [License](#), [GrapeCity](#)

More Info: Published 2016-06-02, 05:18:06; Last released 2023-08-15, 23:08:36; Identifier grapecity.gc-excelviewer

0 0 0

5.3 运行snakemake

- 激活环境

```
1 mamba activate snakemake
```

- 流程图

```
1 # Silva流程
2 snakemake -s workflow/16S_rRNA_amplicon_Silva.smk --dag | dot -Tsvg >
  workflow_Silva.svg
3 # EZ流程
4 snakemake -s workflow/16S_rRNA_amplicon_EZ.smk --dag | dot -Tsvg > workflow_EZ.svg
```

- 运行Silva流程

```
1 # 试运行
2 # -c 线程数量
3 # --use-conda 使用conda环境
4 snakemake -s workflow/16S_rRNA_amplicon_Silva.smk -np -c all --use-conda
5 # 真运行
6 snakemake -s workflow/16S_rRNA_amplicon_Silva.smk -p -c all --use-conda
```

- 运行EZ流程

优点:

- 1数据库更全，能注释到更多种水平的ASV
- 2各个分类水平的单词是最新的命名单词

```
1 # 试运行
2 snakemake -s workflow/16S_rRNA_amplicon_EZ.smk -np -c all --use-conda
3 # 真运行
4 snakemake -s workflow/16S_rRNA_amplicon_EZ.smk -p -c all --use-conda
```

5.4 结果文件说明

```
1 # benchmarks
2 |─ dada2_single_end_reads_benchmark.tsv # 计算资源使用情况
3 # dada2_single_end_output
4 |─ seqs
5 |   |─ ASV_seqs_nochim.fasta # 去掉嵌合体后的ASV序列
6 |   |─ ASV_seqs_no_chloro_mito.fasta # 去掉嵌合体、叶绿体和线粒体后的ASV序列
7 |   |─ ASV_seqs_with_chloro_mito.fasta # 叶绿体和线粒体的ASV序列
8 |   |─ filtered_reads [73 entries exceeds filelimit, not opening dir]
9 |─ tables
10 |   |─ ASV_taxon_no_chloro_mito.tsv # 去掉嵌合体、叶绿体和线粒体后的ASV物种分类
```

```

11 | └─ ASV_taxon_with_chloro_mito.tsv # 叶绿体和线粒体ASV物种分类
12 | └─ ASV_prevalence_totalAbundance.tsv # ASV流行度和总丰度表
13 | └─ seq_table_no_chloro_mito.biom # 用于导入qiime2的特征表
14 | └─ seq_table_no_chloro_mito.tsv # 去掉嵌合体、叶绿体和线粒体后的特征表
15 | └─ seq_table_prop_no_chloro_mito.tsv # 去掉嵌合体、叶绿体和线粒体后的相对丰度特征表，可
    用于后续的beta多样性分析
16 |   └─ seq_table_with_chloro_mito.tsv # 叶绿体和线粒体序列的特征表
17 # summary_output
18 | └─ beta_diversity # 不同降维和距离算法的beta多样性
19 |   | └─ ord_nmds_bray_stress.tsv
20 |   | └─ ord_nmds_bray.tsv
21 |   | └─ ord_nmds_unifrac_stress.tsv
22 |   | └─ ord_nmds_unifrac.tsv
23 |   | └─ ord_nmds_wunifrac_stress.tsv
24 |   | └─ ord_nmds_wunifrac.tsv
25 |   | └─ ord_pcoa_bray_explained_variance_ratio.tsv
26 |   | └─ ord_pcoa_bray.tsv
27 |   | └─ ord_pcoa_unifrac_explained_variance_ratio.tsv
28 |   | └─ ord_pcoa_unifrac.tsv
29 |   | └─ ord_pcoa_wunifrac_explained_variance_ratio.tsv
30 |   | └─ ord_pcoa_wunifrac.tsv
31 | └─ cutadapter
32 |   | └─ R1_length_frequency_distribution_after_cutting.tsv # R1 裁剪引物后的长度分布矩阵
33 |   | └─ R1_length_frequency_distribution_before_cutting.tsv # R1 裁剪引物前的长度分布矩阵
34 |   | └─ R2_length_frequency_distribution_after_cutting.tsv # R2 裁剪引物后的长度分布矩阵
35 |   | └─ R2_length_frequency_distribution_before_cutting.tsv # R2 裁剪引物前的长度分布矩阵
36 | └─ dada2
37 |   | └─ assign_taxonomy_stats.tsv # 每一步处理所剩序列数量表 + 最终去除叶绿体和线粒体的所剩数
    量
38 |   | └─ denoising_stats.tsv # 每一步处理所剩序列数量表
39 | └─ alpha_diversity
40 |   | └─ alpha_diversity.tsv # 多种alpha多样性指标值表
41 | └─ stacked_bar_plot # 不同分类水平绘制堆叠柱状图的数据
42 |   | └─ Class.tsv
43 |   | └─ Family.tsv
44 |   | └─ Genus.tsv
45 |   | └─ Kingdom.tsv
46 |   | └─ Order.tsv
47 |   | └─ Phylum.tsv
48 |   | └─ Species.tsv
49 # figure_output
50 | └─ alpha_diversity # 多种alpha多样性图
51 |   | └─ ACE.png
52 |   | └─ Chaol.png
53 |   | └─ InvSimpson.png
54 |   | └─ Observed.png
55 |   | └─ Shannon.png
56 |   | └─ Simpson.png
57 | └─ cutadapter # 引物裁剪前后长度分布热图
58 |   | └─ R1_length_frequency_distribution_after_cutting.png
59 |   | └─ R1_length_frequency_distribution_before_cutting.png
60 |   | └─ R2_length_frequency_distribution_after_cutting.png

```

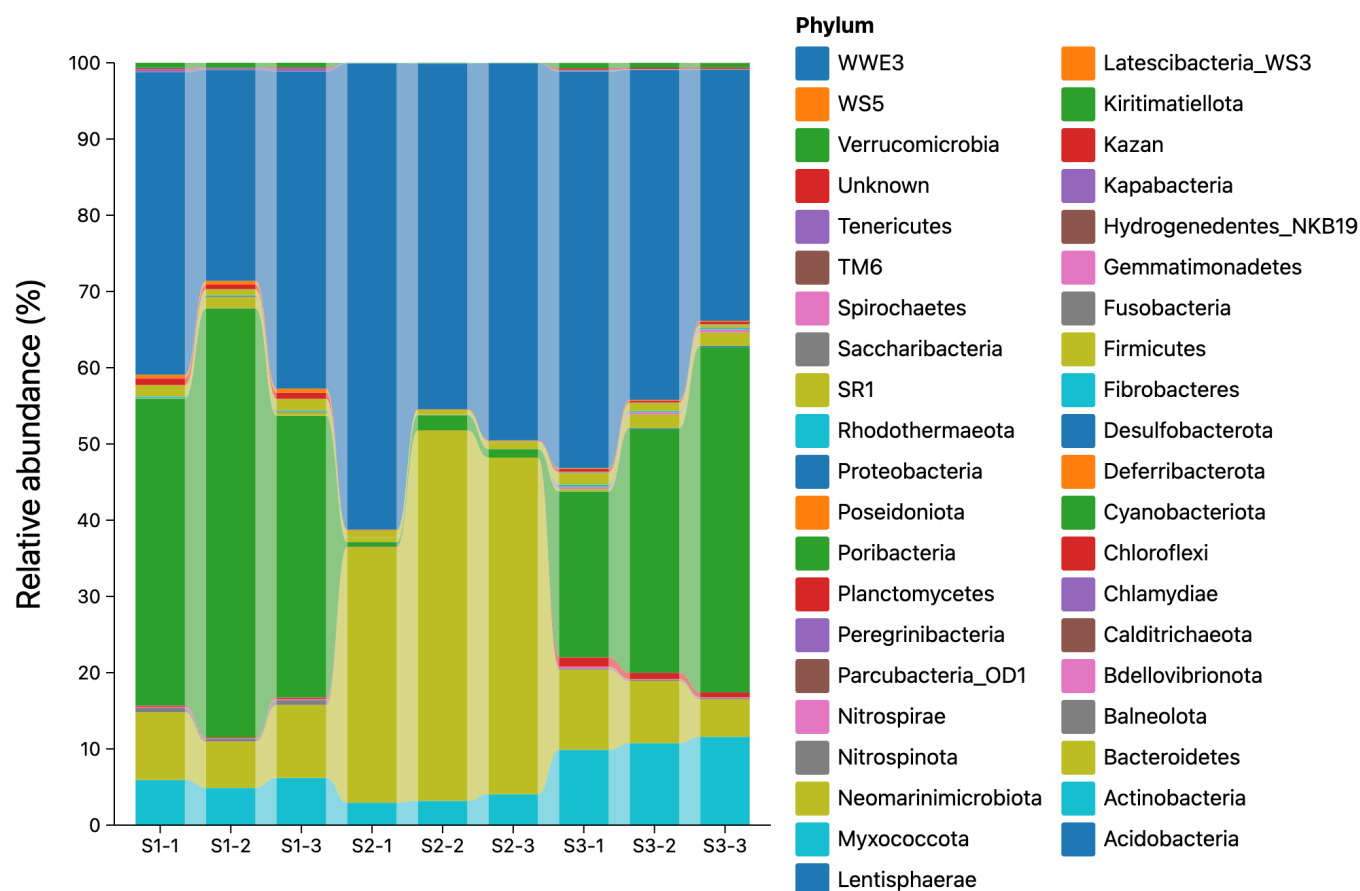
```

61 | └─ R2_length_frequency_distribution_before_cutting.png
62 | └─ prevalence # 不同分类水平的ASV流行度与总丰度的散点图
63 |   └─ Class.png
64 |   └─ Family.png
65 |   └─ Genus.png
66 |   └─ Kingdom.png
67 |   └─ Order.png
68 |   └─ Phylum.png

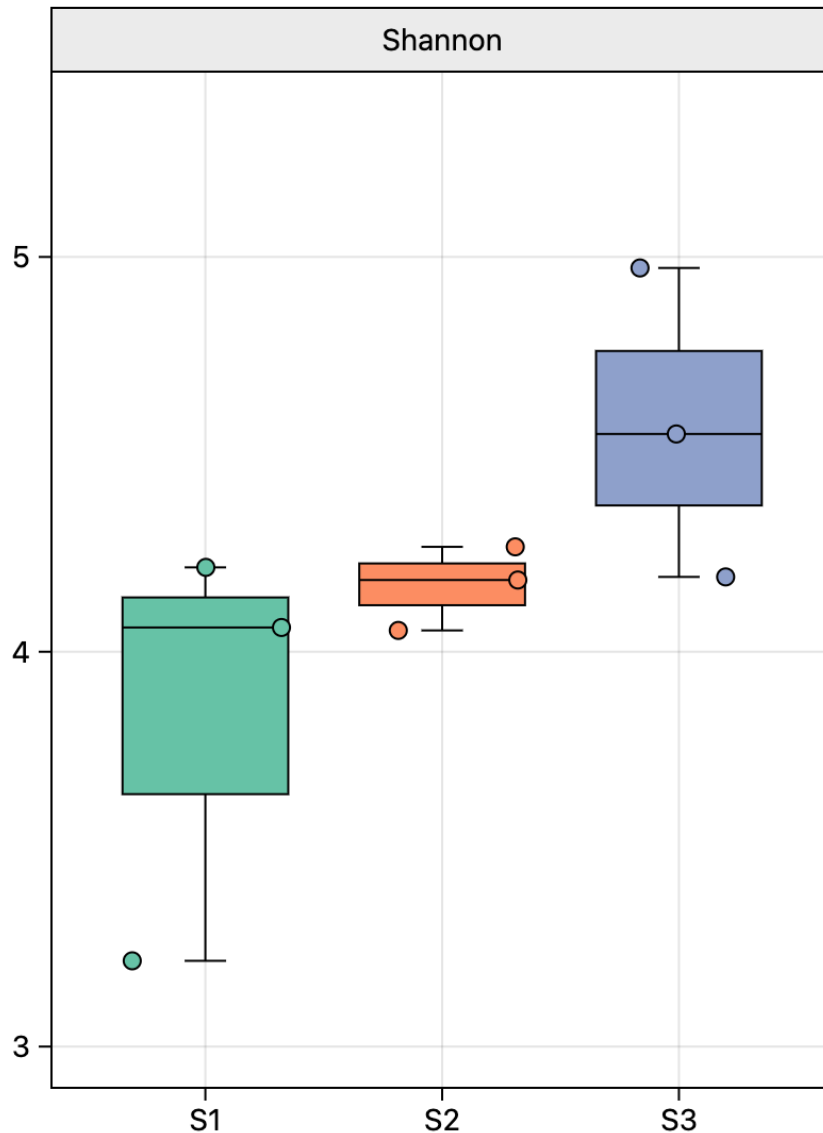
```

6. 下游数据分析与可视化

6.1 绘制物种组成丰度堆叠柱状图



6.2 alpha多样性可视化



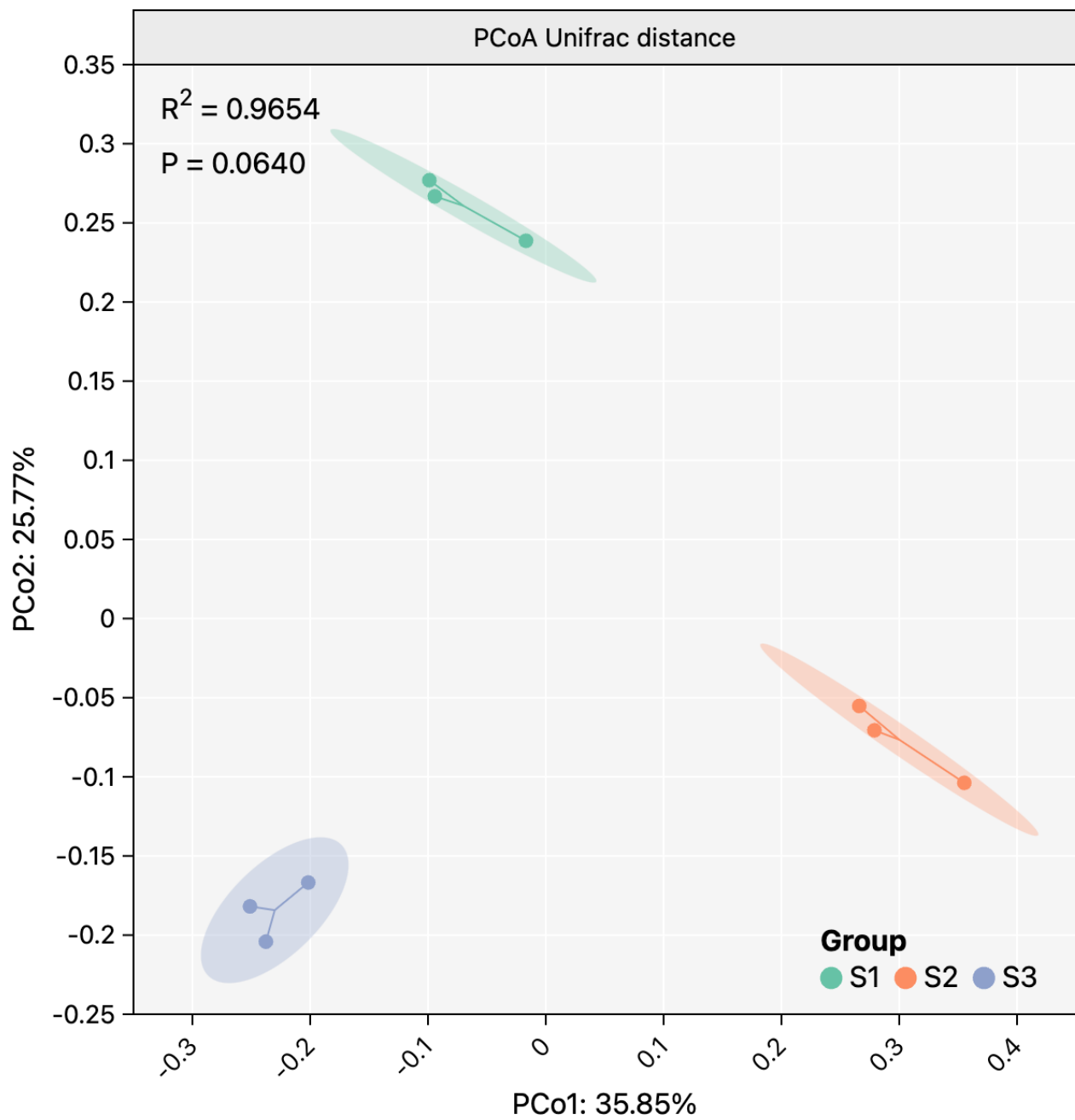
6.3 beta多样性可视化

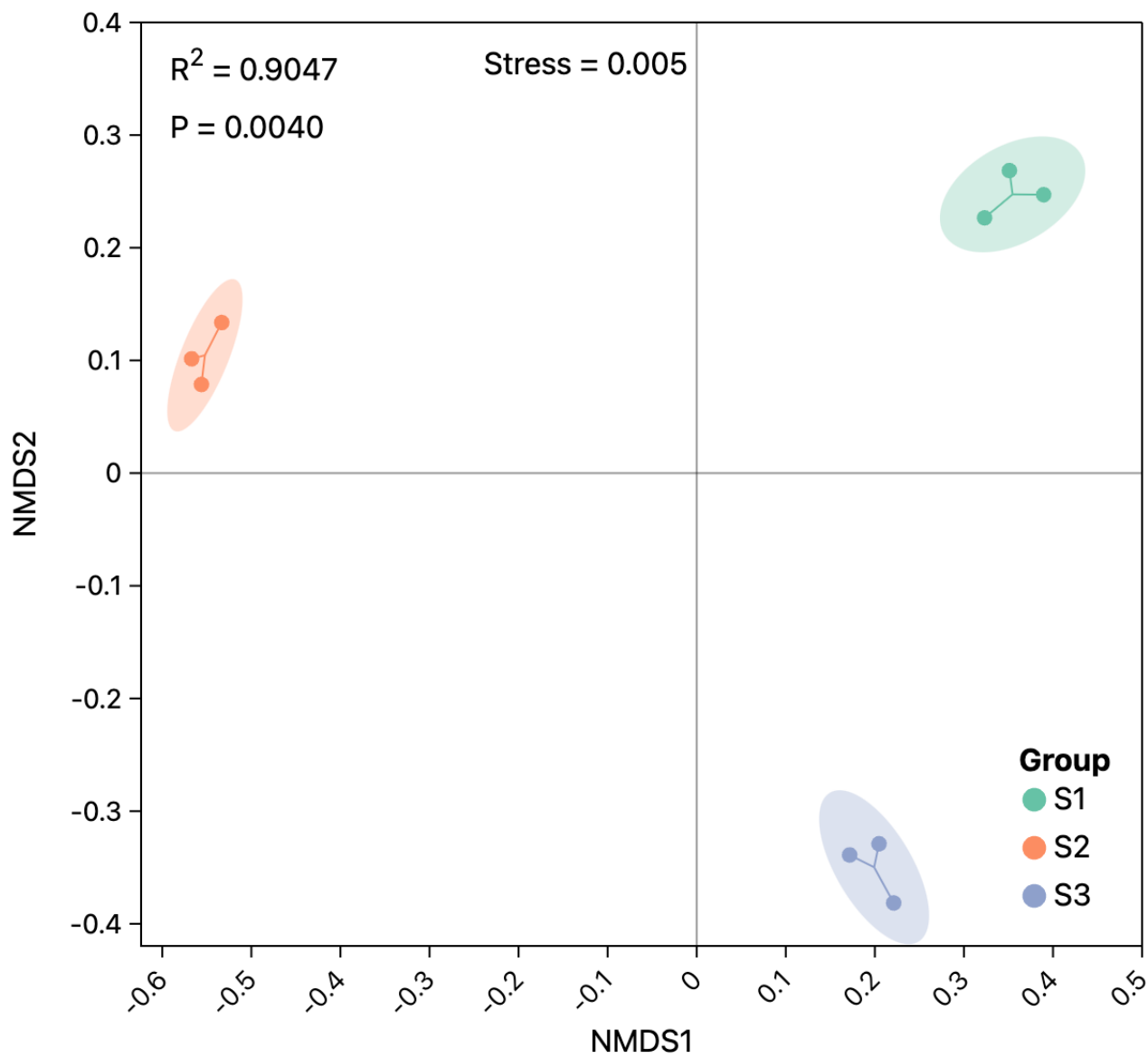
- **PERMANOVA** (Permutational Multivariate Analysis of Variance)

PERMANOVA是一种用于比较多组间差异的统计方法，其中 R^2 表示解释变量对响应变量的方差解释程度。在PERMANOVA中， R^2 值越接近1表示解释变量对样本间差异的解释程度越高，即样本间的差异可以较好地由解释变量来解释。这个值可以帮助我们理解解释变量对于样本结构的影响程度。通常情况下， R^2 值越高，说明解释变量对于样本间差异的解释能力越强。

- **NMDS** (Non-metric multidimensional scaling)

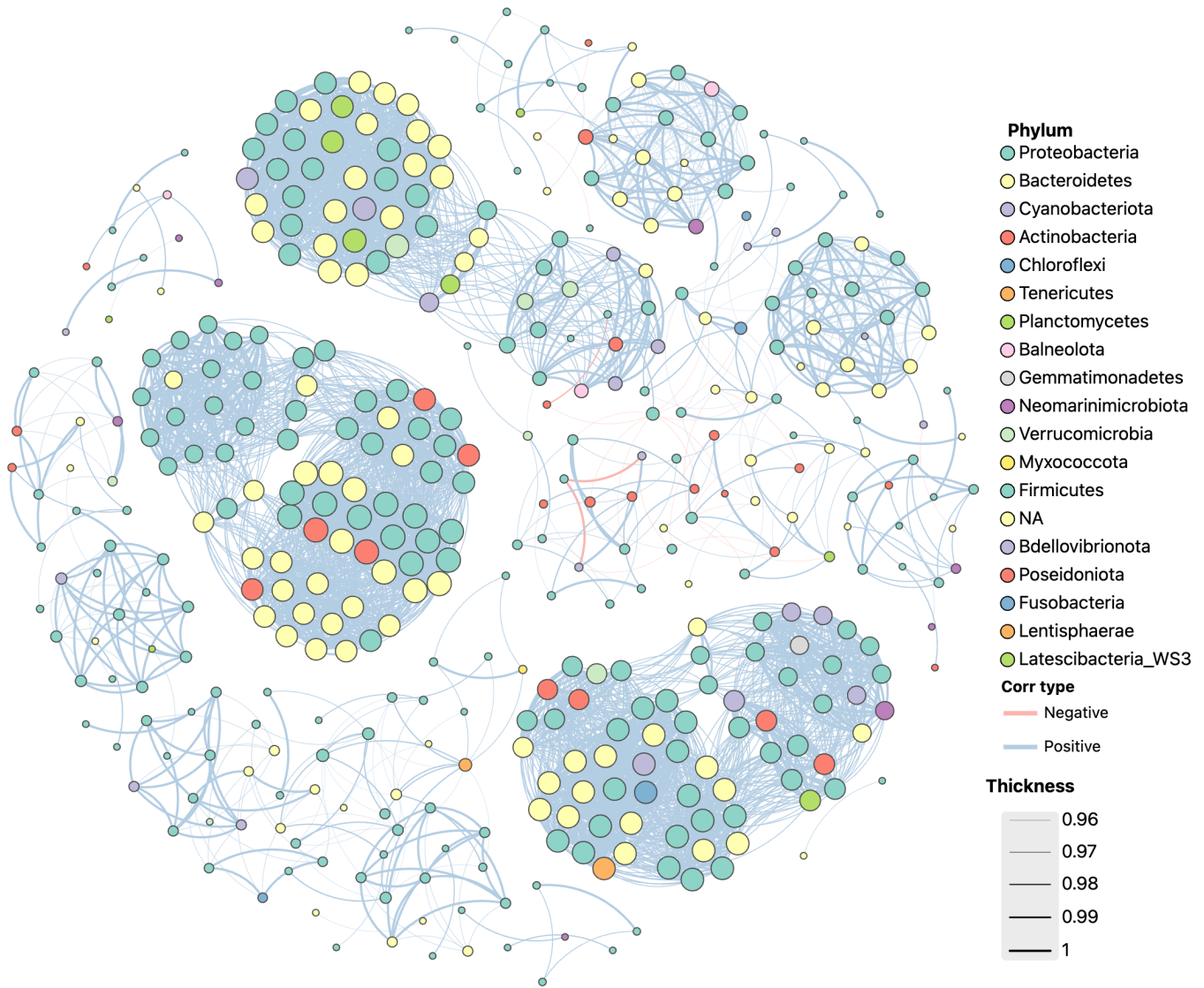
非度量多维尺度分析（NMDS）中的“stress”是一种用来衡量数据在低维空间中的拟合程度的指标。它表示数据在降维过程中失真的程度，即原始数据点之间的距离在降维后是否得到了很好的保留。Stress值越接近0，表示降维后的结果越能够准确地反映原始数据点之间的距离关系。通常来说，Stress值在0-0.2之间被认为是一个比较好的拟合度量。在NMDS分析中，通过调整数据点在低维空间中的位置，以最小化stress值来找到最佳的降维结果。





6.4 微生物共现网络分析

Microbial Co-occurrence Network



6.4.1 常见的微生物共现网络指标

- (1) **节点数 (Number of nodes)**: 网络中的节点个数
- (2) **边数 (Number of edges)**: 网络中所有连接微生物的边的数量
- (3) **节点度 (Node Degree)**: 节点的度是指与该节点直接相连的边的数量。在微生物共现网络中，节点通常代表微生物物种。节点度表示了一个微生物与多少个其他微生物有共现关系，高度连接的微生物可能在微生物群落中起着重要的作用。
- (4) **平均度 (Average Degree)** 是一个节点度 (Node Degree) 的统计性指标，用于描述微生物网络中每个节点（通常代表微生物物种）平均连接到多少个其他节点。更具体地说，平均度是通过计算网络中所有节点的度数（与其他节点相连的边的数量）并取平均值得到的。这个指标提供了关于微生物共现网络的整体连接密度的信息。平均度的值越高，表示微生物物种之间的相互关系更为密集；而平均度的值越低，则表示微生物物种之间的相互关系相对稀疏。
- (5) **网络密度 (Network Density)**: 网络密度衡量了网络中实际存在的边数量与可能存在的最大边数量之间的比例。高密度网络表示微生物之间的相互作用较为密集，低密度网络则反映了微生物之间的相互作用较为稀疏。
- (6) **平均最短路径长度 (Average Shortest Path Length)**: 这个指标衡量了网络中任意两个节点之间的平均最短路径长度。较短的平均路径长度可能表示微生物群落内的信息传递较为迅速。

(7) **网络直径 (Network diameter)**: 网络中任意两个节点之间最短路径长度的最大值, 即网络中最远的微生物间的距离

(8) **聚类系数 (Clustering coefficient)**: 用于描述网络中节点的聚集程度或者说节点之间的紧密程度。在微生物共现网络中, 聚类系数是用来衡量微生物物种之间共现关系的密度, 即它们是否倾向于与彼此共存形成子图 (聚类)。

(9) **网络中心性指标 (Centrality Measures)**: 包括介数中心性 (Betweenness Centrality)、接近度中心性 (Closeness Centrality) 和特征向量中心性 (Eigenvector Centrality) 等。这些指标用于识别在网络中具有重要影响力的微生物物种, 可能在信息传递或控制微生物群落结构方面发挥关键作用。

(10) **社区结构 (Community Structure)**: 社区检测算法用于将微生物网络分割成不同的社区或模块, 每个社区包含相互密切相关的微生物。社区结构有助于识别共现模式并揭示微生物群落的功能模块。

(11) **网络中心节点 (Hubs)**: 中心节点是在网络中具有极高节点度的微生物物种, 它们通常在微生物群落中具有关键的生态功能。

(12) **共现和互斥关系 (Co-occurrence and Co-exclusion)**: 共现关系指的是微生物在相同样本中同时出现的情况, 而互斥关系指的是微生物在相同样本中很少同时出现的情况。这些关系可以揭示微生物之间的相互作用模式, 如协作或竞争。

(13) **模块间联系 (Inter-module Connectivity)**: 这个指标用于了解不同社区之间的联系程度, 有助于理解微生物群落的整体结构。

(14) **模块度 (modularity)**: 是衡量网络社区结构的一个指标。模块度描述了网络中节点聚集成社区的程度, 数值范围通常在-1到1之间。当模块度接近1时, 表示网络中存在明显的社区结构; 而当模块度接近0时, 则意味着网络的社区结构不明显。

6.4.2 共现网络计算

- 安装R包 `Hmisc`、`igraph`

```
1 install.packages("Hmisc")
2 install.packages("igraph")
```

- 计算代码

```
1 # 参考贴子1: https://mp.weixin.qq.com/s/Ew7lUaPBQbuT84vYQOMLNA
2 # 参考贴子2: https://mp.weixin.qq.com/s/AAfrCaVobIfgbRtrM7E5dg
3 library(Hmisc)
4 library(igraph)
5
6 # ASV表: 行为样本, 列为ASV
7 ASV_table =
  read.table("/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/dada2_sing
    ngle_end_output/seq_table_no_chloro_mito.tsv", fill = T, check.names = F, row.names
    = 1, header=T, sep="\t")
8 # 物种分类表
9 tax_table =
  read.table("/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/dada2_sing
    ngle_end_output/ASV_taxon_no_chloro_mito.tsv", fill = T, check.names = F, row.names
    = 1, header=T, sep="\t")
```

```

10
11
12 # ASV过滤阈值参考文献: https://doi.org/10.1038/s41396-020-0621-7
13
14
15 # 取流行度比例大于20%的ASV
16 prev_prop_df = apply(X = ASV_table,
17                       MARGIN = 2,
18                       FUN = function(x){sum(x > 0)/dim(ASV_table)[1]})
19
20 total_relative_abundance = colSums(ASV_table)/sum(colSums(ASV_table))
21
22 is.above.prev = prev_prop_df > 0.2
23 table(is.above.prev)
24
25 # 取总相对丰度大于0.01%的ASV
26 is.above.abundance = total_relative_abundance > 0.0001
27 table(is.above.abundance)
28
29 table(is.above.prev & is.above.abundance)
30 target_ASV = colnames(ASV_table)[is.above.prev & is.above.abundance]
31
32 # 转换成相对丰度
33 ASV_table.prop = as.data.frame(t(apply(ASV_table, 1, function(ASV) ASV/sum(ASV))))
34
35 target_ASV_table.prop = ASV_table.prop[, target_ASV]
36
37 print(rowSums(target_ASV_table.prop))
38
39
40 # Calculate pairwise correlation (Spearman correlation)
41 correlation_matrix <- rcorr(as.matrix(target_ASV_table.prop), type="spearman")
42
43 # 相关性过滤阈值参考文献: https://doi.org/10.1186/s40168-023-01708-6
44
45 # Thresholding: Retain only significant correlations
46 correlation_threshold <- 0.7 # Adjust as needed
47 p_threshold <- 0.001 # Adjust as needed 0.05 有点太大
48
49
50
51 CorrDF <- function(cormat, pmat) {
52   ut <- upper.tri(cormat) # 由于相关性矩阵是对称的, 取上三角矩阵计算即可
53   data.frame(
54     source = rownames(cormat)[col(cormat)[ut]],
55     target = rownames(cormat)[row(cormat)[ut]],
56     corr = (cormat)[ut],
57     weight = abs((cormat)[ut]), # 相关性系数绝对值作为权重
58     p = pmat[ut],
59     p.adjust = p.adjust(pmat[ut], method="BH"), #Multiple testing correction using
60     Benjamini-Hochberg standard false discovery rate correction ("FDR-BH"), to minimize
61     false-positive signals before network construction

```

```

60     cor_type = ifelse((cormat)[ut]> 0, "Positive", "Negative")
61   )
62 }
63
64
65 cor_df <- CorrDF(correlation_matrix$r , correlation_matrix$p) # 整理ASV之间的连接关系
66
67 significant_cor_df <- cor_df[cor_df$weight >= correlation_threshold &
68   cor_df$p.adjust < p_threshold,] # 保留spearman相关性绝对值>0.7且q-value < 0.001的边
69
70 # 非加权无向
71 g = graph_from_data_frame(significant_cor_df, directed=F)
72
73 # plot_network(g, ps.prop_above_prev_abund, type = 'taxa', color = "Phylum")
74
75 significant_tax = tax_table[V(g)$name,]
76 V(g)$Kingdom = significant_tax$Kingdom #界
77 V(g)$Phylum = significant_tax$Phylum #门
78 V(g)$Class = significant_tax$Class #纲
79 V(g)$Order = significant_tax$Order #目
80 V(g)$Family = significant_tax$Family #科
81 V(g)$Genus = significant_tax$Genus #属
82 V(g)$Species = significant_tax$Species #种
83
84 # 创建节点列表
85 node_list = data.frame(node_id = names(V(g)),significant_tax)
86
87 # 创建边列表
88 edge_list = data.frame(edge_id = paste0("edge_", c(1:length(E(g)))),
89   significant_cor_df)
90
91 table(edge_list$cor_type)
92
93 # Network analysis
94
95 nodes_num = length(V(g)) #节点数
96 nodes_num
97
98 edges_num = length(E(g)) #边数
99 edges_num
100
101 positive_corr_num = sum(E(g)$corr>0) #正相关的数量
102 positive_corr_num
103
104 positive_corr_prop = positive_corr_num/edges_num #正相关的比例
105 positive_corr_prop
106
107 negative_corr_num = sum(E(g)$corr<0) #负相关的数量
108 negative_corr_num
109

```

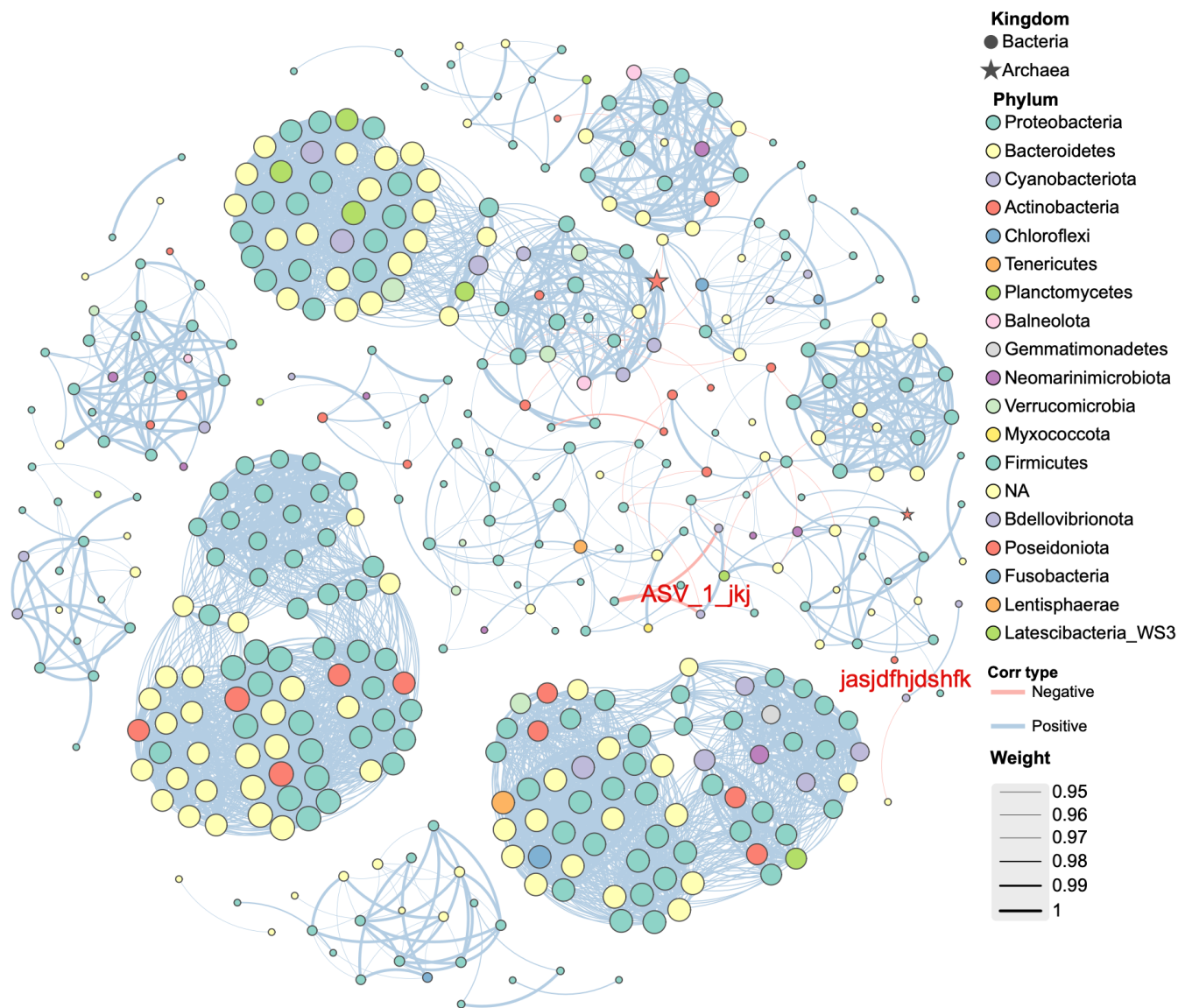
```

110 negative.corr_prop = negative.corr_num/edges_num #负相关的比例
111 negative.corr_prop
112
113 average_degree = mean(degree(g)) #平均度
114 average_degree
115
116 average_shortest_path_length = mean_distance(g, directed = FALSE) #平均最短路径长
    度
117 average_shortest_path_length
118
119 network_diameter = diameter(g, directed = FALSE) #网络直径
120 network_diameter
121
122 network_density = edge_density(g) #网络密度
123 network_density
124
125
126 modularity = modularity(g, membership(cluster_louvain(g)), directed = F) # 模块度
127 modularity
128
129 # 传递性: Transitivity measures the probability that the adjacent vertices of a
    vertex are connected. This is sometimes also called the clustering coefficient.
130 # 聚类系数
131 clustering_coefficient <- transitivity(g, type="global")
132 clustering_coefficient
133
134 network_parameter = data.frame(nodes_num,
135                                edges_num,
136                                positive.corr_num,
137                                positive.corr_prop,
138                                negative.corr_num,
139                                negative.corr_prop,
140                                average_degree,
141                                average_shortest_path_length,
142                                network_diameter,
143                                network_density,
144                                clustering_coefficient,
145                                modularity
146 )
147
148 output_dir =
    "/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/summary_output/netw
    ork"
149
150 write.table(node_list, file = file.path(output_dir, "network.node_list.tsv"), sep =
    "\t", quote = F, row.names = F)
151 write.table(edge_list, file = file.path(output_dir, "network.edge_list.tsv"), sep =
    "\t", quote = F, row.names = F)
152 write.table(network_parameter, file = file.path(output_dir,
    "network_parameter.tsv"), sep = "\t", quote = F, row.names = F)
153 write_graph(g, file.path(output_dir, 'network.graphml'), format = 'graphml') #后续在
    Gephi中可视化

```

6.4.3 共现网络可视化

- ChiPlot



- Gephi: [下载地址](#)

[b站教程](#)

缺点：没有图注

6.5 Mantel test环境因子与群落组成相关性分析

6.5.1 Mantel test计算

- 安装R包 `linkET`

```
1 install.packages("devtools")
2 devtools::install_github("Hy4m/linkET", force = TRUE)
```

- 计算代码

```
1 library(linkET)
2
3 # ASV表: 行为样本, 列为ASV
4 ASV_table =
5   read.table("/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/summary_o
6   utput/mantel_test/data/seq_table_no_chloro_mito.tsv", fill = T, check.names = F,
7   row.names = 1, header=T, sep="\t")
8
9 # 物种分类表
10 tax_table =
11   read.table("/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/summary_o
12   utput/mantel_test/data/ASV_taxon_no_chloro_mito.tsv", fill = T, check.names = F,
13   row.names = 1, header=T, sep="\t")
14
15 # 理化因子表
16 chem_table =
17   read.table("/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/summary_o
18   utput/mantel_test/data/varechem.tsv", fill = T, check.names = F, row.names = 1,
19   header=T, sep="\t")
20
21 # 转换成相对丰度
22 ASV_table.prop = as.data.frame(t(apply(ASV_table, 1, function(ASV) ASV/sum(ASV))))
23
24 # 计算理化因子相关性
25 correlation_matrix = correlate(chem_table, engine = "Hmisc", method = "pearson")
26
27 corr_p = correlation_matrix$p
28
29 # p值校正 Multiple testing correction using Benjamini-Hochberg standard false
30 discovery rate correction ("FDR-BH"), to minimize false-positive signals
31 correlation_matrix = adjust_pvalue(correlation_matrix, method = "BH")
32
33 corr_r = correlation_matrix$r
34 corr_p_adjust = correlation_matrix$p
35
36 # 目标计算分类水平
37 tax_level = "Phylum"
38 target_tax_name = c("Actinobacteria", "Bacteroidetes", "Cyanobacteriota",
39 "Proteobacteria")
40
```

```

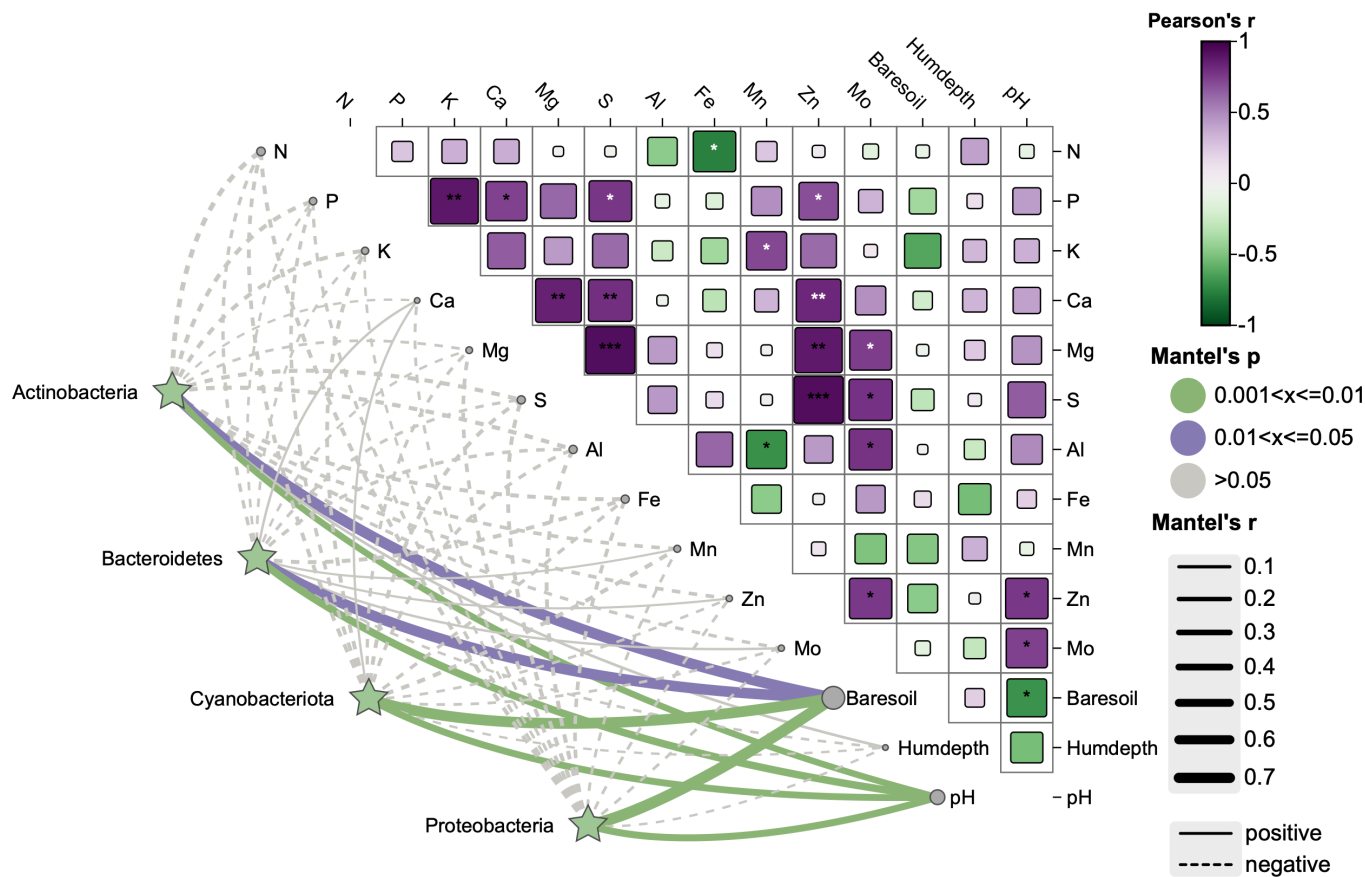
30 output_dir =
  "/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/summary_output/mante
  l_test/res"
31
32 for(target_tax in target_tax_name){
33   # 根据目标分类抽取相应的ASV子表
34   target_ASV = rownames(tax_table)[tax_table[[tax_level]] %in% target_tax]
35   sub_ASV_table.prop = ASV_table.prop[, target_ASV]
36
37   # 计算Mantel test
38   mantel_res <- mantel_test(sub_ASV_table.prop, chem_table, mantel_fun = "mantel",
spec_dist = "bray", env_dist = "euclidean")
39   colnames(mantel_res)[1] <- target_tax
40
41   write.table(mantel_res, file = file.path(output_dir, paste0(target_tax, ".tsv")),
sep = "\t", quote = F, row.names = F)
42
43
44 }
45
46 # 计算全部物种的Mantel test
47 mantel_res <- mantel_test(ASV_table.prop, chem_table, mantel_fun = "mantel",
spec_dist = "bray", env_dist = "euclidean")
48
49 colnames(mantel_res)[1] <- "All species"
50
51 write.table(mantel_res, file = file.path(output_dir, "All species.tsv"), sep = "\t",
quote = F, row.names = F)
52
53
54 # 导出相关性表格
55 # 修改第一列名称为rValue, 方便chiplot识别
56 write.table(cbind(rValue=rownames(corr_r),corr_r), file = file.path(output_dir,
"rValue.tsv"), sep = "\t", quote = F, row.names = F)
57
58 # 修改第一列名称为pValue, 方便chiplot识别
59 write.table(cbind(pValue=rownames(corr_p),corr_p), file = file.path(output_dir,
"pValue.tsv"), sep = "\t", quote = F, row.names = F)
60 write.table(cbind(pValue=rownames(corr_p_adjust), corr_p_adjust), file =
file.path(output_dir, "pValue_adjust.tsv"), sep = "\t", quote = F, row.names = F)
61
62

```

Mantel test相关性r的生物学意义

- 可以理解为两个矩阵内距离变化的相关性，传统的是两列数值大小变化的相关性
- 就微生物生态方向的研究，一个环境因素对某种生物类群起作用应该是环境因素差异增大时，生物类群差异变大。不存在环境因素差异增大，让生物类群变得更一致的情况。即使结果是负值，也只是基于计算得出的
- 生态领域，一般计算环境距离和群落距离负的情况不多，即使负也是接近0. 所以一般论文那样处理原则性问题没有， $r < 0.25$ 和 $|r| < 0.25$ 一样的。因为负的都比较小，Mantel基本不会是负，如果是负，也不会显著

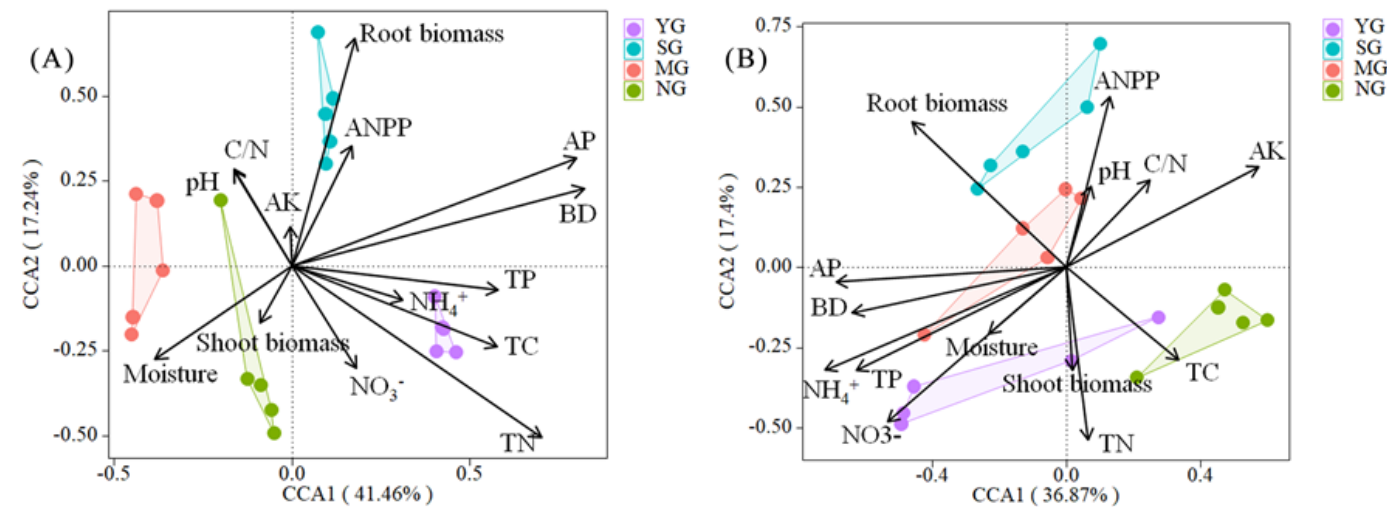
6.5.2 Mantel test可视化



6.6 CCA、RDA环境因子与物种分布相关性分析

6.6.1 CCA、RDA图解读

[参考推文](#)



CCA、RDA是一种用于研究生态学数据的多元统计技术。它结合了对物种丰富度和环境因素的分析，帮助研究人员理解生态系统中不同物种与环境条件之间的关系。可以揭示物种分布与环境变量之间的相关性，帮助解释生态系统中的生物多样性模式。该方法常用于揭示环境因素对生物多样性空间分布的影响，以及预测不同环境条件下物种的分布。通过对数据进行排序和约简，有助于识别主要驱动生态系统变化的因素。

解读

- RDA和CCA的结果图中使用点代表不同的样本，从原点发出的箭头代表不同的环境因子。
- 箭头的长度代表该环境因子对群落变化影响的强度，箭头的长度越长，表示环境因子的影响越大。
- 箭头与坐标轴的夹角代表该环境因子与坐标轴的相关性，夹角越小，代表相关性越高。
- 样本点到环境因子箭头极其延长线的垂直距离表示环境因子对样本的影响强度，样本点与箭头距离越近，该环境因子对样本的作用越强。
- 样本位于箭头同方向，表示环境因子与样本物种分布的变化正相关，样本位于箭头的反方向，表示环境因子与样本物种分布的变化负相关。

算法选择

我们进行RDA或CCA之前先进行DCA分析(Detrended correspondence analysis),根据结果中的Lengths of gradient的数值来进行判断，以下是一条不成文的规则：

结果会给出4个Lengths of gradient的数值，如果其中最大的数值大于4，则应选择CCA，如果最大的数值小于3，则选择RDA，如果最大的数值在3-4之间，则两种分析方法都可以。

其实，这只是经验定律，在实际的分析过程中，我们同时进行两个分析，然后比较哪个分析的结果更符合我们的预期，我们就用哪个，毕竟统计是为我们服务的，不能反被统计绑架。

6.6.2 CCA、RDA计算

- 安装R包 `vegan`

```
1 install.packages("vegan")
```

- 计算代码

```
1 #Canonical Correspondence Analysis (CCA)
2 library(vegan)
3
4
5
6 # ASV表：行为样本，列为ASV
7 ASV_table =
8   read.table("/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/summary_o
9   utput/CCA/data/seq_table_no_chloro_mito.tsv", fill = T, check.names = F, row.names =
10  1, header=T, sep="\t")
11
12
13 # 理化因子表
14 chem_table =
15   read.table("/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/summary_o
16   utput/CCA/data/varechem.tsv", fill = T, check.names = F, row.names = 1, header=T,
17   sep="\t")
18
19 # 标准化成相对丰度
20 # 标准化参考帖子: https://r.qcbs.ca/workshop09/book-en/transformations.html
```

```

15 ASV_table.std.total = decostand(ASV_table, MARGIN = 1, method = "total")
16
17 # 标准化成相对丰度的平方根
18 #ASV_table.std.hellinger = decostand(ASV_table, MARGIN = 1, method = "hellinger")
19
20
21 #先进行DCA分析
22 dca <- decorana(veg = ASV_table.std.total)
23
24 # DCA分析
25 GL = as.data.frame(apply(dca$rproj, 2, max))
26
27 colnames(GL) <- c("Gradient length")
28
29 #You do not give a reproducible case, but if you really have 31 sampling units
  (observations) and 55 predictors, you are overfitting your data. You cannot have more
  predictors than observations – or you can, but 31 random variables will completely
  explain 31 observations. Probably the problem is the same with your "only 10 first
  variables": these were enough to predict exactly your observations and the later ones
  were dumped (we call that "aliasing"). As a summary: the number of predictor
  variables cannot be higher than the number of observations. You need either more data
  or you need to reduce the number of your predictors.
30 # 环境因子数不能少于样本数，要不只会取前（样本数-1）个环境因子去计算
31
32 # 计算RDA
33 # rda_result <- rda(ASV_table.std.total, chem_table)
34
35 # 计算CCA
36 cca_result <- cca(ASV_table.std.total, chem_table)
37
38 plot(cca_result)
39
40 cca_score <- scores(cca_result)
41
42 # 坐标轴解释比例
43 cca_summary = summary(cca_result)
44 explained_proportion1 <- round(cca_summary$cont$importance["Proportion Explained", 1]
  * 100,2)
45 explained_proportion2 <- round(cca_summary$cont$importance["Proportion Explained", 2]
  * 100,2)
46
47 output_dir =
  "/Users/xjm/Documents/sampling/resampling_data_analysis/16S_test/summary_output/CCA/r
  es"
48
49 sites_table = cca_score$sites
50
51 write.table(cbind(sample=rownames(sites_table), sites_table), file =
  file.path(output_dir, "cca_sample.tsv"), sep = "\t", quote = F, row.names = F)
52
53 envs_table = cca_score$biplot
54

```

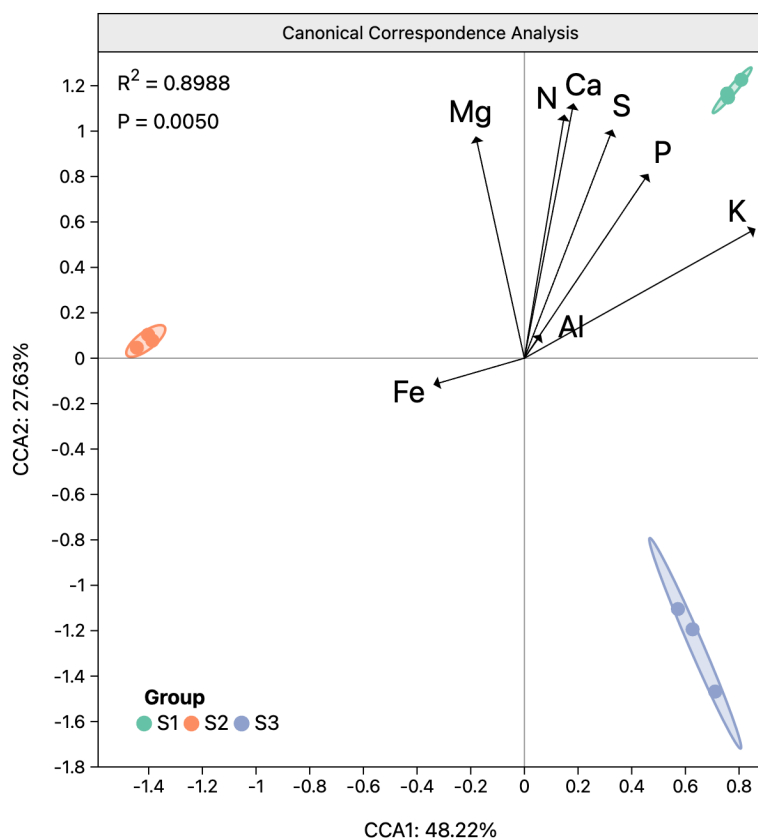
```

55 write.table(cbind(env=rownames(envs_table), envs_table), file = file.path(output_dir,
56 "cca_env.tsv"), sep = "\t", quote = F, row.names = F)
57 explained_proportion_table = data.frame(type=c("explained variance ratio (%)"),
58 CCA1=explained_proportion1, CCA2=explained_proportion2)
59 write.table(explained_proportion_table, file = file.path(output_dir,
60 "explained_proportion.tsv"), sep = "\t", quote = F, row.names = F)
61
62

```

6.6.3 CCA、RDA可视化

[参考推文](#)



7. 导入qiime2

- 创建目录

```

1 mkdir -p qiime2_output/qza
2 mkdir -p qiime2_output/qzv

```

- 导入代表序列

```

1  # 导入
2  qiime tools import \
3    --input-path dada2_single_end_output/seqs/ASV_seqs_no_chloro_mito.fasta \
4    --output-path qiime2_output/qza/rep-seqs.qza \
5    --type 'FeatureData[Sequence]'
6
7  # 可视化
8  qiime feature-table tabulate-seqs \
9    --i-data qiime2_output/qza/rep-seqs.qza \
10   --o-visualization qiime2_output/qzv/rep-seqs.qzv

```

- 导入特征表

```

1  # 导入
2  qiime tools import \
3    --input-path dada2_single_end_output/tables/seq_table_no_chloro_mito.biom \
4    --type 'FeatureTable[Frequency]' \
5    --input-format BIOMV100Format \
6    --output-path qiime2_output/qza/feature-table.qza
7  # 可视化
8  qiime feature-table summarize \
9    --i-table qiime2_output/qza/feature-table.qza \
10   --o-visualization qiime2_output/qzv/feature-table.qzv

```